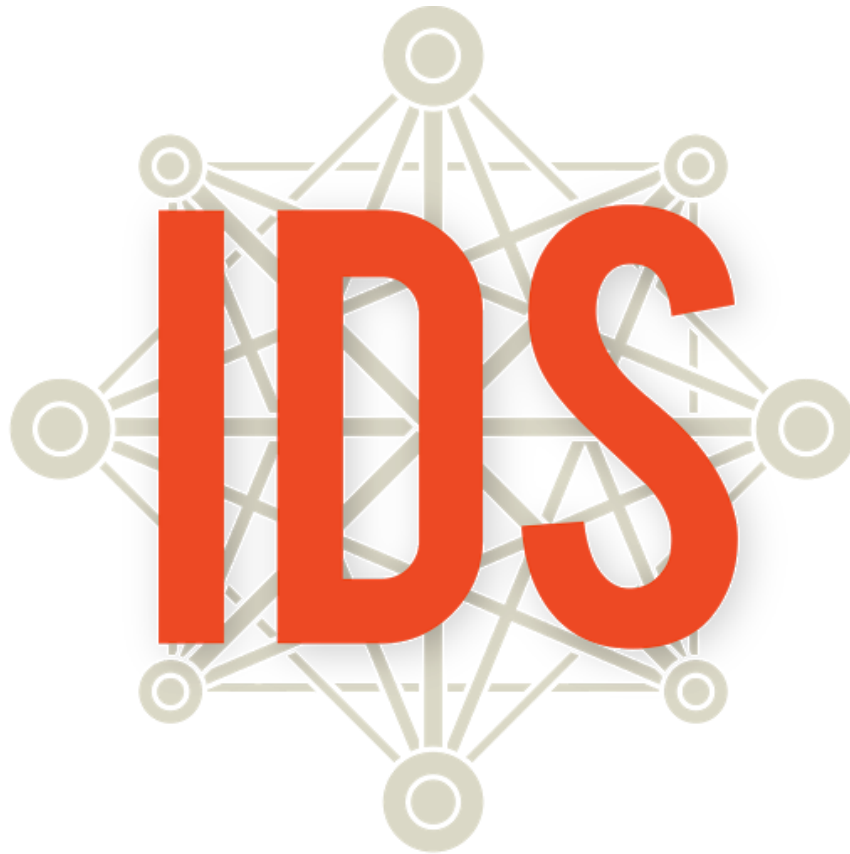




Introduction to Data Science

Robert Gould

Suyen Machado

Terri Anna Johnson

James Molyneux

Sponsors & Supporters

This curriculum was created under the auspices of the National Science Foundation, Mathematics and Science Partnership grant, “MOBILIZE: Mobilizing for Innovative Computer Science Teaching and Learning.” Lead Principal Investigator: Robert Gould (UCLA, Statistics).



Contributing Authors

LAUSD: Monica Casillas, Heidi Estevez, and Carole Sailer

UCLA: Amelia McNamara and Linda Zanontian

Acknowledgments and Special Thanks

Co-Principal Investigators: Deborah Estrin (UCLA, CENS), Joanna Goode (University of Oregon), Mark Hansen (UCLA, Statistics), Jane Margolis (UCLA, Center X), Thomas Philip (UCLA, Center X), Jody Prisela (UCLA, GSEIS), Derrick Chau (LAUSD), Gerardo Loera (LAUSD) and Todd Ullah (LAUSD); Mobilize Project Director: LeeAnn Trusela

LAUSD IDS Pilot Teachers

Robert Montgomery	Carole Sailer	Joy Lee	Monica Casillas
Roberta Ross	Velia Valle	Jose Guzman	Pamela Amaya
Arlene Pascua	Chris Marangopoulos		

This material is based upon work supported by the National Science Foundation under Grant Number 0962919.

Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

For additional information related to IDS visit: www.thinkdataed.org



Mobilize, an innovative partnership between the University of California, Los Angeles (UCLA) and the Los Angeles Unified School District (LAUSD), was funded in 2010 by the National Science Foundation to develop barrier-breaking curriculum in science, mathematics, and computer science to teach students to think creatively, constructively, and critically about the role of data in science and in everyday life. The Mobilize curricula center around Participatory Sensing campaigns, through which students use their mobile devices to collect and share data about their community and their lives and analyze these data to gain a greater understanding about their world. Mobilize broke barriers by teaching students to apply concepts and practices from computer science and statistics to learning science and mathematics, and it was uniquely dynamic in that each Mobilize class collects its own data, and each class has the opportunity to make unique discoveries. Across all Mobilize curricula, mobile devices are used not as gimmicks to capture students' attention, but as legitimate tools that bring scientific enquiry into their everyday lives. From 2011 to 2014, the Mobilize grant funded summer institutes for LAUSD high school mathematics, science, and computer science teachers to learn to use the Participatory Sensing (PS) methods, tools, and materials to deepen their knowledge of computer science (CS) concepts and to support student CS, math, and science learning.

First implemented in 2014 under the auspices of the Mobilize grant, Introduction to Data Science (IDS) began as a pilot program with 10 LAUSD mathematics teachers. Since, the IDS curriculum has been offered across the United States and internationally. In addition to addressing the Common Core State Standards (CCSS) for High School Statistics and Probability, IDS leads students to:

- understand how data are used by professionals to address real-world problems;
- understand that data are used in all facets of modern life;
- understand how data support science to identify and tackle real-world problems in our communities;
- analyze statistical graphics to identify patterns in data and to connect these patterns back to the real world;
- understand that by treating photos, words, numbers, and sounds as data, we can gain insight into the real world;
- learn to analyze data, including: posing questions that can be answered by considering relations among variables in a dataset, using collected data to generate hypotheses for future data collection, critically evaluating shortcomings and strengths in the data and the data collection process, and informally evaluating hypotheses using data at hand.

Table of Contents

Content	Page
Overview & Philosophy	8
Scope and Sequence	14
Unit 1	18
Daily Overview	19
Essential Concepts	20
Section 1: Data Are All Around	22
Lesson 1: Data Trails	Defining data, consumer privacy 24
Lesson 2: Stick Figures	Organizing & collecting data 26
Lesson 3: Data Structures	Organizing data, rows & columns, variables 28
Lesson 4: The Data Cycle	Data cycle, statistical questions 30
Lesson 5: So Many Questions	Statistical questions, variability 35
Lesson 6: What Do I Eat?	Data cycle, collecting data [Food Habits] 39
Lesson 7: Setting the Stage	Participatory Sensing [Food Habits] 42
Section 2: Visualizing Data	47
Lesson 8: Tangible Plots	Dotplots, minimum/maximum, frequency [Food Habits] 48
Lesson 9: What is Typical?	Typical value, center [Food Habits] 52
Lesson 10: Making Histograms	Histograms, bin widths [Food Habits] 55
Lesson 11: What Shape Are You In?	Shape, center, spread [Food Habits] 58
Lesson 12: Exploring Food Habits	Single & multi-variable plots [Food Habits] 60
Lesson 13: RStudio Basics	Intro to RStudio [Food Habits] 62
Lab 1A: Data, Code & RStudio	RStudio basics [Food Habits] 65
Lab 1B: Get the Picture?	Variable types, bar graphs, histograms [Food Habits] 69
Lab 1C: Export, Upload, Import	Importing data [Food Habits] 72
Lesson 14: Variables, Variables, Variables	Multi-variable plots 77
Lab 1D: Zooming Through Data	Subsetting 81
Lab 1E: What's the Relationship?	Multi-variable plots 85
Practicum: The Data Cycle & My Food Habits	Data cycle, variability [Food Habits] 88
Section 3: Would You Look at the Time	90
Lesson 15: Americans' Time on Task	Evaluating Claims [Time Use] 91
Lab 1F: A Diamond in the Rough	Cleaning names, categories, and strings [Time Use] 96
Lesson 16: Categorical Associations	Joint relative frequencies in 2-way tables [Time Use] 101
Lesson 17: Interpreting Two-Way Tables	Marginal & conditional relative frequencies [Time Use] 103
Lab 1G: What's the FREQUENCY?	2-way tables, tally [Time Use] 108
Practicum: Teen Depression	Statistical question, interpreting plots [Time Use] 111
Lab 1H: Our Time	Data cycle, synthesis 113
End of Unit 1 Project: Evaluating Claims from the Media	Data cycle 114

Table of Contents (continued)

Content		Page
Unit 2	Topics [Campaign]	115
Daily Overview		116
Essential Concepts		117
Section 1: What is Your True Color?		119
Lesson 1: What is Your True Color?	Subsets, relative frequency [Personality Color]	121
Lesson 2: What Does Mean Mean?	Measures of center – mean [Personality Color]	124
Lesson 3: Median in the Middle	Measures of center – median [Personality Color]	128
Lesson 4: How Far is it from Typical?	Measures of spread – MAD [Personality Color]	131
Lab 2A: All About Distributions	Measures of center & spread [Personality Color]	135
Lesson 5: Human Boxplots	Boxplots, IQR	137
Lesson 6: Face Off	Comparing distributions	140
Lesson 7: Plot Match	Comparing distributions	143
Lab 2B: Oh the Summaries...	Numerical summaries, custom functions [Personality Color]	145
Practicum: The Summaries	Data cycle, comparing distributions [Food Habits/Time Use]	148
Section 2: How Likely Is It?		150
Lesson 8: How Likely Is It?	Probability, simulations	151
Lesson 9: Dice Detective	Simulations to detect unfairness	154
Lesson 10: Marbles Marbles	Probability, with replacement	158
Lab 2C: Which Song Plays Next?	Probability of simple events, do loops, set.seed()	160
Lesson 11: This AND/OR That	Compound probabilities	163
Lab 2D: Queue It Up!	Probability with/without replacement, sample()	167
Practicum: Win Win Win	Probability estimation through repeated simulations	170
Section 3: Are You Stressing or Chilling?		171
Lesson 12: Don't Take My Stress Away	Introduction to campaign [Stress/Chill]	173
Lesson 13: The Horror Movie Shuffle	Chance differences – categorical [Stress/Chill]	177
Lab 2E: The Horror Movie Shuffle	Inference for categorical variables [Stress/Chill]	181
Lesson 14: The Titanic Shuffle	Chance differences – numerical [Stress/Chill]	184
Lab 2F: The Titanic Shuffle	Inference for numerical variables [Stress/Chill]	188
Lesson 15: Tangible Data Merging	Merging datasets [Stress/Chill]	190
Lab 2G: Getting It Together	Stacking vs. joining datasets [Stress/Chill & Personality Color]	192
Practicum: What Stresses Us?	Analyzing merged data [Stress/Chill & Personality Color]	194
Section 4: What's Normal?		195
Lesson 16: What is Normal?	Introduction to normal curve	196
Lesson 17: A Normal Measure of Spread	Measures of spread – SD	199
Lesson 18: What's Your Z-score?	z-scores, shuffling	202
Lab 2H: Eyeballing Normal	Normal curves overlaid on distributions & simulated data	206
Lab 2I: R's Normal Distribution Alphabet	Normal probability, rnorm(), pnorm(), qnorm()	208
End of Unit 2 Project: Comparing Groups Using Our Own Data	Synthesis of Unit 2 [Stress/Chill, Personality Color, FoodHabits, or TimeUse]	210

Table of Contents (continued)

Content	Page
Unit 3	211
Daily Overview	212
Essential Concepts	213
Section 1: Testing, Testing...1, 2, 3...	215
Lesson 1: Anecdotes vs. Data	217
Lesson 2: What is an Experiment?	220
Lesson 3: Let's Try an Experiment!	223
Lesson 4: Predictions, Predictions	225
Lesson 5: Time Perception Experiment	227
Lab 3A: The Results Are In!	229
Practicum: TB or Not TB	230
Section 2: Would You Look at That?	232
Lesson 6: Observational Studies	234
Lesson 7: Observational Studies vs. Experiments	236
Lesson 8: Monsters that Hide in Observational Studies	238
Lab 3B: Confound it All!	242
Section 3: Are You Asking Me?	244
Lesson 9: Survey Says...	245
Lesson 10: We're So Random	248
Lesson 11: The Gettysburg Address	252
Lab 3C: Random Sampling	257
Lesson 12: Bias in Survey Sampling	259
Lesson 13: The Confidence Game	262
Lesson 14: How Confident Are You?	265
Lab 3D: Are You Sure About That?	267
Practicum: Let's Build a Survey	270
Section 4: What's the Trigger?	271
Lesson 15: Ready, Sense, Go!	272
Lesson 16: Does It Have a Trigger?	275
Lesson 17: Creating Our Own PS Campaign	277
Lesson 18: Evaluating Our Own PS Campaign	280
Lesson 19: Implementing Our Own PS Campaign	282
Section 5: Webpages	284
Lesson 20: Online Data-ing	285
Lab 3E: Scraping Web Data	288
Lab 3F: Maps	290
Lesson 21: Learning to Love XML	292
Lesson 22: Changing Format	297
End of Unit 3 Project: Analyzing Our Own Campaign Data	300

Table of Contents (continued)

Content		Page
Unit 4	Topics [Campaign]	301
Daily Overview		302
Essential Concepts		303
Section 1: Campaigns and Community		305
Lesson 1: Trash	Modeling to answer real world problems, official datasets	307
Lesson 2: Drought	Exploratory data analysis, campaign creation	310
Lesson 3: Community Connection	Community topic research, campaign creation	312
Lesson 4: Evaluate and Implement the Campaign	Evaluate & mock-implement campaign	315
Lesson 5: Refine and Create the Campaign	Revise and edit campaign, data collection [Team Campaign]	317
Section 2: Predictions and Models		318
Lesson 6: Statistical Predictions in One Variable	One variable predictions using a rule [Team Campaign]	320
Lesson 7: Statistical Predictions Applying the Rule	Predictions applying MSE, MAE [Team Campaign]	322
Lesson 8: Statistical Predictions Using Two Variables	Two-variable statistical predictions [Team Campaign]	326
Lesson 9: The Spaghetti Line	Estimate line of best fit, linear regression [Team Campaign]	329
Lab 4A: If the Line Fits...	Estimate line of best fit [Team Campaign]	331
Lesson 10: What's the Best Line?	Predictions based on linear models [Team Campaign]	333
Lab 4B: What's the Score?	Comparing predictions to real data [Team Campaign]	336
Lab 4C: Cross-Validation	Use training & test data for predictions [Team Campaign]	338
Lesson 11: What's the Trend?	Trend, associations, linear model [Team Campaign]	341
Lesson 12: How Strong Is It?	Correlation coefficient, strength of trend [Team Campaign]	345
Lab 4D: Interpreting Correlations	Correlation coefficient, best model [Team Campaign]	348
Lesson 13: Improving Your Model	Non-linear regression [Team Campaign]	351
Lab 4E: Some Models Have Curves	Non-linear regression [Team Campaign]	353
Practicum: Predictions	Linear regression [Team Campaign]	355
Section 3: Piecing it Together		356
Lesson 14: More Variables to Make Better Predictions	Multiple linear regression [Team Campaign]	358
Lesson 15: Combination of Variables	Multiple linear regression [Team Campaign]	361
Lab 4F: This Model is Big Enough for All of Us!	Multiple linear regression [Team Campaign]	364
Section 4: Decisions, Decisions		365
Lesson 16: Football or Futbol?	Multiple predictors, classifying into groups [Team Campaign]	366
Lesson 17: Grow Your Own Decision Tree	Decision trees based on training/test data [Team Campaign]	372
Lab 4G: Growing Trees	Decision trees to classify observations [Team Campaign]	376
Section 5: Ties that Bind		379
Lesson 18: Where do I belong?	Clustering, k-means [Team Campaign]	380
Lab 4H: Finding Clusters	Clustering, k-means [Team Campaign]	386
Lesson 19: Our Class Network	Clustering, k-means [Team Campaign]	388
End of Unit 4 Project: Modeling a Community Issue	Synthesis of Unit 4 [Team Campaign]	391

Introduction to Data Science: Overview & Philosophy

Course Overview

Goals

Introduction to Data Science (IDS) is designed to introduce students to the exciting opportunities available at the intersection of data analysis, computing, and mathematics through hands-on activities. Data are everywhere, and this curriculum will help prepare students to live in a world of data. The curriculum focuses on practical applications of data analysis to give students concrete and applicable skills. Instead of using small, tailored, curated datasets as in a traditional statistics curriculum, this curriculum engages students with a wider world of data that fall into the “Big Data” paradigm and are relevant to students’ lives. In contrast to the traditional formula-based approach, in IDS, statistical inference is taught algorithmically, using modern randomization and simulation techniques. Students will learn to find and communicate meaning in data, and to think critically about arguments based on data.

This curriculum was developed in partnership with the Los Angeles Unified School District for a culturally, linguistically, and socially diverse group of students. Upon first publication of the IDS curriculum in 2015, the district-wide student ethnicities included .3% American Indian, 3.7% Asian, .4% Pacific Islander, 2.3% Filipino, 73.0% Latino, 10.9% African American, 8.8% White, and .6% other/multiple responses. Over 38% of students were English-language learners – most of whom spoke Spanish as their primary language – and 74% of students qualified for free or reduced lunches.

Standards

The standards used for the IDS curriculum are based on the High School Probability and Statistics Mathematics Common Core State Standards (**CCSS-M**) and include the Standards for Mathematical Practice (**SMP**). Specific standards are delineated in the scope and sequence section. The Computer Science Teachers Association (**CSTA**) K-12 Computer Science Standards were also consulted and incorporated. Applied Computational Thinking Standards (**ACT**) delineate the application of Data Science concepts using technology.

Hardware

An ideal laboratory environment has a 1:1 computer to student ratio. The computers can be either Apple, PC, or Chromebook, depending upon availability. Internet access is required for the use of RStudio on an external server. The IDS instructor must have access to a computer and a projector for daily use.

Software

Each computer (tablets are not recommended) in the classroom should have a modern, updated web browser installed (such as Firefox or Google Chrome). This will allow students to access RStudio via the RStudio Cloud platform, and to perform searches and make use of a variety of websites and internet tools. RStudio is accessible at <https://posit.cloud/> or through the IDS home page at <https://portal.thinkdataed.org>. The IDS team will provide the remainder of the software used in the IDS curriculum, also available at <https://portal.thinkdataed.org>.

This software includes the IDS ThinkData Ed app, which is deployed for Android and iOS (Apple) smartphones and tablets, as well as through a web browser on a desktop or laptop computer via the IDS home page. The app allows students to collect the Participatory Sensing data that is a motivational foundation for the course. In addition to the app, students will use the IDS software to access and manipulate their Participatory Sensing data, and to author their own campaigns.

All computer-based assignments are intended to be completed in class to avoid the assumption that students have access to computers at home. However, if a student misses a lab assignment, they will need to make it up on their own time. All the software required for the curriculum is available via the Internet, so students can complete the assignment on any Internet-enabled computer (e.g., at the school or public library).

Prerequisites

It is recommended that students successfully complete a first-year Algebra course prior to taking IDS. With this background, the curriculum provides a rigorous but accessible introduction to data science and statistics. No previous statistics or computer science courses are required to take this course.

The Instructional Philosophy of Introduction to Data Science

IDS uses a project-based learning approach to instruction. Finkle and Torp (1955) define Project-Based Learning (PBL) as a curriculum development and instructional system that simultaneously develops both problem-solving strategies and disciplinary knowledge bases and skills by placing students in the active role of problem solvers confronted with an ill-structured problem that mirrors real-world problems. PBL, therefore, is a model for teaching and learning that focuses on the main concepts and principles of a discipline, involves students in problem-solving investigations and other meaningful tasks, allows students to construct their own knowledge through inquiry, and culminates in a project.

Because IDS is a mathematical science, the BSCS 5-E Instructional Model provides a planned sequence of instruction that places students at the center of their learning experiences. This model encourages students to explore, create their own meaning of concepts, and relate their understanding to other concepts. The units in IDS contain lessons that, together, fit the 5-E Instructional Model:

Stage of Inquiry in an Inquiry-Based Science Program	Possible Student Behavior	Possible Teacher Strategy
Engage	Asks questions such as, Why did this happen? What do I already know about this? What can I find out about this? How can I solve this problem? Shows interest in the topic.	Creates interest. Generates curiosity. Raises questions and problems. Elicits responses that uncover student knowledge about the concept/topic.
Explore	Thinks creatively within the limits of the activity. Tests predictions and hypotheses. Forms new predictions and hypotheses. Tries alternatives to solve a problem and discusses them with others. Records observations and ideas. Suspends judgment. Tests ideas.	Encourages students to work together without direct instruction from the teacher. Observes and listens to students as they interact. Asks probing questions to redirect students' investigations when necessary. Provides time for students to puzzle through problems. Acts as a consultant for students.
Explain	Explains their thinking, ideas, and possible solutions or answers to other students. Listens critically to other students' explanations. Questions other students' explanations. Listens to and tries to comprehend explanations offered by the teacher. Refers to previous activities. Uses recorded data in explanations.	Encourages students to explain concepts and definitions in their own words. Asks for justification (evidence) and clarification from students. Formally provides definitions, explanations, and new vocabulary. Uses students' previous experiences as the basis for explaining concepts.
Elaborate	Applies scientific concepts, labels, definitions, explanations, and skills in new, but similar situations. Uses previous information to ask questions, propose solutions, make decisions, and design experiments. Draws reasonable conclusions from evidence. Records observations and explanations.	Expects students to use vocabulary, definitions, and explanations provided previously in new context. Encourages students to apply the concepts and skills in new situations. Reminds students of alternative explanations. Refers students to alternative explanations.
Evaluate	Checks for understanding among peers. Answers open-ended questions by using observations, evidence, and previously accepted explanations. Demonstrates an understanding or knowledge of the concept or skill. Evaluates his or her own progress and knowledge. Asks related questions that would encourage future investigations.	Refers students to existing data and evidence and asks, What do you know? Why do you think...? Observes students as they apply new concepts and skills. Assesses students' knowledge and/or skills. Looks for evidence that students have changed their thinking. Allows students to assess their learning and group process skills. Asks open-ended questions such as, Why do you think...? What evidence do you have? What do you know about the problem? How would you answer the question?

IDS is designed to develop students' computational and statistical thinking skills. Computationally, students will learn to write code to enhance analyses of data, to break large problems into smaller pieces, and to understand and employ algorithms to solve problems. Statistical thinking skills include developing a data "habit of mind" in which one learns to seek data to answer questions or support (or undermine) claims; thinking critically about the ability of particular data to support claims; learning to interpret analyses of data; and learning to communicate findings.

IDS employs Participatory Sensing to give students control of the data collection process, and to enable them to collect data about things that are important to them. The curriculum is organized around a series of Participatory Sensing "campaigns" in which students engage in all stages of the statistical process, which we call the Data Cycle: posing questions, examining and collecting data, analyzing data, interpreting data and, if necessary, beginning again. As students progress, they engage in the Data Cycle in a deeper way. Initially, analysis and interpretation are purely descriptive. Later, randomization-based algorithms and simulations are used to develop notions of inference and to make students more critical of the data collection process. By engaging in the Data Cycle repeatedly in different contexts - some of which include the students' own designs - students will learn to think like data scientists.

Student Team Collaboration

Many of the activities in the IDS curriculum are based on students collaborating with each other. Activities may call on pairs or teams of students. **It is imperative that teams and team roles be established as close to the beginning of the course as possible.** Expectations about teamwork should be introduced as soon as teams are formed. The ideal team comprises four students. The Teacher Resources section provides a list of instructional strategies and a description of team roles to use for effective student team collaboration. If student teams are unfamiliar with these instructional strategies, it is important for the instructor to take the time to model each strategy.

Classroom Discussions

Because this is an inquiry-based curriculum, classroom discussion will be especially important. It is important to set classroom discussion norms from the beginning of the course. All students should be encouraged to contribute to the classroom discussion, and the learning environment should be as non-judgmental and as open as possible. Instead of one right answer, most questions in this class have many right answers. In fact, even yes/no questions could have two right answers, both with valid supporting evidence. Teachers should create an environment to help students hold each other accountable so that all voices are heard, meaning that if there are a few students who tend to share a lot, invite them to encourage their peers so other voices can be heard. If there are students who tend to avoid contributing to the class discussion, encourage them to share so that their voices are heard.

Assignments & Homework

As much as possible, IDS work will take place in the classroom. Lessons are designed for a 50–60-minute class period. Classes on block schedule will need to complete two lessons; however, it is up to the teacher to decide where to stop in each lesson. There will be open-ended assignments that are sent home. Assignments that require the computer will be completed in class, to avoid the assumption that students have access to computers at home. The exception to this is if a student misses lab time, in which case they will need to find a time to complete the assignment outside of class. As discussed in the software section above, they can use an Internet-enabled computer to do their make-up work.

IDS assignments will not be drill-based. Instead, they will follow the inquiry-based instructional model. Again, most questions will not have one right answer. Instead, students will learn to support their claims with evidence and to participate in data-based discussions. Newspaper or other periodical or digital articles are available via links in the lessons. If desired, articles may be downloaded and printed.

On average, students will complete a lab assignment in RStudio approximately once per week. It will be at the discretion of the teacher whether or not to collect lab assignments. Calculators should be available every day for students to use.

Every day, students will be expected to bring their Data Science (DS) journal, a notebook where they record their notes, work on small assignments, and sketch plots. Teachers may choose to check DS journals and other assignments in the curriculum for credit.

End of Unit Projects, oral presentations, performance-task-based assessments, and Practicums are designed as application exercises. Scoring guides are provided as an aid for student performance expectations. It will be up to the teacher to score or attach a grade to these assignments.

Overview of Instructional Topics

The purpose of IDS is to introduce students to dynamic data analysis. The four major components of this curriculum are based on the conceptual categories called upon by the Common Core State Standards High School - Statistics and Probability:

- I. Interpreting Categorical and Quantitative Data**
- II. Making Inferences and Justifying Conclusions**
- III. Conditional Probability and the Rules of Probability**
- IV. Using Probability to Make Decisions**

IDS will emphasize the use of statistics and computation as tools for creative work, and as a means of telling stories with data. Seen in this way, its content will also prepare students to “read” and think critically about existing data stories. Ultimately, this course will be about how we discern good stories from bad through a practice that involves compiling evidence from one or more sources, and which often requires hands-on examination of one or more datasets.

IDS will develop the tools, techniques, and principles for reasoning about the world with data. It will present a process that is iterative and authentically inquiry-based, comparing multiple “views” of one or more datasets. Inevitably, these views are the result of some kind of computation, producing numerical summaries or graphical displays. Their interpretation relies on a special kind of computation known as simulation to describe the uncertainty in each view. This kind of reasoning is exploratory and investigatory, sometimes framed as hypothesis evaluation, and sometimes as hypothesis generation.

Interpreting Categorical and Quantitative Data

A handful of data interpretations are standard. Some, including summaries of shape, center, and spread of one or more variables in a dataset – as well as graphical displays like histograms and scatterplots – are standard in the sense that they provide interpretable information in a number of research contexts. They are portable from one set of data to the next, and the rules for their use are simple. And yet, our interpretation of data is rarely “standard”. Data have no natural look - even a spreadsheet or a table of numbers embeds within it a certain representational strategy. We construct multiple views of data in an attempt to uncover stories about the world.

In addition to numerical data, this course will consider time, location, text, and image as data types, and will examine views that uncover patterns or stories. Throughout the course, simulation will be used to calibrate our interpretation of a view, or of a numerical or graphical summary, so that we understand what “story-less” data (i.e., pure noise, no association) look like.

In addition to summaries and simple graphics, students will engage in a modeling practice aligned with the CCSS mathematical practices in order to learn how statistical analyses can explain and describe real-world phenomena. Students will practice fitting and evaluating standard mathematical and statistical models, such as the least-squares regression line. Modeling comes into play when students are asked to design and implement probabilistic simulations in order to test and compare hypothetical chance processes to real-world data.

Making Inferences and Justifying Conclusions

Data are becoming increasingly plentiful, supported by a host of new “publication” techniques or services. Post-Web 2.0, data are interoperable, flowing out of one service and into another, helping us easily build a detailed data version of many phenomena in the world. Reasoning with data, then, starts with the sources and the mechanics of this flow. Which sources do we trust? How do data from different organizations compare? What stories have been told previously with these data, and by whom?

This course answers these questions, in part, by using the tools and techniques already mentioned. The ability to read and critique published stories and visualizations are additions to these tools and techniques. Finally, as an act of comparison, students should also be able to formulate questions, identify existing datasets, and evaluate how the new stories stack up against the old. To support this cycle of inquiry, students will examine the basic publication mechanisms for data and develop a set of questions to ask of any data source - computation meets critical thinking. In some cases, data will exhibit special structures that can be used to aid in inference. The simulation techniques for calibrating different views of a dataset take on new life when some form of random process was followed to generate the data. Polls, for example, rely on random samples of the population, and clinical trials randomly assign patients to treatment and control groups. A simulation strategy that repeats these random mechanisms can be used to assess uncertainty in the data, assigning a margin of error to poll results, or identifying new drugs that have a “significant” effect on some health outcome.

In many cases, data will not possess this kind of special origin story. A census, for example, is meant to be a complete enumeration of a population, and we can reason in a very direct way from the data. In other cases, no formal principle was applied, perhaps being a sample “of convenience.” The techniques for telling stories from these kinds of data will also rely on a mix of simulation and subsetting or filtering. Finally, this course will introduce Participatory Sensing as a technique for collecting data. The idea of a data collection campaign will be introduced as a means of formalizing a question to be addressed with data. Campaigns will be informed by research and data analysis, and will build on, augment, or challenge existing sources. The “culture” behind the existing sources and the summaries or views they promote will be part of the classroom discussions.

It is worth noting that everything described so far depends on computation, using a piece of statistical software on a computer. Students will be taught simple programming tools for accessing data, creating views or fitting models, and then assessing their importance via simulation. Computation becomes a medium through which students learn about data. The more expressive the language, the more elaborate the stories we can tell.

Probability

Since simulation is our main tool for reasoning with data, interpreting the output of simulations requires understanding some basic rules of probability. First and foremost, this course will discuss the ways in which a computer can generate random phenomena (e.g., How does a computer toss a coin?). Simple probability calculations will be used to describe what we expect to see from random phenomena, then students will compare their results to simulations. The point is to both rehearse these basic calculations and to make a formal tie between simulation and theory in simple cases.

In that vein, this course will motivate the relationship between frequency and probability. Students will essentially be simulating independent trials and creating summaries of those simulations. In turn, they should understand that the frequency with which an event occurs in a series of independent simulations tends to the probability for that event as the number of simulations gets large (the Law of Large Numbers, a topic that is often taught in introductory statistics courses).

From here, students will simulate a variety of random processes to aid in formal statistical inference when some random mechanism was applied as part of the data design. In short, probability becomes a ruler of sorts for assessing the importance of any story we might tell. In this approach to probability, a combination of direct mathematical calculation and computer simulations will be used in order to give students a deep sense of the underlying statistical concepts.

Topic Outline

This outline describes only the scope of the course; the sequence is described in each unit.

I. Interpreting Data

- A. Types of data
- B. Numerical and graphical summaries
 - 1. Measures of center and spread, boxplots
 - 2. Bar plots
 - 3. Histograms
 - 4. Scatterplots
 - 5. Graphical summaries of multivariate data
- C. Simulation and visual inference
 - 1. Side-by-side bar plots and association
 - 2. Scatterplots
- D. Models
 - 1. Linear models
 - 2. k-means
 - 3. Smoothing
 - 4. Learning and tree-based models

II. Making Inferences and Justifying Conclusions

- A. Aggregating data
 - 1. Identification of sources
 - 2. Mechanics of Web 2.0
 - 3. Comparison of sources
- B. Data with special structures
 - 1. Random sampling
 - 2. Random assignment and A/B testing
 - 3. Simulation-based inference
- C. Participatory Sensing
 - 1. Designing a campaign
 - 2. Participation as a data collection strategy

III. Probability

- A. Computers and randomness
 - 1. Web services
 - 2. Pseudo-random numbers (optional)
- B. Frequency and probability
- C. Probability calculations

IV. Algebra in RStudio

- 1. Vectors
- 2. Algorithms
- 3. Functions
- 4. Evaluating and fitting models to data
- 5. Graphical representations of multivariate data
- 6. Numerical summaries of distributions and interpreting in context

Scope and Sequence

Unit 1

This unit will introduce the idea of “data,” fundamental to the rest of the course. While most people think of data simply as a spreadsheet or a table of numbers, almost anything can be considered data, including images, text, GPS coordinates, and much more. Our world has become increasingly data-centric, and we are constantly generating data, whether we know it or not. From posts on social media, to shopping records created when you swipe your credit card, to driving over sensors embedded in highway on-ramps, we leave behind a stream of data wherever we go. These data are used to generate stories about our world, whether it is for political forecasting, marketing, scientific research, or even Netflix recommendations. Traditional statistics courses consist of understanding data from only a small subset of data generation processes, namely those collected through random sampling or random assignment in scientific experiments. This unit exposes students to a wider world of data and will help students see how to make sense of these ubiquitous data types.

This unit will motivate the idea that data and data products (charts, graphs, statistics) can be analyzed and evaluated just like other arguments, such as those used by journalists. We want to know how the evidence was collected, what the perspective or bias of the creator might be and look behind the scenes to the process used to create the product. Even the way data are represented embeds within it decisions on the part of the data creator.

Using the techniques of descriptive statistics, students will begin learning how to construct multiple views of data in an attempt to uncover new insights about the world. This will require the introduction of the computational tool R through the interface of RStudio. Standard graphical displays like histograms and scatterplots will be introduced in RStudio, as well as measures of center and spread.

Focus Statistics/Mathematics Standards

- S-ID 1: Represent data with plots on the real number line (dotplots, histograms, and boxplots).
- S-ID 2: Use statistics appropriate to the shape of the data distribution to compare center (median, mean) of two or more different data sets (measures of spread will be studied in Unit 2).
- S-ID 3: Interpret differences in shape, center, and spread in the context of the data sets, accounting for possible effects of extreme data points (outliers).
- S-ID 5: Summarize categorical data for two categories in two-way frequency tables. Interpret relative frequencies in the context of the data (joint, marginal, and conditional relative frequencies). Recognize possible associations and trends in the data.
- S-ID 6: Represent data on two quantitative variables on a scatterplot and describe how the variables are related.
- S-IC 6: Evaluate reports based on data.*

*This standard is woven throughout the course. It is a recurring standard for every unit.

Focus Standards for Mathematical Practices

- SMP-3. Construct viable arguments and critique the reasoning of others.
- SMP-5. Use appropriate tools strategically.

Upon completion of Unit 1, students will be able to:

- Give examples of where they leave data traces.
- Understand that rows and columns are a form of data structure.
- Explain why the relationship between the variables might exist, or, if there is no relationship, why that might be so.
- Construct and interpret a frequency table.
- Critically read reports from media sources to evaluate their claims.
- Read plots (identify the name of the plot, interpret the axes, look for trends, identify confounding factors).
- Calculate conditional and marginal probabilities using frequency tables.
- Provide a real-world explanation for why the conditional or independent probabilities make sense, using critical thinking skills and background knowledge.

- Communicate their evaluations in written or verbal form using different types of media.
- Load data into RStudio.
- Create basic plots in RStudio.
- Create frequency tables in RStudio.

Unit 2

This unit deepens the informal reasoning skills developed in Unit 1 by enriching students' technical vocabulary and developing more precise analytical tools. Most importantly, this unit introduces the formal concept of probability as a tool for understanding that sometimes patterns observed in data are not "real." Traditional courses attempt to develop this understanding through the development of abstract mathematical probability concepts, but IDS creates enduring understanding by teaching students to design and implement simulations using pseudo-random number generators. This activity also develops computational thinking by teaching students about some basic programming structures. Then, the use of models will come to the foreground. Students will be introduced to linear models - the most common form of modeling in introductory statistics classes - which will serve as the foundation to learn more complex modeling techniques that use the computer technology available to them later in the course, including smoothing techniques and tree-based models.

Focus Statistics/Mathematics Standards

A-SSE A: Interpret the structure of expressions.

S-ID 2: Use statistics appropriate to the shape of the data distribution to compare center (median, mean) and spread (interquartile range, standard deviation) of two or more different data sets.

S-ID 3: Interpret differences in shape, center, and spread in the context of the data sets, accounting for possible effects of extreme data points (outliers).

S-ID 4: Use the mean and standard deviation of a data set to fit it to a normal distribution and to estimate population percentages. Understand that there are data sets for which such a procedure is not appropriate. Use calculators, spreadsheets, and tables to estimate areas under the normal curve.

S-IC 2: Decide if a specified model is consistent with results from a given data-generating process, e.g., using simulation.

S-IC 6: Evaluate reports based on data.*

*This standard is woven throughout the course. It is a recurring standard for every unit.

S-CP 2: Understand that two events A and B are independent if the probability of A and B occurring together is the product of their probabilities, and use this characterization to determine if they are independent.

S-CP 9: (+) Use permutations to perform [informal] inference.

*This standard will be addressed in the context of data science.

Focus Standards for Mathematical Practices

SMP-4. Model with mathematics.

SMP-5. Use appropriate tools strategically.

Upon completion of Unit 2, students will be able to:

- Create a boxplot by calculating the five-number summary, upper and lower fences, and determining outliers.
- Explain what "standard deviation" means in context.
- Explain why the measures of central tendency and spread may or may not be accurate descriptions of the data from which they came.
- Use permutations of data to solve problems.
- Read/interpret a normal curve/distribution.
- Explain where the normal distribution came from.
- Describe situations where the normal distribution may model the phenomena, and others where it may not.
- Simulate normal distribution.

- Simulate from a model.
- Compare real data to simulation.
- Determine if model and data appear consistent.
- Merge data by columns/rows and verify that merging is successful.
- Create functions.

Unit 3

Unit 3 focuses on data collection methods, including traditional methods of designed experiments and observational studies and surveys. It introduces students to sampling error and bias, which cause problems in analysis made from survey data. Participatory Sensing is presented as another method of data collection, and students learn to design Participatory Sensing campaigns that will allow them to address particular statistical questions. Participatory Sensing is a unique data collection method because it uses sensors. Furthermore, this method emphasizes the involvement of citizens and community groups in the process of sensing and documenting where they live, work, and play. Triggers play an important role in the Participatory Sensing data collection process. The response to the triggers may or may not be the same each time. Data takes on a variety of forms online and requires a different style of representation. Students enhance computing skills by learning about modern data structures, and by learning to “scrape” data stored in HTML format.

Focus Statistics/Mathematics Standards

S-IC 1: Understand statistics as a process for making inferences about population parameters based on a random sample from that population.

S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.

S-IC 6: Evaluate reports based on data.*

*This standard is woven throughout the course. It is a recurring standard for every unit.

Focus Standards for Mathematical Practices

SMP-1. Make sense of problems and persevere in solving them.

SMP-4. Model with mathematics.

SMP-8. Look for and express regularity in repeated reasoning.

Upon completion of Unit 3, students will be able to:

- Provide a loose definition of “statistics” in their own words.
- Compare and contrast population vs. sample.
- Compare and contrast parameter vs. statistic.
- Explain the difference between special data structures, particularly as they relate to inference.
- Exploit special data structures for re-randomization analysis.
- Explain situations where one measure of central tendency or spread may be more appropriate than others.
- Read/interpret boxplots (in-depth look into samples size and their relationship to the population parameters).
- Identify reports that use special data structures (census, survey, observational study, and randomized experiment).
- Do data scraping.
- Use HTML and XML formats.
- Use RStudio to re-randomize data.
- Compute measures of central tendency and spread in RStudio.

Unit 4

This unit will develop modeling skills, beginning with learning to fit and interpret least squares regression lines and learning to use regression to make predictions. Students will learn to evaluate the success of these predictions and so compare models for their predictive accuracy. Modern algorithmic approaches to

regression are presented, and students will strengthen algorithmic thinking skills by understanding how and why these algorithms help data scientists make accurate predictions from data. Students engage in a complete modeling experience in which they apply the skills and concepts learned in the previous units. The modeling experience is designed to make students' thinking visible and audible by encouraging them to be metacognitive about the process of inventing and testing a model, ask questions as they go through the process, and recognize the iterative nature of modeling.

Focus Statistics/Mathematics Standards

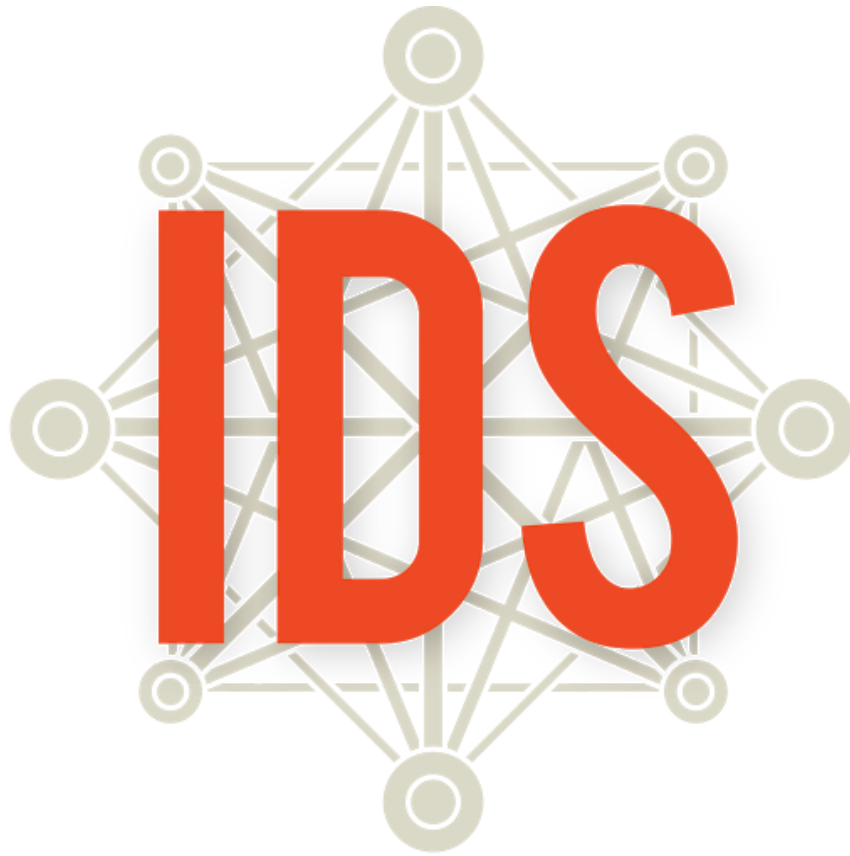
- A-SSE A: Interpret the structure of expressions.
- A-REI D-10: Understand that the graph of an equation in two variables is the set of all its solutions plotted in the coordinate plane, often forming a curve (which could be a line).
- F-IF A-2: Use function notation, evaluate functions for inputs in their domains, and interpret statements that use function notation in terms of context.
- F-BF A-1: Write a function that describes a relationship between two quantities
- S-IC 2: Decide if a specified model is consistent with results from a given data-generating process, e.g., using simulation.
- S-ID 6: Represent data on two quantitative variables on a scatter plot and describe how the variables are related.
- Fit a function to the data; use functions fitted to data to solve problems in the context of the data. *Use given functions or choose a function suggested by the context. Emphasize linear models.*
 - Informally assess the fit of a function by plotting and analyzing residuals.
 - Fit a linear function for a scatter plot that suggests a linear association.
- S-ID 7: Interpret the slope (rate of change) and the intercept (constant term) of a linear model in the context of the data.
- S-ID 8: Compute (using technology) and interpret the correlation coefficient of a linear fit.
- S-IC 6: Evaluate reports based on data.*
- *This standard is woven throughout the course. It is a recurring standard for every unit.

Focus Standards for Mathematical Practices

- SMP-2. Reason abstractly and quantitatively.
- SMP-4. Model with mathematics.
- SMP-7. Look for and make use of structure.

Upon completion of Unit 4, students will be able to:

- Describe how well the linear model fits the data (or does not).
- Provide a real-world explanation of why the model may or may not fit, using critical thinking skills and background knowledge.
- Interpret the slope and intercept on a plot.
- Compute the correlation coefficient using RStudio.
- Interpret linear models in reports, including the correlation coefficient.
- Determine if a trend is “real” or if it could have arisen from randomness.
- Use critical thinking skills to explain why a trend may or may not make sense.
- Fit a regression line.
- Extract the slope, intercept, correlation coefficient, coefficient of determination, and residuals using RStudio.
- Use RStudio to predict y given an x value.
- Explore what happens to the line and the response variable if we multiply (divide) or add (subtract) a constant from the predictor.
- Design and execute their own Participatory Sensing Campaigns.
- Use RStudio to compute permutations and combinations.
- Create Classification and Regression Tree (CART) models.
- Understand non-linear models.



Introduction to Data Science

Unit 1

Introduction to Data Science

Daily Overview: Unit 1

Theme	Day	Lessons and Labs	Campaign	Topics	Page
Data Are All Around (7 days)	1	Lesson 1: Data Trails		Defining data, consumer privacy	24
	2	Lesson 2: Stick Figures		Organizing & collecting data	26
	3	Lesson 3: Data Structures		Organizing data, rows & columns, variables	28
	4	Lesson 4: The Data Cycle		Data cycle, statistical questions	30
	5	Lesson 5: So Many Questions		Statistical questions, variability	35
	6	Lesson 6: What Do I Eat?	Food Habits	Data cycle, collecting data	39
	7^	Lesson 7: Setting the Stage	Food Habits – data	Participatory Sensing	42
Visualizing Data (14 days)	8	Lesson 8: Tangible Plots	Food Habits – data	Dotplots, minimum/maximum, frequency	48
	9	Lesson 9: What Is Typical?	Food Habits – data	Typical value, center	52
	10	Lesson 10: Making Histograms	Food Habits – data	Histograms, bin widths	55
	11	Lesson 11: What Shape Are You In?	Food Habits – data	Shape, center, spread	58
	12	Lesson 12: Exploring Food Habits	Food Habits – data	Single & multi-variable plots	60
	13	Lesson 13: RStudio Basics	Food Habits – data	Intro to RStudio	62
	14	Lab 1A: Data, Code & RStudio	Food Habits – data	RStudio basics	65
	15+	Lab 1B: Get the Picture?	Food Habits – data	Variable types, bar graphs, histograms	69
	16	Lab 1C: Export, Upload, Import	Food Habits – data	Importing data	72
	17	Lesson 14: Variables, Variables, Variables		Multi-variable plots	77
	18	Lab 1D: Zooming Through Data		Subsetting	81
	19	Lab 1E: What's the Relationship?		Multi-variable plots	85
	20	Practicum: The Data Cycle & My Food Habits	Food Habits	Data cycle, variability	88
	21	Practicum Presentations	Food Habits	Data cycle, variability	-
Would You Look at the Time? (9 Days)	22^	Lesson 15: Americans' Time on Task	Time Use – data	Evaluating claims	91
	23	Lab 1F: A Diamond in the Rough	Time Use - data	Cleaning names, categories, and strings	96
	24	Lesson 16: Categorical Associations	Time Use - data	Joint relative frequencies in 2-way tables	101
	25	Lesson 17: Interpreting Two-Way Tables	Time Use - data	Marginal & conditional relative frequencies	103
	26+	Lab 1G: What's the FREQ?	Time Use – data	2-way tables, tally	108
	27	Practicum: Teen Depression	Time Use	Statistical question, interpreting plots	111
	28	Practicum Presentations		Statistical question, interpreting plots	-
	29-30	Lab 1H: Our Time		Data cycle, synthesis	113
End of Unit Project (5 Days)	31-35	End of Unit 1 Project: Evaluating Claims from the Media		Data cycle	114

^=Data collection window begins.

+ =Data collection window ends.

IDS Unit 1: Essential Concepts

Lesson 1: Data Trails

Data are a collection of recorded observations. Data are gathered by people and by sensors. Patterns in data can reveal previously unknown patterns in our world. Data play a large, and sometimes invisible, role in our lives.

Lesson 2: Stick Figures

Data consist of records of particular characteristics of people or objects. Data can be organized in many different ways, and some ways make it easier than others for achieving particular purposes.

Lesson 3: Data Structures

Variables record values that vary. By organizing data into rectangular format, we can easily see the characteristics of observations by reading across a row, or we can see the variability in a variable by reading down the column. Computers can easily process data when it is in rectangular format.

Lesson 4: The Data Cycle

A statistical investigation consists of cycling through the four stages of the Data Cycle. The term statistical questions encompasses the variety of questions asked during the statistical problem-solving process which support statistical thinking and reasoning. Statistical questions are perhaps the most important because they are challenging to learn and are the types of questions that determine whether an analysis is productive or not. Statistical questions are questions that address variability and are productive in that they motivate data collection, analysis, and interpretation. The Data Collection phase might consist of collecting data through Participatory Sensing or some other means, or it might consist of examining previously collected data to determine the quality of the data for answering the statistical questions. Data Analysis is almost always done on the computer and consists of creating relevant graphics and numerical summaries of the data. Data Interpretation is involved with using the analysis to answer the statistical questions.

Lesson 5: So Many Questions

Statistical questions typically begin with a vague general question, then develop into a precise question. The process of developing or creating a good statistical question is iterative and requires time and effort to get right. In her 2021 paper, What Makes a Good Statistical Question, Dr. Pip Arnold identified the following as features of a good statistical question:

- (1) The variable(s) of interest is/are clear
- (2) The group or population we are interested in is clear
- (3) The question can be answered with the data
- (4) The question asks about the whole group, not an individual or portion of the group
- (5) The intention is clear (e.g., summary, comparison, association, time series)
- (6) The question is one that is worth investigating, is interesting, and has a purpose

Lesson 6: What Do I Eat? [The Data Cycle: Consider Data]

After raising statistical questions, we examine and record data to see if the questions are appropriate.

Lesson 7: Setting the Stage [The Data Cycle: Collect Data]

In Participatory Sensing, we humans behave as if we are robot sensors, collecting data whenever a “trigger” event occurs. Our ability to learn about the patterns in our life through these data depends on our being reliable data collectors.

Lesson 8: Tangible Plots [The Data Cycle: Analyze Data]

Distributions organize data for us by telling us (a) which values of a variable were observed, and (b) how many times the values were observed (their frequency).

Lesson 9: What Is Typical?

The “center” of a distribution is a deliberately vague term, but it is one way to answer the subjective question “What is a typical value?”. The center could be the perceived balancing point or the value that approximately cuts the area of the distribution in half.

Lesson 10: Making Histograms

Histograms can be created through the use of an algorithm. The distributions displayed in a histogram can be classified using the technical terms for the shapes of distributions. Learning to describe routine tasks through an algorithm is an important component of computational thinking.

Lesson 11: What Shape Are You In?

Identifying the shape of a histogram is part of the **interpret** step of the Data Cycle.

Lesson 12: Exploring Food Habits

Once Participatory Sensing data has been collected, the Dashboard and PlotApp perform the analysis step of the Data Cycle, though humans need to tell the computer which plots to examine.

Lesson 13: RStudio Basics

The computer has a syntax, and it can only understand if you speak its language.

Lesson 14: Variables, Variables, Variables

To examine whether two (or more) variables are related, we can plot their distributions on the same graph.

Lesson 15: Americans' Time on Task

Learning to examine other analyses is an important part of statistical thinking.

Lesson 16: Categorical Associations

A two-way table is a summary of the association/relationship between two categorical variables. Joint relative frequencies answer questions of the form “What proportion of the people/objects had *this* value on the first variable and *this* value on the second?”

Lesson 17: Interpreting Two-Way Tables

Marginal (relative) frequencies tell us about the distribution of a single variable. Conditional relative frequencies tell us about the distribution of one variable when “subsetting” the other.

Data Are All Around

Instructional Days: 7

Enduring Understandings

Data play an important role in our everyday lives. Organizing it can provide evidence about real-life events and people. The data collected by answering survey questions produce variability. Distributions, graphs, and plots are useful tools for organizing data to understand variability. Statistical questions address people, processes, and/or events that contain variability. Situations with variability can sometimes be simplified with some basic statistics.

Engagement

The Target Story will introduce students to the idea that data are ubiquitous. The advent of computers has transformed the way data are collected, used, and analyzed. Video can be found at:

<https://www.youtube.com/watch?v=XvSA-6BJkx4&feature=youtu.be>

Note: Pre-loading the video on your computer prior to the beginning of class is highly recommended to avoid any technical difficulties.

Learning Objectives

Statistical/Mathematical:

S-ID: Summarize, represent, and interpret data on a single count or measurement variable.

S-ID 1: Represent data with plots on the real number line (dotplots, histograms, bar plots, and boxplots).

S-ID 2: Use statistics appropriate to the shape of the data distribution to compare center (median, mean) of two or more different data sets. (Measures of spread will be studied in unit 2.)

S-ID 6: Represent data on two quantitative variables on a scatterplot, and describe how the variables are related.

Focus Standards for Mathematical Practice:

SMP-3: Construct viable arguments and critique the reasoning of others.

SMP-5: Use appropriate tools strategically.

Data Science:

Experience data handling using ubiquitous data and organize data using rectangular or spreadsheet format as data storage structures.

Everyday activities can be observed and recorded as data. Become aware of the difference between plots used for categorical and numerical variables. Interpret and understand graphs of distributions for numerical and categorical variables.

Applied Computational Thinking using RStudio:

- Work effectively in teams.
- Explain how data, information, and knowledge are represented for computational use.
- Collect, upload, and share personal data via a Participatory Sensing campaign.
- Learn about different representations of distributions using software.
- Utilize software to begin to analyze plots of data collected via Participatory Sensing.

Real-World Connections:

Students begin to develop an awareness that data are all around us. Information can be collected and organized. Computers are powerful tools that make organizing, storing, retrieving, and analyzing data accessible to use in problem solving and decision making. Students will begin to see the relevance of data collection to their own lives. They will begin to understand that data on its own is just collected; but once interpreted, it can lead to discoveries or understandings.

Language Objectives

1. Students will use complex sentences to construct summary statements about their understanding of data, how it is collected, how it is used, and how to work with it.
2. Students use mathematical vocabulary to describe data and how it is organized in various forms.
3. Students will use complex sentences to write informative short reports that use data science concepts and skills.

Data File or Data Collection Method

Data Collection Method:

1. Students will keep a Data Diary for 24 hours to track their daily data output.
2. Students will gather data from the cards in the Stick Figures file.
3. As a class, students will determine how to organize the Stick Figures data.
4. Students will collect data using paper and pencil on a *Food Habits Data Collection* activity sheet.
5. **Food Habits Participatory Sensing Campaign:** Students will collect data about their snacking habits.

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Note: Students should have completed the IDS Survey and LOCUS Assessment before teaching this lesson. See Addendum_IDS_Survey_LOCUS in the Documents tool on [Portal](#) for more information.

Lesson 1: Data Trails

Objective:

Students will understand what are data, how they are collected, and possible effects of sharing data.

Materials:

1. Video: *The Target Story* found at:
<https://www.youtube.com/watch?v=XvSA-6BJkx4>
2. Data Science (DS) journal (quad-ruled composition book or similar)
Note: MUST be available for every lesson
3. *Data Diary* handout (LMR_U1_L1)
4. Video: *Terms and Conditions* found at:
<https://www.youtube.com/watch?v=ZcjtEKNP05c>

Vocabulary:

data, observations, data trails, privacy

Essential Concepts: Data are a collection of recorded observations. Data are gathered by people and by sensors. Patterns in data can reveal previously unknown patterns in our world. Data play a large, and sometimes invisible, role in our lives.

Lesson:

Before implementing the IDS curriculum, ensure that:



- Students have been placed in teams and each student understands his or her role in the team.
- Each student knows who his/her partner is within each team.
- Expectations regarding collaborative teamwork are discussed and understood (see Team Roles in Teacher Resources).



1. Introduce the lesson by showing *The Target Story* video:
<https://www.youtube.com/watch?v=XvSA-6BJkx4>



2. In pairs, ask students to discuss the following question using the *TPS* strategy (see Instructional Strategies in Teacher Resources):
 - a. How do you think Target knew about the daughter? In other words, how did Target know the daughter was pregnant before her father did? *Answer: Target used the information gathered from the daughter's Red Card and compared it to information about other shoppers. Typically, women who bought those particular products were pregnant.*



3. After students have had time to share their responses, engage in a whole class discussion regarding:
 - a. What are **data**? *Answer: Data are information, or observations, that have been gathered and recorded.*
 - b. Where do data come from? *Answer: Data can come from a variety of places. Some examples might include: cell phones, computers, school records, surveys, etc.*
 - c. Give an example of data. *Answers will vary. One example might be information about a person – including their age, height, weight, eye color, etc.*
 - d. Give an example of something that is not data (e.g., something that was never written down). *Answers will vary. One example might be just watching an event happen. If it wasn't recorded in some way, it cannot be counted as data.*
4. Explain to the students that we create “**data trails**” as we go through life. A data trail is the data collected about us as individuals that could be used to see the patterns in our personal lives.

Inform the students that they will learn about their own data trails by keeping a data diary and logging entries over the next 24 hours. It is likely that students do not realize how often they leave a data trail or what information is being collected about them on a regular basis.

5. Distribute the *Data Diary* handout (LMR_U1_L1) and be sure to go over the instructions, along with the first example to give the students an idea of how to proceed.

Name: _____ Date: _____

Data Diary

Instructions:

You will keep a data diary for 24 hours. You will write down everything you do that could potentially provide someone with electronic personal data about you without you necessarily choosing to give him/her your information.

Do not include events such as how long you brush your teeth, for example, unless you are transmitting this data to an outside source.

Some good examples might include using Google to do a search online, using TikTok, shopping with a credit card, using a GPS, using an app on your phone, watching a movie on Netflix, texting, etc. An example is given to you in the first line.

Time	Activity	Type of data collected from you
4:00-4:45 pm	Watched "Mad Men" on Netflix.	Viewing interests, time watched, possibly geographic location, account information (name, e-mail address, and credit card number)

LMR_U1_L1



6. Inform the students that you will collect the handouts during the next class in order to assess their understanding of data.



7. To get students thinking about what happens to their data, show the *Terms and Conditions* video: <https://www.youtube.com/watch?v=ZcjtEKNP05c>



8. Engage the students in a whole class discussion about the video, particularly noting:
 - a. What terms in the **privacy** statements were concerning or worrisome?
 - b. Do you read the agreements when you download phone apps?

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were. Be prepared to facilitate a good discussion and to ask probing questions in order for students to elaborate on their thinking so that vague responses such as "we learned about data" can be avoided.

Homework



Students will complete the *Data Diary* handout. When grading the homework, be aware of whether the data really could be collected, and whether the students' ideas about how the data might be used are reasonable. For instance, students will often imagine that there is a "spy" watching them; this is not what we are after. We are after actual instances in which sensors or electronic surveillance records their actions or records information about them. For example, "someone saw me going into the store" is not valid data for this exercise, but "a camera recorded me entering the store" is valid data.

Lesson 2: Stick Figures

Objective:

Students will learn how to observe, record, and organize data.

Materials:

1. *Stick Figures* cutouts (LMR_U1_L2)
Advance preparation required: Needs to be cut into cards (see step 3 in lesson)
2. Poster paper
3. Markers
4. Sticky notes

Vocabulary:

collect, record, organize, representations, variables

Essential Concepts: Data consist of records of particular characteristics of people or objects. Data can be organized in many different ways, and some ways make it easier than others for achieving particular purposes.

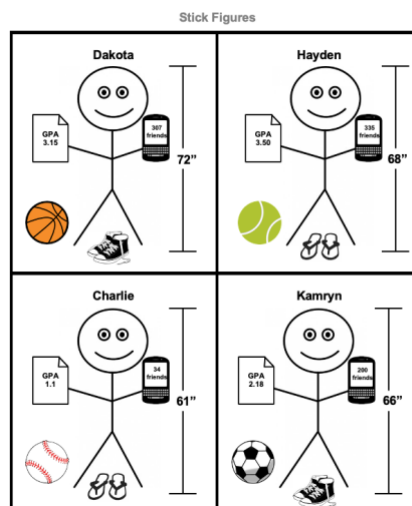
Lesson:



1. Engage students in a *Think-Pair-Share* (see Instructional Strategies) of the *Data Diary* handout that the students completed for homework. Ask them to think about the following questions as they reflect on their data collection homework:
 - How many observations did you make?
 - Where do you leave the most data trails?
 - What could someone learn about you if that person had all of this data?
2. Explain to students that they are going to act as researchers and **collect** data on a strange group of people who appear to be completely two-dimensional. Their goal is to **record** as much information as possible about these people, and to then **organize** the information in any way they choose.
3. Distribute one full set of 8 cards from the *Stick Figures* file (LMR_U1_L2) to each student team.

Advance preparation required:

Print the *Stick Figures* file (LMR_U1_L2). The handout can then be cut into the 8 cards. You will need enough sets of the cards for each student team to share one full set. For example, if there are 5 student teams in a class, then 5 copies of the file will need to be printed so that each team gets all 8 cards.



LMR_U1_L2

4. Every student from the team will select one of the cards from the team's pile of 8 and should record all possible information in their DS journal. Once each student has completed this, the team should come together to share individual findings.
5. Distribute one piece of poster paper and a set of markers to each team. The students will then begin to organize the data from all 8 cards into a visual that they think represents the data. It is important that no guidance is given during this portion of the lesson. Students should be free to come up with their own schema for organizing the data.
6. Display all the posters around the room and allow students to participate in a Gallery Walk (see Instructional Strategies in Teacher Resources) to view other teams' **representations** of the Stick Figure data. For each poster, the teams should write either a comment or a question on a sticky note and add it to the poster to provide feedback for the original team.
7. Afterwards, engage the students in a discussion with the following questions:
 - a. Describe some similarities among the team posters. Were the data organized in similar ways? *Answers will vary by class.*
 - b. Describe some differences among the team posters. How were the data organized differently across teams? *Answers will vary by class.*
 - c. What information was available about the stick figures on each card? *Answer: The person's name, height, GPA, shoe style, sport, and number of friends on social media.*
Note: The names of the stick figures were purposely chosen as gender-neutral, as we never want make assumptions about our data unless it's explicitly stated.
 - d. Which representations made it easy to see what (or who) the objects were that were observed? Which representations made it easy to see whether different stick figures had different characteristics? *Answers will vary by class.*
 - e. Which representation makes it easiest to see which stick figure is tallest? *Answers will vary by class.*
 - f. If you were handed a blank stick figure and knew only the person's name, could you fill in the rest of the information? *Answer: No. You would not know a person's height, GPA, shoe preference, etc. just by knowing their name.*
8. Explain to the students that the general categories of information, such as a person's height, are called **variables**. Variables are simply characteristics of an object or person. As statisticians, we use variable names to organize data into a simplified form so that a computer can read them. This will be discussed further in Lesson 3.



Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Lesson 3: Data Structures

Objective:

Students will learn that data can be represented in rectangular format.

Materials:

1. DS journals (must be available during every lesson)
2. *Stick Figures* cutouts (see Lesson 2)

Vocabulary:

variables, numerical variables, categorical variables, rows, columns, rectangular or spreadsheet format, variability

Essential Concepts: Variables record values that vary. By organizing data into rectangular format, we can easily see the characteristics of observations by reading across a row, or we can see the variability in a variable by reading down the column. Computers can easily process data when it is rectangular format.

Lesson:

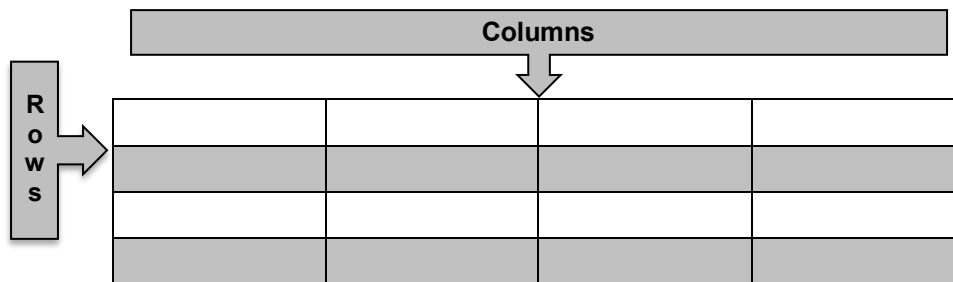
1. Remind students that they briefly learned what **variables** are during the previous lesson. Have students create their own definitions of the term “variables” and share their responses with their teams. Select a few students in the class to share out their definitions and discuss what could be modified (if anything) to create a more complete definition.
2. Using the *Stick Figure* information from Lesson 2, allow the class to come up with a set of variable names that describe the different categories of information. Note that it is best when variable names are short (one to three words). The variable names for the *Stick Figures* data could possibly be:
 - a. Name
 - b. Height
 - c. GPA
 - d. Shoe or Shoe Type
 - e. Sport
 - f. Friends or Number of Friends
3.  Next, have a class discussion about how the values from “Shoe” are different than the values from “Height”.
 - a. The values from “Shoe” are either “sneakers” or “sandals”.

Note: Other terms for these shoes are acceptable – e.g., tennis shoes, flip flops, closed-toe, open-toe, etc.
 - b. The values from “Height” are 72, 68, 61, 66, 65, 61, 67, and 64.
4. Students should notice that the “Shoe” variable consists of categories or groupings, and the “Height” variable consists of numbers. Therefore, we can classify variables into two types: **categorical variables** and **numerical variables**. Typically, categorical variables represent values that have words, while numerical variables represent values that have numbers.

Note: Categorical variables can sometimes be coded as numbers (e.g., our *Stick Figures* variable “Shoe” could have values 0 and 1, where 0 = Sneakers and 1 = Sandals).
5. As a class, determine which variables from the *Stick Figures* data are numerical, and which variables are categorical. The students should create two lists in their DS journals similar to the ones below (the correct classifications are in grey):

<u>Numerical</u>	<u>Categorical</u>
1. <i>Height</i>	1. <i>Name</i>
2. <i>GPA</i>	2. <i>Shoe</i>
3. <i>Friends</i>	3. <i>Sport</i>

6. Explain that although we can understand many different representations of data (as evidenced by the posters from Lesson 2), computers are not as capable. Instead, we need to organize data in a structured way so that a computer can read and interpret them.
7. One way to organize the data is to create a **data table** that consists of **rows** and **columns**. We can define this type of organization as **rectangular format**, or **spreadsheet format**.
8. Display a generic table on the board (see example below) and explain that the columns are the vertical portions of the table, while the rows are the horizontal portions. Another way to think of it is that columns go from top to bottom, and rows go from left to right.



9. Ask students:
 - a. What should each row represent? *Answer: Each row should represent one observation, or one stick figure person in this case.*
 - b. What should each column represent? *Answer: Each column should represent one variable. As you go down a column, all the values represent the same characteristic (e.g., Height).*
10. On the board, draw the following table and have the students copy it into their DS journals (be sure to use variable names agreed upon by the class):

Name	Height	GPA	Shoe	Sport	Friends

11. In teams, students should complete the data table using all 8 of the *Stick Figures* cards. Each row of the table should represent one person on a card.
12. Engage the class in a discussion with the following questions:
 - a. Do any of the people in the data have the same value for a given variable? In other words, does a value appear more than once in a column? Give two examples. *Answers will vary. One example could be that Dakota, Kamryn, Emerson, and London all wear sneakers. Another example could be that Charlie and Jessie are both 61 inches tall.*
 - b. Do any of the people in the data have different values for a given variable? *Answer: Absolutely. There are many instances of this in the data table.*
13. Discuss the term **variability**. As in question (b) above, the values for each variable vary depending on which person we are observing. This shows that the data has variability, and the first step in any investigation is to notice variability. We can see the relationship between the terms **variable** and **variability**. The word “variable” indicates that values vary.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Lesson 4: The Data Cycle

Objective:

Students will learn about the stages of the Data Cycle.

Materials:

1. *The Data Cycle* file (LMR_U1_L4_A)
2. Computer, projector, or board and markers/chalk
3. Printed description of each stage of the Data Cycle (refer to step 3 in the lesson)
4. *The Data Cycle Spinners* handout (LMR_U1_L4_B)
5. RStudio: <https://portal.thinkdataed.org>
6. Article headline: *People Who Order Coffee Black Are More Likely To Be Psychopaths*
7. *Dude Map* found at (open in new tab):
<https://qz.com/316906/the-dude-map-how-american-men-refer-to-their-bros/>
8. *Bros & Dudes Graphics* handout (LMR_U1_L4_C)
9. Sticky notes
10. Poster paper

Vocabulary:

data cycle, statistical questions, data collection, data analysis, data interpretation

Essential Concepts: A statistical investigation consists of cycling through the four stages of the Data Cycle. The term statistical questions encompasses the variety of questions asked during the statistical problem-solving process which support statistical thinking and reasoning. Statistical questions are perhaps the most important because they are challenging to learn and are the types of questions that determine whether an analysis is productive or not. Statistical questions are questions that address variability and are productive in that they motivate data collection, analysis, and interpretation. The Data Collection phase might consist of collecting data through Participatory Sensing or some other means, or it might consist of examining previously collected data to determine the quality of the data for answering the statistical questions. Data Analysis is almost always done on the computer and consists of creating relevant graphics and numerical summaries of the data. Data Interpretation is involved with using the analysis to answer the statistical questions.

Lesson:

If you haven't done so already, create your IDS class and import users.

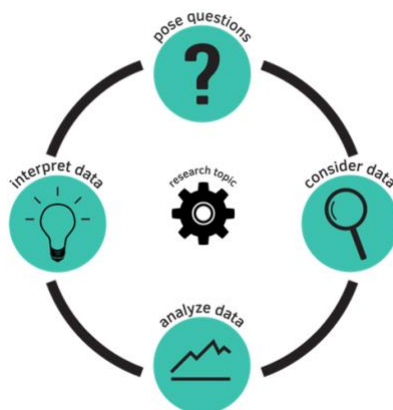


Create
your IDS
class!

- Watch this video to learn how to set up and manage a class:
<https://www.youtube.com/watch?v=M7u0yNRsAtg>
 - Students will need their usernames and passwords in lesson 7.
- If students' first and last names are merged in one column separated by commas, watch this video to learn how to split them into two columns:
<https://www.youtube.com/embed/dtWF291XwzE>

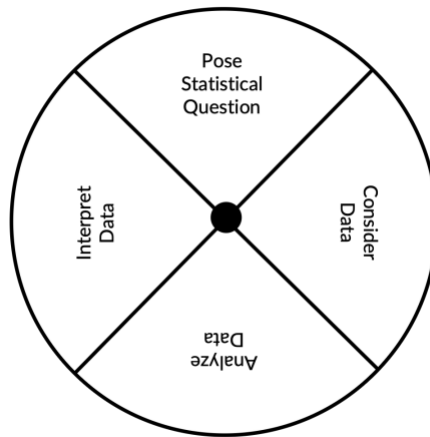
1. During the past few lessons, we have discussed what data are, how to collect and organize them, and how their values can vary. But what do we do with all this data? How can we navigate it and turn it into something useful to us?
2. Inform students that they will be learning about the **Data Cycle** today. The Data Cycle is a guide we can use when learning to think about data. We usually start with posing statistical questions. Display the graphic from *The Data Cycle* file (LMR_U1_L4_A):

The Data Cycle



LMR_U1_L4_A

3. Display the Data Cycle on the board or on a projector and give a brief explanation of the 4 components (listed below).
Note: we will explore each component of the Data Cycle more explicitly throughout the course.
 - a. **Pose Statistical Questions:** Statistical questions are questions that address variability and can be answered with data.
 - b. **Consider Data:** This is the process of observing and recording data, or of examining previously collected data to make sure it meets the needs of the investigation.
 - c. **Analyze Data:** During analysis, tables, graphs, and summaries of the data are produced to help us find patterns and relationships.
 - d. **Interpret Data:** The statistical questions are answered by referring to the tables, graphs, and summaries made in the Data Analysis phase.
4. Almost all statistical investigations begin with statistical questions. There are times when the questions may be given to us, so we might start at the data collection step, but this should ideally be our starting point.
5. As an example, explain that you might ask a person “How old are you?” Although this is a question, it is NOT a statistical question because we are only asking one person so there is no variability in the data. The question “How old are you?” is a survey question that you might ask if you were trying to answer the statistical question “How old are the students in my school?” We would need to collect data to answer the question and we would expect student ages to vary.
6. To help students get a firm understanding of the Data Cycle and how each component is connected, they will participate in a *Four Corners* strategy (see Instructional Strategies in Teacher Resources). Write down the name of each stage in the Data Cycle on a sheet of paper and include the description of that particular stage (see step 3 above for descriptions). Then tape each sheet on a different corner of your room.
7. Explain to the students that you are going to display different artifacts from statistical investigations on the projector. For each artifact, they will move to the corner of the room they feel that artifact represents (posing a statistical question, consider data, analyze data, interpret data). If you have limited space in your classroom or for students that cannot physically participate, you may consider printing LMR_U1_L4_B. Students can participate by pointing to the spinner.



LMR_U1_L4_B

8. Once students have chosen a corner of the room (stage of the Data Cycle) they will discuss the following with their classmates in that same corner:
 - a. What part of the data cycle does the artifact represent (posing a statistical question, consider data, analyze data, interpret data)? Why do we think that?
 - b. What questions or wonderings do we have about the artifact?
9. Allow each group time to discuss the questions and have one member from each team (corner) share the answers to the questions. This activity is not about having a correct answer. It is about having students begin to think critically about statistical artifacts that they are constantly consuming. Data are encountered through visualizations, reports from scientific studies, and journalists' articles and websites. This activity is meant to begin to develop students' statistical habits of mind.
10. Artifact 1: Spreadsheet of the CDC data.

Note: The CDC uses sex as a binary variable, which may bring up classroom conversations about sex and gender. While medical science has moved to a non-binary definition of gender and also allows for more than two sexes, this notion has not yet made its way into many official agencies. We feel students need experience working with data from these government agencies since they play a very important role in the US economic and cultural life.

	age	sex	grade	height	weight	seat_belt	drive_text
1	16 years old	Female	11th grade	1.57	70.31	Always	Did not drive
2	15 years old	Male	NA	NA	NA	Most of the time	Did not drive
3	15 years old	Male	10th grade	1.70	61.69	Always	Did not drive
4	15 years old	Male	10th grade	1.75	60.33	Most of the time	0 days
5	14 years old	Male	10th grade	1.90	88.91	Always	Did not drive
6	15 years old	Female	10th grade	NA	NA	Always	Did not drive
7	15 years old	Male	10th grade	1.73	84.82	Most of the time	Did not drive
8	14 years old	Male	10th grade	1.60	49.90	Always	Did not drive
9	15 years old	Male	10th grade	1.78	61.24	Most of the time	0 days
10	15 years old	Female	10th grade	1.47	54.43	Most of the time	0 days

Note: You can display this spreadsheet below using RStudio by running the following commands:
`data(cdc)` `View(cdc)`

- a. What part of the data cycle does the artifact represent (posing a statistical question, consider data, analyze data, interpret data)? Why do we think that? *Answers will vary.*
- b. What questions or wonderings do we have about the artifact? *Students should begin developing statistical habits of mind. They should be interrogating the data by asking questions such as: Who is this data about? What was the purpose of collecting the data? What was the survey question asked to collect the data?*



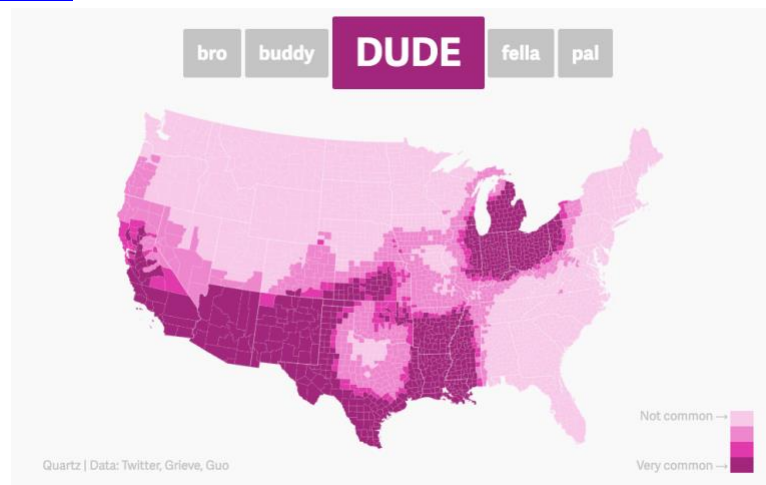
11. Artifact 2: Headline from Huffington Post *People Who Order Coffee Black Are More Likely To Be Psychopaths*.



- What part of the data cycle does the artifact represent (posing a statistical question, consider data, analyze data, interpret data)? Why do we think that? *Answers will vary.*
- What questions or wonderings do we have about the artifact? *Answer: Students should be interrogating this headline with questions like: What type of study was this? Who funded the study? What was the purpose of the study? How was the variable measured?*



12. Artifact 3: The Dude map found at: <https://qz.com/316906/the-dude-map-how-american-men-refer-to-their-bros/>



- What part of the data cycle does the artifact represent (posing a statistical question, consider data, analyze data, interpret data)? Why do we think that? *Answers will vary.*
 - What questions or wonderings do we have about the artifact? *Students should be asking questions like: What was the purpose of this study? What variables were measured and how were they measured?*
13. Inform students that the *Dude Map* was created for the Quartz website by Nikhil Sonnad as a data visualization. He collected the data via Twitter. The graphic shows how common the terms: bro, buddy, dude, fella, and pal are when referring to friends throughout the United States.
14. Ask students to return to their seats, take out their DS journal and make a sketch of the Data Cycle, making sure to include the names of the four stages (Pose statistical question, consider data, analyze data, interpret data).
15. Ask students to write *Dude Map* under the analyze data of their data cycle and the information about where the data came from (see #13) in the consider data part of their Data Cycle sketch.
16. Have each team discuss a possible statistical question that could be answered using the *Dude Map* graphic. Have the Recorder/Reporter write the question on a sticky note and the resource manager bring it up to the board.
17. Lead a class discussion around the statistical questions the student teams created, and as a class, choose one to write down as an example. *Example: Where in the United States is the term dude more common to use when referring to a friend?*



18. Allow the teams to work together to answer the statistical question. Ask the Recorder/Reporter to share their team's interpretation. Have students write down the answer that resonated the most with them under the interpret part of the data cycle.

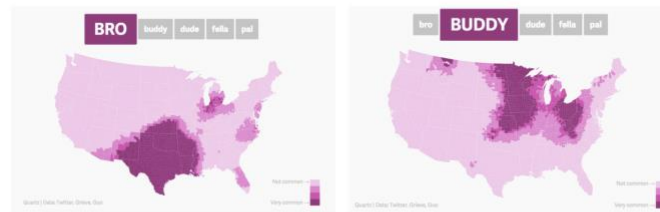


19. Assign ONE of the pages from the *Bros & Dudes Graphics* handout (LMR_U1_L4_C) to each team. There are 10 different versions of word pairings (10 combinations of 2 words chosen from the 5 options), so multiple teams will have the same graphic if there are more than 10 teams in a class.

Team Members: _____

Date: _____

The dude map: How Americans refer to their bros
(<http://qz.com/316906/the-dude-map-how-american-men-refer-to-their-bros/>)
Created by Nikhil Sonnad, December 23, 2014, for Quartz website.
The data originally came from accessing Twitter feeds.



What statistical question(s) would you ask given the above graphics? Give 2 examples.

(1) _____

LMR_U1_L4_C

20. The goal of this activity is for each team to complete a full statistical investigation with the *Bros & Dudes Graphics* assigned to them. Tell the teams that they need to create a Data Cycle poster using their assigned graphic for the analyze stage. The cycle should be clearly labeled and have appropriate responses for each of the 4 components.

For example, a team given the “Bro” and “Buddy” graphics might come up with the following questions: Which region of the US is most likely to use the term “Bro” when referring to a friend? Do the coastal areas prefer different terms than the Midwest? Is there a difference between the northern states versus southern states?

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students reformulate any statistical questions generated by their team about the *Bros & Dudes Graphics* that could not be answered so that they can be answered.

Lesson 5: So Many Questions [The Data Cycle: Pose Questions]

Objective:

Students will learn the features of a good statistical question.

Materials:

1. Post-its
2. *Statistical Questions Scenarios* handout (LMR_U1_L5)

Vocabulary:

statistical questions (summary, comparison, association)

Essential Concepts: Statistical questions typically begin with a vague general question, then develop into a precise question. The process of developing or creating a good statistical question is iterative and requires time and effort to get right. In her 2021 paper, *What Makes a Good Statistical Question*, Dr. Pip Arnold identified the following as features of a good statistical question:

1. The variable(s) of interest is/are clear
2. The group or population we are interested in is clear
3. The question can be answered with the data
4. The question asks about the whole group, not an individual
5. The intention is clear (e.g., summary, comparison, association, time series)
6. The question is one that is worth investigating, is interesting, and has a purpose

Note: You may see statistical questions referred to as statistical investigative questions, which may be used interchangeably by statisticians and data scientists.

Lesson:

1. Entrance Slip (see Instructional Strategies in Teacher Resources): Each student should submit a ticket that displays the 4 components of the Data Cycle.
2. Inform students that they will learn about what makes a good statistical question. Ask them to recall the definition of a statistical question:
Statistical questions are questions that address variability and can be answered with data. They are questions we ask of the data. A good way to determine this is to ask: *Do we need to see the data to answer the question?*
3. Remind the students of the two questions from the previous lesson, noting that one of the questions was a statistical question, and the other was not:
 - a. How old am I?
 - b. How old are the students in my school?
4. In pairs, ask students to analyze each question using the definition of a statistical question and come to an agreement about which one is a statistical question.
5. Using Agree/Disagree (see Instructional Strategies in Teacher Resources), ask a pair of students for their results. Discuss why the first question **IS NOT** a statistical question (*there is only one possible value so there is no variability in the data*) and why the second question **IS** a statistical question (*not all students are the same age. The ages vary, so there is variability in the data*).
6. Ask students to think of the data they collected about the stick figures (name, GPA, friends, sport, height, shoe). Inform them that the researchers used the following survey questions to collect the data:
 - a. What is your name?
 - b. What is your GPA?
 - c. How many friends do you have?
 - d. What sport do you play?



- e. How tall are you in inches?
- f. What type of shoe do you mostly wear?

Survey questions are questions we ask to get the data.

7. Tell students that it is important to know exactly what survey questions were asked to collect the data before asking statistical questions. For example, we saw an image of a ball next to each stick figure but we don't know if that represents a sport they like to watch, their favorite sport, or a team they are on.
8. In teams, ask students to create statistical questions that could be answered using the data collected about the stick figures. Introduce the sentence stem "I wonder..." to help students get started. Have the Recorder/Reporter record the questions on post-its.
9. Ask the teams to identify which variable(s) each question is investigating by having them circle the variable name(s) within their statistical questions.
10. Have the Task Manager organize their group's statistical questions on the board, placing statistical questions that incorporate only one variable on the left-hand side of the board and statistical questions that incorporate two or more variables on the right-hand side.

Note: This sorting activity will help students begin to distinguish between different types of statistical questions.

- **Summary statistical questions** are questions about a single variable.
- **Comparison statistical questions** compare a numerical variable across groups.
- **Association statistical questions** look for a relationship between paired numerical or paired categorical variables.

11. As a class, begin the process of transforming some of the summary statistical questions so that they have all of the features of a good statistical question. Here is an example to get you started:

Initial statistical question: *I wonder... Who has the most friends?*

Feature	Explanation
The variable(s) of interest is/are clear	Yes. The variable of interest is the number of friends.
The group or population we are interested in is clear	No. Teacher should ask: "Who did the researchers want to learn about?" <i>These stick figures.</i>
The question can be answered with the data	Yes. The researchers collected data on the number of friends.
The question asks about the whole group, not an individual or portion of the group	No. This question is about an individual stick figure. Teacher should ask: "How can we reword the question to include the whole group?" <i>How many friends do ... have?</i>
The intention is clear (e.g., summary, comparison, association)	It is clear that this is a summary statistical question (single variable), specifically the number of friends.
The question is one that is worth investigating, is interesting, and has a purpose	For students, this might be something interesting.

Reworded statistical questions after going through the criteria: *How many friends does this group of stick figures have?*

12. As a class, apply the same process to a few of the comparison and association questions.

Initial statistical question: *I wonder... Does someone who plays a specific sport have more friends?*

Feature	Explanation
The variable(s) of interest is/are clear	Yes. This seems to be a comparison statistical question comparing the number of friends within the sport played.
The group or population we are interested in is clear	No. Teacher should ask: "Who did the researchers want to learn about?" <i>These stick figures.</i>
The question can be answered with the data	Yes. The researchers collected data on the number of friends and the sport the stick figures played.
The question asks about the whole group, not an individual or portion of the group	No. The word <i>someone</i> gives the impression that we are interested in one observation.
The intention is clear (e.g., summary, comparison, association)	The intent is somewhat clear. This seems to be a comparison statistical question between the sport the stick figures played and the number of friends each stick figure had. Teacher should ask: "What is the variable that is being compared? Which groups within the sport variable are you comparing (all groups, specific groups)?"
The question is one that is worth investigating, is interesting, and has a purpose	For students, this might be something interesting.

Reworded statistical question after going through the criteria:

Is there a difference in the typical number of friends the stick figures have based on the sport they play?

Do the stick figures who play soccer tend to have more friends than the stick figures who play tennis?



13. Using the criteria of a good statistical question, student teams will go back and modify their statistical questions. Facilitators will ensure the team goes through the criteria for each statistical question. Task Managers will encourage everyone to contribute. Resource Managers will ensure all materials are easily accessible for recording and reporting. Recorders in each team will capture team members' responses while the teacher circulates the room to check for understanding.



14. Ask the Reporters of selected teams to share their revised statistical question(s). Students in the audience will listen to the presentations and provide feedback about each team's statistical question(s). Be sure to discuss disagreements before moving on to different questions.

Additional practice: Students can practice identifying types of statistical questions and rewriting them to be good statistical questions using the *Statistical Questions Scenarios* handout (LMR_U1_L5). This can be done in groups during class or individually for homework.

Statistical Questions Scenarios

Scenario #1

Diane Ma, data scientist for the Lakers Basketball team, is analyzing data to help players win games. She's come up with some initial statistical questions but would like to transform them into good statistical questions.

1. Identify the type of statistical question (summary, comparison, or association) and then go through the process of transforming it into a good statistical question.

I wonder:

- a. Which player scores the most points?
- b. Does a certain position (center, power forward, small forward, point guard, shooting guard) score more points?

2. Write another statistical question you can ask that might help to win basketball games.

Scenario #2

The president comes to visit, and you want to take the opportunity to ask him some good statistical questions. Your friends offer the following initial statistical questions to get you started.

1. Identify the type of statistical question (summary, comparison, or association) and then go through the process of transforming it into a good statistical question.

I wonder:

- a. Is one state's spending more than another?
- b. Does someone in a big family spend a lot on healthcare?

2. Write another statistical question that you could ask the president.

Scenario #3

LMR_U1_L5

15. Inform students that in the next lesson, they will begin using the Data Cycle to learn about their food habits. To prepare for this, students should begin collecting the "Nutrition Facts" labels from foods/snacks they typically eat.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Ask students to bring at least 2 cutouts of the "Nutrition Facts" labels of the snacks they typically eat (e.g., chips, yogurt, blended drinks, etc.).

Note: An alternative to collecting "Nutrition Facts" labels is to print them from an online source and bring the printouts to class.

Lesson 6: What Do I Eat? [The Data Cycle: Consider Data]

Objective:

Students will collect data using paper and pencil to understand the challenges of organizing, storing, and sharing data. They will learn that there must be an agreement about the variables that need to be recorded in order to attain consistency.

Materials:

1. Video: Jamie Oliver's *Food Revolution* found at:
<https://youtu.be/IOvYwqkktM>
Note: If the video is unavailable, search for "Jamie Oliver's Food Revolution What's In a Sundae"
The video should be 5-6 minutes in length.
2. Nutrition Facts labels or pictures (collected previously by students)
Note: If needed, use *Nutrition Facts Cutouts* handout (LMR_U1_L6_A)
3. *Food Habits Data Collection* handout (LMR_U1_L6_B)

Vocabulary:

dataset(s)

Essential Concepts: After raising statistical questions, we examine and record data to see if the questions are appropriate.

Lesson:



1. Tell students that today's lesson will focus on the Data Collection component of the Data Cycle.
2. To motivate this, the students will watch a short video of an episode of Jamie Oliver's show titled *Food Revolution* found at: <https://youtu.be/IOvYwqkktM>. This video was recorded at a Los Angeles high school.
 - a. As the students watch the video, they should use their DS journals to write down their comments and/or reactions to what they see and hear and be ready to share out.
 - b. After sharing out some of their responses, have the students respond to the following question in their DS journals: **Why should I care about what I eat?**
 - c. Student teams will share their reactions and responses by engaging in a Silent Discussion (see Instructional Strategies in Teacher Resources).
3. Have students recall the *Stick Figures* activity from Lesson 2. During that activity, they collected data about other people. But today, they are going to be collecting data about themselves and the foods they eat.
4. Students should have Nutrition Facts labels available from food/snacks they consumed at home between the previous lesson and today.
Note: If some students forgot to bring any, then you can pass out some of the *Nutrition Facts Cutouts* (LMR_U1_L6_A) for them to use instead.

Nutrition Facts Cutouts

Fruits & Nuts

Bananas, raw

Serving size: 1 cup, mashed (225g)

Nutrition Facts	
Serving Size 225 g	
Amount Per Serving	
Calories 200	Calories from Fat 6
	% Daily Value*
Total Fat 1g	1%
Saturated Fat 0g	1%
Trans Fat	
Cholesterol 0mg	0%
Sodium 2mg	0%
Total Carbohydrate 51g	17%
Dietary Fiber 6g	23%
Sugars 28g	
Protein 2g	
Vitamin A	3% • Vitamin C 33%
Calcium	1% • Iron 3%

*Percent Daily Values are based on a diet of other people's secrets. Your daily values may be higher or lower depending on your calorie needs.

Apples, raw, with skin

Serving size: 1 cup, quartered or chopped

Nutrition Facts	
Serving Size 125 g	
Amount Per Serving	
Calories 65	Calories from Fat 2
	% Daily Value*
Total Fat 0g	0%
Saturated Fat 0g	0%
Trans Fat	
Cholesterol 0mg	0%
Sodium 1mg	0%
Total Carbohydrate 17g	6%
Dietary Fiber 3g	12%
Sugars 13g	
Protein 0g	
Vitamin A	1% • Vitamin C 10%
Calcium	1% • Iron 1%

*Percent Daily Values are based on a diet of other people's secrets. Your daily values may be higher or lower depending on your calorie needs.

Oranges, raw

Serving size: 1 cup, sections (180g)

Nutrition Facts	
Serving Size 180 g	
Amount Per Serving	
Calories 85	Calories from Fat 2
	% Daily Value*
Total Fat 0g	0%
Saturated Fat 0g	0%
Trans Fat	
Cholesterol 0mg	0%
Sodium 0mg	0%
Total Carbohydrate 21g	7%
Dietary Fiber 4g	17%
Sugars 17g	
Protein 2g	
Vitamin A	8% • Vitamin C 160%
Calcium	7% • Iron 1%

*Percent Daily Values are based on a diet of other people's secrets. Your daily values may be higher or lower depending on your calorie needs.

LMR_U1_L6_A

- For 3-5 minutes, allow students to collect any data they can from the label and record it in their DS journals. This should be done individually.
- Once they have collected their facts, ask students to compare and contrast their data with their team members. They need to respond to:
 - How are their **datasets** similar?
 - How are their **datasets** different?
- Gather the students as a whole group and ask them to share out the similarities and differences they discussed. Be sure to draw responses that show that while some facts collected were the same, there were others that were collected by some students and not by others. Also point to differences in the variables collected and the data structure used.
 - Ask students to engage in the following individual Quickwrite (see Instructional Strategies in Teacher Resources): How can the data you just gathered be quickly displayed and easily read?
 - Distribute the *Food Habits Data Collection* handout (LMR_U1_L6_B). They can use their own label for the first observation and then fill out the rest of the table with the food/snacks they consume during the Food Habits campaign.

Name: _____

Date: _____

Food Habits Campaign Data Collection

What is the name of the snack?	When did you eat the snack? (morning, afternoon, evening, night)	Is the snack salty or sweet? (Salty, Sweet)	How healthy is the snack? (1=Very unhealthy, 5=Very healthy)	How many calories per serving?	How many grams of protein per serving?	How many grams of sugar per serving?	How many milligrams of sodium per serving?	How many ingredients are in the snack?	Why are you eating the snack? (availability, craving, emotional, energy, hungry/thirsty, social, other)	How much does the snack cost (in dollars)? (\$0 to < \$1, \$1 to < \$3, \$3 to < \$7, \$7 or more)

LMR_U1_L6_B

- Once they are finished, in pairs, ask students to give a one-word identifier to each variable. For example: "What's the name of your snack?" = Name
- Share the one-word variable identifiers with the class by conducting a quick team Whip Around (see Instructional Strategies in Teacher Resources).

10. For homework, students will begin to formulate statistical questions based on their *Food Habits* data.
11. Inform students that they are permitted to bring mobile devices to the next class.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Ask students to examine the data in their *Food Habits Data Collection* handout (LMR_U1_L6_B) and to generate three statistical questions (summary, comparison, association) that they think can be answered by the Food Habits data.

Students may bring their mobile devices to the next class for data collection purposes.

Lesson 7: Setting the Stage [The Data Cycle: Collect Data]

Objective:

Students will begin to collect and record data to learn more about their own eating habits, as well as those of their classmates. They will learn about data that is collected by a Participatory Sensing campaign, and also about privacy issues and photo ethics when collecting and sharing data.

Materials:

1. Students' own mobile devices (smartphone or tablet compatible with iOS or Android)
2. Access to App Store or Google Play Store in student devices to download IDS ThinkData Ed app
3. Login information (username and password) for each student – **generated and ready for distribution prior to the lesson**
4. *Food Habits Campaign* guidelines (LMR_U1_Campaign_Food_Habits)
5. **Optional:** You can show this video to demonstrate to students how to take a survey using the browser: <https://www.youtube.com/embed/hksbF5QWY2E> or this video for the app: <https://www.youtube.com/watch?v=tGK3Am6-BQo>

Vocabulary:

Participatory Sensing, campaign, surveys, images, GPS, ethics, photo ethics

Essential Concepts: In Participatory Sensing, we humans behave as if we are robot sensors, collecting data whenever a “trigger” event occurs. Our ability to learn about the patterns in our life through these data depends on our being reliable data collectors.

Lesson:



Before students begin collecting data, ensure that:

- The campaign settings are set to your preference.
- Watch this video to learn about campaign settings: <https://www.youtube.com/embed/ZTGpbgqlc0>

1. Become familiar with the *Food Habits Campaign* guidelines (shown at the end of this lesson), especially the big questions found under “1. The Issues”, to help guide students during the campaign (see Campaign Guidelines in Teacher Resources).
2. Distribute the usernames and passwords to student team leaders (make sure safeguards are in place so that only the owner of the username and password is able to see this information). Ask team leaders to distribute their team’s information when you are ready to download the IDS ThinkData Ed app.
3. Ask students to think about the Nutritional Facts labels from which they collected data in the previous lesson, and answer the following in their DS journals:
 - a. What questions would you want answered about eating habits?
 - b. What can you do to find out about your own eating habits?
4. Ask students to refer back to their reactions and comments from Jamie Oliver’s video and have them ponder the question: **What are we really eating?**
5. Over the next 9 days, they will engage in a **Participatory Sensing campaign** in which they will act as human sensors to collect data about themselves. The data collected will be used to analyze their classmates’ and their own snacking habits.
6. For this unit, they will collect data about every snack they eat.

Note: Students should **NOT** collect data for full meals like breakfast, lunch, or dinner. Data should only be collected for anything eaten in between meals, like fruit, chips, cookies, nuts, sodas, etc.



7. Ask students why it makes sense to study snacks specifically. Brainstorm some questions that could be answered using the snack-only data that would be hard to answer if the data included meals as well.
8. Inform students that they will be taking part in a specific data collection method known as **Participatory Sensing** via a mobile application. This application can gather data via **surveys**, **images**, and **GPS** tracking. For students who choose to share their location, their precise location will not appear on the shared IDS map. In the interest of protecting student privacy, our GPS is configured to show a generalized location tag, not a specific address.
9. Make it clear to students that the reason they are collecting the data is to learn more about themselves and their classmates, NOT to provide data for an external data collection team. Students occasionally have the misconception that when they use the Participatory Sensing app, they are providing data to external researchers, such as UCLA.
10. Inform students that they are now going to engage in their own first Participatory Sensing data collection experience, in which they will collect their own data using a smart device. Depending on the device, there are 3 different options available:
 - a. **Android:** A native Android application called “IDS ThinkData Ed” is available from the Google Play Store.
 - b. **iOS (Apple devices):** The mobile application called “IDS ThinkData Ed” is available from the iOS App Store.
 - c. **No mobile device - browser-based version:** For students that do not have a mobile device (or an unsupported device, such as a Windows phone or Blackberry), a browser-based version to perform data collection is available at <https://portal.thinkdataed.org>
Click on the **Survey Taking icon** on the page.
11. Once students have downloaded the app or have found the website, ask team leaders to distribute the login information. Students will need to keep this information in a safe place for the entire duration of this course. **Emphasize the importance of keeping their username and password confidential.** When students receive their login information, they can log in to the app. If students have trouble with their logins, the teacher has the ability to reset a student's password.
12. Once logged into the app or the browser-based version, students will see the **campaigns** in which they will participate. They will then select the campaign by tapping the name of the campaign. If no campaigns are visible, ask them to tap the refresh option, located on the top right-hand side of the screen.
13. Using one of their nutrition facts cutouts or pictures, ask students to complete their first survey by going through the questions in the app.
14. After every student has had the opportunity to complete at least one survey, ask students the meaning of the word **ethics**. For this course, they will need to understand **photo ethics**. They may NOT take pictures of any person's identifying features such as faces, hair, hands, tattoos, etc. For this campaign, they may only take pictures of their snacks and/or the nutrition facts labels.
Note to teacher: Inspect students' data collection photos throughout the data collection period and before each data collection, monitoring to ensure that no inappropriate images are shared. If you believe a photo is inappropriate, please delete the data entry immediately.
15. Setting reminders: The IDS ThinkData Ed app has a reminder feature to help students in their data collection journey. Show students that they can set reminders directly on the app by tapping the menu button on the top left-hand side of the screen and selecting **Notifications** from the menu.
16. Data collection norms: Ask students how many snacks they think they eat a day. From this, come up with an approximate number of surveys they think each student should complete during the data collection period (days 7 through 15).

17. Inform students that you will be monitoring their data collection to make sure that everyone is submitting surveys regularly.
18. Go over the previous day's homework. Ask the Facilitator from each table group to conduct a round robin during which each team member shares their statistical questions. The Recorder/Reporter will select and share the team's statistical questions (one of each type; summary, comparison, association) with the class.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

For the next 9 days, students will collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Campaign Guidelines – Food Habits

1. The Issue:

Although we might take its existence for granted today, the Nutrition Facts label was not always required to be on food packages. It wasn't until 1990 that the Nutrition Labeling Education Act mandated food companies to provide information on food labels to help consumers make wiser choices about what they eat. This raises some interesting questions:

- 1) Does knowing nutritional information about my snacks help me change my habits?
- 2) What is my snacking pattern?
- 3) Do I tend to eat healthy? How do I compare to my class? How does my class compare to the rest of the country?

2. Objectives:

Upon completing this campaign, students will have the enduring understanding that interpreting graphs provides useful information about the real world as viewed through the data represented in the graphs. We can explore the relationship between two variables, and if there is a relationship, if it is driven by the change in the independent variable, x , which causes a change in the dependent variable, y .

3. Survey Questions: (students will enter data about the snacks they consume):

Consider Data: Before students submit a survey for their first snack, a class consensus of the meaning of the variables must be reached so that proper analysis and interpretations can be made. Two examples are listed below:

when - If students have different definitions of "evening", it might make it hard to compare snacking patterns across students. As a class, come to a consensus about what time intervals are considered *morning*, *afternoon*, *evening* and *night*.

cost - If a student has a bowl of cereal as a snack, are they going to include the cost of the entire box or are they going to calculate the unit cost for one serving? This needs to be a class decision.

Prompt	Variable	Data Type
What's the name of your snack?	name	text
When did you eat the snack?	when	categorical morning afternoon evening night
Is your snack salty or sweet?	salty_sweet	categorical Salty Sweet
How healthy is the snack? (1 = Very unhealthy, 5 = Very healthy)	healthy_level	numerical 1 2 3 4 5
How many calories per serving?	calories	numerical
How many grams of protein per serving?	protein	numerical
How many grams of sugar per serving?	sugar	numerical
How many milligrams of sodium per serving?	sodium	numerical
How many ingredients are in the snack?	ingredients	numerical
Why are you eating the snack?	why	categorical availability

		craving emotional energy hungry/thirsty social other
How much does the snack cost (in dollars)?	cost	categorical \$0 to < \$1 \$1 to < \$3 \$3 to < \$7 \$7 or more
Take a picture (optional).	snack_image	photo
AUTOMATIC	location	lat, long
AUTOMATIC	time	time
AUTOMATIC	date	date
AUTOMATIC	user	user id

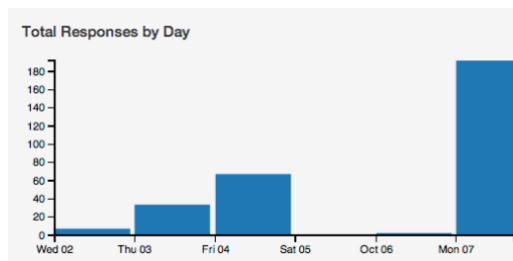
When should you take the survey? If possible, take the survey every time you eat a snack or at the end of the day. Reminders can be set to ensure survey completion.

How long should the campaign last? About nine days. Ideally, two of these days will include a weekend.

4. Motivation:

As a class, come to an agreement about how many surveys each student should submit. All students should submit roughly the same number of surveys, and each student should submit at least four surveys. After the first day, use the Campaign Monitoring tool to see who has collected data. After two to three days, direct students' attention to the Total Responses by Day plot and comment on any patterns. For example, if they see a plot like the one below, ask "What story does this tell us about our data collection?"

Story: They collected a lot of data together in class. Data collection increased every day from Wednesday to Friday. There was little to no data collection over the weekend. Data collection peaked on Monday – there were over 180 responses!



Discuss data collection issues. What makes it hard? Does this affect the quality of data? What sort of snacks are you less likely to enter?

5. Technical Analysis:

Students will use the Dashboard and Plot App as well as RStudio.

6. Guiding Questions:

- 1) At what time of day do we eat the healthiest snacks?
- 2) When did you snack? How does this compare to the rest of the class?
- 3) Typically, how healthy were your snacks? How does this compare to the class as a whole?
- 4) What were some of the characteristics of healthy snacks? What about unhealthy snacks?

7. Report:

Students will complete a practicum in which they answer a statistical question based on the Food Habits data collected.

Visualizing Data

Instructional Days: 14

Enduring Understandings

Data collection methods affect what we can know about the real world. Visual representations help tell stories with data. Distributions of numerical and categorical variables help describe variability in the data. Technology and computers allow us to visualize complex relationships in data.

Engagement

Students will view the video called *The Value of Data Visualization* to help them understand the importance of graphical representations of data. Discussion questions will allow students to begin to think about how they would want to see a dataset visualized. The video can be found at:

<https://www.youtube.com/watch?v=xekEXMOVonc>.

Learning Objectives

Statistical/Mathematical:

S-ID 1: Represent data with plots on the real number line (dotplots, histograms, bar plots, and boxplots).

S-ID 3: Interpret differences in shape, center, and spread in the context of the data sets, accounting for possible effects of extreme data points (outliers).

S-ID 6: Represent data on two quantitative variables on a scatterplot and describe how the variables are related.

Data Science:

Create visualizations with data. Learn the difference between plots used for categorical and numerical variables. Interpret and understand graphs of distributions for numerical and categorical variables.

Applied Computational Thinking Using RStudio:

- Learn to download, load, upload, and work with data using RStudio syntax and structure.
- Create appropriate graphical displays of data.
- Differentiate between observations and variables.
- Learn to use objects, functions, and assignments.

Real-World Connections:

Students will continue to understand that data on its own is just collected; but once interpreted, it can lead to discoveries or understandings.

Language Objectives

1. Students will describe orally and in writing the center, shape, and spread of distributions.
2. Students will discuss the appropriateness of dotplots, bargraphs, histograms, and scatterplots for various data analyses and the pros and cons of each.
3. Students will compare and contrast their hand-drawn data visualizations to computer-generated RStudio/Posit Cloud data visualizations.

Data File or Data Collection Method

Data Files:

1. Students' *Food Habits* Campaign Data
2. Centers for Disease Control and Prevention (CDC): `data(cdc)`

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Lesson 8: Tangible Plots [The Data Cycle: Analyze Data]

Objective:

Students will learn how distributions help us organize and visualize data values, and that the shapes of the distributions give us information about the variability of the data.

Materials:

1. Computer and projector for Campaign Monitoring
2. Video: *Value of Data Visualization* found at: <https://www.youtube.com/watch?v=xekEXM0Vonc>
3. Nutrition facts labels or pictures (collected previously by students)
4. *Food Habits Data Collection* handout (from Lesson 6, LMR_U1_L6_B)
5. 3 pieces of tape per student
6. Poster paper
7. Dot stickers or sticky notes
8. *Tangible Plot* handout (LMR_U1_L8)

Vocabulary:

x-axis, y-axis, visualization, frequency, data points, minimum, maximum, distribution, range, dotplot

Essential Concepts: Distributions organize data for us by telling us (a) which values of a variable were observed, and (b) how many times the values were observed (their frequency).

Lesson:

In preparation for this lesson, ensure that:



- You watch this video on the Campaign Monitoring Tool: https://www.youtube.com/embed/EV-uEh_OogM
- Each day, preview student responses to ensure appropriate inputs. Watch this video to learn how to manage student responses: <https://www.youtube.com/embed/RLIOoLhakg8>

1. Food Habits Campaign Data Collection Monitoring:

- a. Display the IDS Campaign Monitoring Tool, found at <https://portal.thinkdataed.org>. Click on **Campaign Monitor** and sign in.
- b. Inform students that you will be monitoring their data collection. This is a good opportunity for teachers to remind students that if their data are not shared, they cannot be used in analysis.
 - i. See *User List* and sort by *Total*. Ask: Who has collected the most data so far?
 - ii. Click on the pie chart. Ask: How many active users are there? How many inactive users are there?
 - iii. See *Total Responses*. How many responses have been submitted?
 - iv. Using TPS, ask students to think about what they can do to increase their data collection.

2. Inform students that today they will be learning how to visualize their data.

3. Show the *Value of Data Visualization* video at <https://www.youtube.com/watch?v=xekEXM0Vonc>, which describes the importance of graphical representations of data. As they watch the video, students should respond to the following in their DS journals:
 - a. What is data visualization? *Answer: Data visualization is the graphical representation of data or information.*
 - b. List one example of how visualization can be used to increase data comprehension. *Answers will vary. Some examples are graphs, maps, charts, and infographics.*





4. Have a whole class discussion regarding the video's last statement: "Your message is only as good as your ability to share it." Ask students:
 - a. What does this statement mean? *Answers will vary.*
 - b. What makes a good message in terms of data and visualizations? *Answers will vary.*
5. Have students take out their nutrition facts labels or pictures, and also their *Food Habits Data Collection* handout (from lesson 6).
6. On poster paper, make the first quadrant of a coordinate plane. Leave the labels for the **x-axis** and **y-axis** blank for now (see step 10 below).
7. Distribute 3 pieces of tape to each student. Make sure they fold each piece of tape to make two sticky sides. Have each student tape one sticky side to the back of each label and ask them to have the labels ready to tape onto the poster paper.
8. As a class, ask students to select 2 numerical variables and 1 categorical variable from the *Food Habits Data Collection* handout whose data they would like to see in a **visualization**, which is a picture of the data. For example, students may vote to see a visualization of the following numerical variables: calories per serving, total fat per serving; categorical variable: salty_sweet
9. Once students select the variables, inform them that they will be creating a plot with the nutrition facts labels for each of the variables they selected.
10. Create a bargraph of the categorical variable chosen by the students. Begin by showing students how to clearly label the x-axis with the categories. For instance, if salty_sweet is the variable, ask students to identify the categories for that variable. Then mark the y-axis with the label **frequency**, which simply means the number of times an outcome occurs. Do not put tick-marks on the y-axis. The frequency will be measured by the number of labels plotted.
11. Have students come up and place their nutrition fact label above the category that describes their snack. Have students stack their nutrition label so that is easy to calculate the frequency. Once all the labels have been placed, create bars with the appropriate height (frequency) for each category. Make sure to leave spaces between the bars, and that bars are the same width.
12. Ask students to respond to the following questions in their DS journals:
 - a. How many **data points** does this distribution have? Why? *Answers will vary by class.*
 - b. What information is this visualization telling us about [insert categorical variable name] in the snacks we consume? *Answers will vary by class.*
13. Use another piece of poster paper to create a distribution for the first numerical variable chosen by the students.
14. Create a dotplot of the first numerical variable chosen by the students. Begin by showing students how to clearly label the x-axis. For instance, if calories per serving is the variable, ask students for the range of values for calories per serving and determine the **minimum** and **maximum** values for the dataset. Clearly label the x-axis with adequate intervals and the variable's name. Then mark the y-axis with the label **frequency**, which simply means the number of times a value occurs. Do not put tick-marks on the y-axis. The frequency will be measured by the number of labels plotted.
15. For each value in the dataset, put a nutrition facts label above that value on the x-axis. When a value occurs more than once, stack the nutrition facts labels. For example, if there are three nutrition facts labels with 120 calories per serving in the dataset, there will be three nutrition facts labels above the 120 mark on the x-axis.
16. Once all the labels have been placed, ask students to observe the **distribution** of the data in the dotplot. Ask students to respond to the following questions in their DS journals:
 - a. What are the minimum and the maximum values of the dataset? *Answers will vary by class.*



- b. **The range** is the largest value minus the smallest value. It is one way of measuring the variability of a variable. What is the range, and why do you think this measures the variability? *Answers will vary by class. The range measures variability because if the values of the variable are really different, the range will be a big number (because the max and min will be far apart); but if there is little variability, the range will be small. For example, if all of the values were the same, we would have no variability and the range would be 0 because the max and min would be the same number.*
- c. How many **data points** does this distribution have? Why? *Answers will vary by class.*
- d. What is the amount of [insert variable name] that appears most often in a snack? Why? *Answers will vary by class.*
- e. What do you think the phrase *distribution of the dataset* means? *Answer: It shows us how values are distributed. We learn where there are many values, where there are only a few values, and where there are no values at all.*
- f. What information is this distribution telling us about the [insert variable name] in the snacks we consume? *Answers will vary. We see how the value of [variable name] varies. For example, we can see whether all foods have the same number of calories, or if they differ.*
- g. A distribution tells us two things: the values of the variable and the frequency of the values. "Frequency" is just another way of saying "the count". Why is this dotplot a picture of the distribution of [variable name]? *Answer: Because the location of the labels on the x-axis tells us the values we saw, and the number of labels at that value tells us the frequency for that value.*



17. Review the students' responses in a class discussion. Ask students to put a check mark next to the answers that were validated, and to revise the answers that need to be corrected.

18. Use another piece of poster paper to create a distribution for the second numerical variable chosen by the students. Repeat steps 14-16 from above with this variable.

19. On the first visualization for the numerical variable, show students how to convert the nutrition facts labels into something more readable. Draw another horizontal line on the plot above the nutrition facts labels. Explain that we can represent each label with an item such as a dot sticker or a sticky note.



20. Then, start with the minimum x-value on the plot and place the new sticker above the second horizontal axis. Do this for each nutrition facts label in the plot. Once all values have been represented, ask the students how this new plot IS or IS NOT better than the original. Explain that we can call this type of plot a **dotplot** since we're using dots to represent each observation.

21. Distribute the *Tangible Plot* handout (LMR_U1_L8). Each student should pick one of the 2 numerical variables plotted on the poster paper. Then, they should complete part 1 of the *Tangible Plot* handout before leaving class. They will complete part 2 of the handout for homework.

Name: _____ Date: _____

Tangible Plot

Part 1:

Make a sketch of the dotplot you and your classmates created in the space provided. Make sure you label your axes and give your plot a title.

LMR_U1_L8

22. Ask students to think about the statistical questions they came up with. Can the visualizations they created in class help answer their statistical question? If yes, answer the question; if not, what visualization would be appropriate?

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework



Students will complete part 2 of the *Tangible Plot* handout (LMR_U1_L8) and bring it to the next class for assessment.

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lesson 9: What is Typical?

Objective:

Students will learn about the typical value when looking at a distribution by finding the “center” and determining any point clusters.

Materials:

1. Nutrition facts dotplot (from Lesson 8)
2. Poster paper
3. Markers, dot stickers, or sticky notes

Vocabulary:

typical, center, shape, spread

Essential Concepts: The “center” of a distribution is a deliberately vague term, but it is one way to answer the subjective question “What is a typical value?” The center could be the perceived balancing point or the value that approximately cuts the area of the distribution in half.

Lesson:



Optional – Advance Preparation Needed:

- Watch this video to learn how to collect data in real time! You may consider using this tool for steps 12-13 in this lesson:
<https://www.youtube.com/embed/N-CkhD33IxE>

1. Food Habits Campaign Data Collection Monitoring:

- a. Display the IDS Campaign Monitoring Tool, found at <https://portal.thinkdataed.org/>. Click on **Campaign Monitor** and sign in.
- b. Inform students that you will be monitoring their data collection again today.
 - i. See *User List* and sort by *Total*. Ask: Who has collected the most data so far?
 - ii. Click on the pie chart. Ask: How many active users are there? How many inactive users are there?
 - iii. See *Total Responses*. How many responses have been submitted?
 - iv. Using TPS, ask students to think about what they can do to increase their data collection.
- c. Encourage students to monitor their own data collection. Show this video:
<https://www.youtube.com/embed/Xg9FI9arETw>

2. Inform students that today they will be learning about a distribution’s **typical** value.
3. Ask the class to brainstorm characteristics of the “typical” student. Does the typical 12th grader differ from the typical 9th grader? How so? *They may say that everyone is different, and that there’s no typical student. Keep pressing them to identify characteristics that are typical. The idea is to get them to recognize that there is variability, and yet we might still have an opinion about what is typical. For instance, not all students walk to school, but this might still be the typical experience.*
4. Give students 3 minutes to write down synonyms to the word “typical” in their DS journals. After time is up, have the students share their responses and keep a record on the board. *Some possible synonyms might be: normal, average, usual, standard, representative, regular, ordinary, natural, etc.*
5. Once students share their synonyms, ask students to think about which terms apply best to categorical variables and which terms apply best to numerical variables. Ask volunteers to share out their thoughts and give a brief explanation of why they categorized the term as either applying best to categorical variables or numerical variables. Create a T-chart on the board to keep track of their categories.





6. Next, display the dotplot created by the class with their nutrition facts labels during the previous class (from Lesson 8). Ask: what value might we consider to be the typical value of this distribution? *Answers will vary by class. Common answers will be to identify the mode (the value with the most labels) or the value in the center. A common wrong answer will be to confuse the frequency with the value. For example, they will say the most typical number of calories was “3” because, perhaps, 100 calories occurred 3 times, and that was more often than any other value. Students may also identify “clumps” of data: “it’s somewhere between 110 and 120.” That’s okay, but probe them as to why they chose that chunk and not another. The point is to get them to see that chunk as being in the middle or center of the distribution.*
 - a. Hopefully, at least one student will choose a value close to the center of the distribution. If not, point to a value near the extreme and ask them if they think this is typical. Then move closer to the center until they agree on which values are typical.
 - b. It is ok to be vague in the definition of typical for today’s lesson. The discussion needs to be very teacher-driven. Some possible points of discussion might be:
 - i. Clustering/clumps of data.
 - ii. Most of the observations are between _____ and _____.
 - iii. Overall range of the data.
7. Ask students to reconsider the typical number of sugar grams. What is the typical amount of sugar (in grams) in our snacks? *For example, students may come up with the same answer for different reasons: “The typical amount of sugar grams is 10.” The reasons may include the data points are half below and half above; it’s the mode; it has plurality.* Then, tie it back to the synonyms they provided. Ask: Which synonym are you using?



8. In pairs, ask students to discuss the question:
 - a. Which synonyms are associated with “center”? Is this concept of **center** useful for numerical or categorical variables? *Center is useful for numerical variables. The center of the distribution often corresponds to our notion of ‘typical value.’ For example, the typical height of the students in our class might be centered around 5'5".*
9. Inform students that the value at the center of the distribution often matches up with our everyday notion of the typical value of a distribution. The middle observation is not always the typical value. Similarly, the middle person would not always be the center value.
10. Defining the center of a distribution depends on many things, such as the placement of points in the distribution (known as the **shape**) and how dense the distribution is at certain values (known as the **spread**).
11. Ask the students to write down the number of hours of sleep they got last night. They will be creating a dotplot of this data, so ask them:
 - a. What do you think the typical value will be?
 - b. What do you think the lowest value will be?
 - c. What do you think the highest value will be?
 - d. What do you think the shape of the distribution will look like?



12. One-by-one, have them come up to the board (or poster paper) and put a dot above the correct value on the dotplot. After each student has placed a dot on the board, have a discussion about the distribution. Is the typical value similar to what they originally thought? The shape? The variability? Why or why not?
13. Next, have the students write down the number of hours of sleep they hope to get this Saturday. Ask: how do they think this plot will differ from the first plot? Focus discussion on the shape, center, and spread of the distributions. Repeat steps 8-9 from above and discuss how this plot is similar and/or different than the first plot.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lesson 10: Making Histograms

Objective:

Students will understand that a histogram represents observations grouped into bins, and that bars are drawn to show how many observations (or what proportion of the observations) lie in each bin, rather than representing individual observations, as in a dotplot.

Materials:

1. Peanut butter
2. Jelly
3. Loaf of sliced bread
4. Butter knife
5. Plate
6. *Sleep* dotplots (from Lesson 9)
7. Poster paper
8. Markers

Vocabulary:

algorithm, histogram, bin(s), bin widths, output, input, left-hand rule, right-hand rule

Essential Concepts: Histograms can be created through the use of an algorithm. The distributions displayed in a histogram can be classified using the technical terms for the shapes of the distributions. Learning to describe routine tasks through an algorithm is an important component of computational thinking.

Lesson:

1. Inform students that they will be telling us how to make a sandwich today. Giving clear, concise instructions to others is an important skill for students to learn. In this activity, students will practice using descriptive vocabulary, communicating ideas to others, recognizing steps in a process, and recognizing the importance of the use of clear language.
2. Prepare for this task by gathering the necessary materials for making a peanut butter and jelly sandwich and arranging them in a way that makes them easy to use. You may want to wear an apron and have a trash bag smock – this can get messy but that’s most of the fun!

Note: Be aware of peanut allergies! If any of your students are allergic to peanut butter, DO NOT ALLOW STUDENTS TO HANDLE THE PEANUT BUTTER! Peanut allergies can be very serious and children can have reactions without even eating it. So be aware and be careful!



3. Ask your students if they have ever followed a recipe before.
 - a. What kinds of things have they made?
 - b. Does anyone know how to make a peanut butter and jelly sandwich?
 - c. Would they teach you how?
 - d. Would they give you all the steps to make a sandwich?
4. Show your students the materials you have for making a sandwich. Have students take out paper and pencils and ask each student (or pair of students) to write down their instructions for making a peanut butter and jelly sandwich. We can also call these instructions an **algorithm**.
5. Explain to students that precise instructions for any process are like a formula to follow in order to get the same results each time. Also, an algorithm is how we communicate with the computer. The teacher will function as the computer. Your job is to give him/her rules so that he/she can carry out and successfully make a PB&J sandwich.
6. Every algorithm needs input and produces output. The output here will be a PB&J sandwich. What is the input? *Answer: Steps, or actions to follow.*

7. Tell students that when they are done you will select someone to share their instructions, and you will make a sandwich following the instructions.
8. Select a student to read their instructions and do EXACTLY what it says. For example, if it says “put the peanut butter on the bread,” you can literally put the jar of peanut butter on the bag of bread. There was no instruction to open the bread or the jar of peanut butter, no instruction to use the knife in any way, etc. Listen for other examples of unclear instructions and think of how you might act them out. If students are not clear about where to spread the peanut butter, put it on the crust. The more literal you are by doing exactly what the instructions say, the funnier the activity will be and the more likely you are to get your point across about the importance of clear instructions.
9. After your first sandwich, ask your students if they think their instructions were clear or not. What are some things they might have done differently?
10. Select another student to read his/her instructions. They will be sure to use clarifications of the instructions you acted upon before – this is a good thing!
11. After you finish the sandwich, ask your students if they think clear instructions are important. Why?
12. Let students know that they will now develop an algorithm for building a histogram to represent the sleep dotplots they created in the previous lesson.
13. Explain that a **histogram**, rather than showing the frequency for each value, shows the frequency (or percent, but we will focus on frequency) of all the values that fall in a certain range, called a **bin**. For example, we might choose bins that go from 0-5, 5-10, 10-15, 15-20, 20-25. **Bin widths will vary by class.**
14. Model how to create a histogram using the data from the dotplot “hours of sleep last night.” On a blank chart, create the x-axis with bin widths 0-3, 3-6, 6-9, etc. and place marks on the plot at those intervals and ask students: “What are the frequencies in each bin?”
Notice that multiples of three appear in more than one bin. Let’s take the value of 6 hours as an example. Should those observations be included in the second bin (3-6) or the third bin (6-9)?
 - If students include the values of 6 hours in the second bin, then they are using the **left-hand rule**.
 - If students include the values of 6 hours in the third bin, then they are using the **right-hand rule**.
15. Once the frequencies have been determined, draw the bars with corresponding heights. Do not include spaces between the bars as time is a continuous variable.
16. Next, student teams will create an algorithm that gives directions for how to construct a histogram for the data from the dotplot for “hours of sleep they hope to get on Saturday.” Remember, an algorithm is a set of rules that can always be applied. Similar to the way they wrote a process for making a PB&J sandwich, students will write a process for creating a histogram. Tell students to continue thinking of the process to transform the data in the dotplot to create a histogram. The algorithm will produce an **output**, which will be a histogram. What’s the **input**? *Answer: Data, or maybe the dotplot.*
17. Inform the students that you will provide a piece of input: how wide the bin will be. For instance, it might be 5 hours, it might be 1 hour, or it might be 10 hours (or half an hour!). Whatever it is, their algorithm should work for any input value.
18. Let students work for a bit. They should write out Step 1, Step 2, etc. Then choose a group and ask them to get you started. Give them a bin width of 4.



19. Teachers should sketch the histogram on the board or chart paper as students read their algorithms. Again, teachers should take things very literally. For example, if they do not tell you exactly where the bins should start, start one way off to the left. If they are vague and say “divide the number line into groups of 10,” then make them arbitrary sizes. If they have to be the same size, ask them how to do that. Points to consider:
- Where do we start drawing the bins? Always at the location of the smallest dotplot? Always at the greatest? A little to the left?
 - What do we do with points that fall exactly on a boundary? Do they go to the bin on the left or on the right? Does it matter? *Answer: No, it doesn't matter as long as they're consistent (all in the left or all on the right).*
 - Can we do it differently every time? *Answer, No. We need to be consistent. This is called either the left-hand rule or the right-hand rule, depending on which is chosen.*
20. After following 2 or 3 algorithms, ask students if they feel their algorithm is precise enough. Allow students time to revise their algorithms.



21. Have a class discussion about the similarities and differences between the original dotplot and a histogram. Ask:
- What have we gained from the histogram? *Answer: We now can see the shape of the distribution as a whole.*
 - What have we lost? *Answer: We lost each individual observation by grouping them into bins.*

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lesson 11: What Shape Are You In?

Objective:

Students will learn to classify distributions in terms of shape and can suggest theories for why a distribution might be one shape or another.

Materials:

1. *Sorting Histograms* handout (LMR_U1_L11) – one copy per group of 4 students. (This activity comes from the AIMS project, University of Minnesota, J. Garfield.)

Advance preparation required: Needs to be cut into sets (see step 1 in lesson)

Vocabulary:

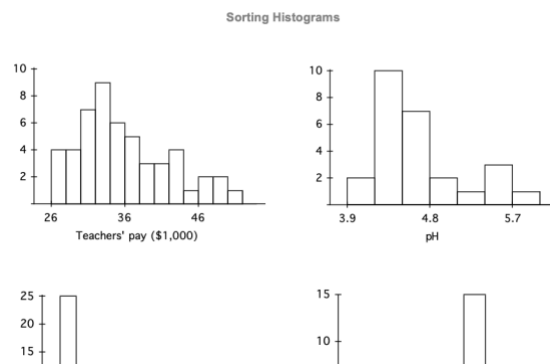
symmetric, left-skewed, right-skewed, unimodal, bimodal

Essential Concepts: Identifying the shape of a histogram is part of the **interpret** step of the Data Cycle.

Lesson:

1. Distribute the cutouts from the *Sorting Histograms* handout (LMR_U1_L11). Give each student team all of the 24 histograms (can be paper clipped together or put in small zippered bags).

Advance preparation required: Print the *Sorting Histograms* file (LMR_U1_L11). Cut each histogram so that it is on its own piece of paper. Create enough sets for each team to have all 24 histograms. They can be paper clipped together, or put in small zippered bags.



LMR_U1_L11

2. Inform students that the type of data being measured is indicated on the horizontal axis, and the vertical axis represents how many observations are in each bar.
3. The students will then sort their stack of plots into different piles according to their shapes. Histograms that have similar shapes should be sorted into the same stack.
4. Once the student teams have agreed upon the histogram shape groupings, they should discuss and write down answers to the following in their DS journals:
 - a. Describe what's similar about the plots in each group. *Answers will vary but should be grouped by the overall shape of the distribution. For example, plots with a higher density of bars on the right side of the plot should all be in the same group.*
 - b. Pick one graph in each group that is the best example of that group. *Answers will vary.*
 - c. Give the group a name that you think describes the general shape. *Answers will vary.*
 - d. If there are graphs that do not fit into any group, try to determine why it was impossible to place them. What is different or confusing about them? *Answers will vary.*

5. After each team has had time to discuss and write down their observations, have a class discussion about the histogram groupings. Do the students agree about the general shapes?
6. In statistics, we use specific terminology when discussing the shapes of distributions, such as **symmetric**, **right-skewed**, **left-skewed**, **unimodal**, **bimodal**, etc. Did any of the teams use these terms? If not, introduce each one and ask which of the 24 histograms could be classified as that shape.
7. Next, introduce the following scenarios and ask students to determine what a corresponding histogram might look like. They should use statistical terms to describe their answer.
 - a. The grades on an easy test. *Answer: Left-skewed, unimodal.*
 - b. The grades on a difficult test. *Answer: Right-skewed, unimodal.*
 - c. The number of times IDS students study during the first week of class. *Answers will vary.*
 - d. The age of cars on a used car lot. *Answer: Right-skewed, probably unimodal.*
 - e. The amount of time spent by students on a difficult test (max time allowed is 50 mins). *Answer: Left-skewed, but may also just be one bar with all observations at 50 mins, unimodal.*
 - f. The heights of students in your high school band. *Answer: Symmetric, bimodal.*
 - g. The salaries of all persons employed at the Los Angeles Unified School District. *Answer: Right-skewed, potentially bimodal (teachers vs. administrators).*

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lesson 12: Exploring Food Habits

Objective:

Students will experience the full Data Cycle, and for the first time will do so with data they have collected. They will use the Dashboard and PlotApp, tools that are easy to learn. This first time, the teacher will “navigate and steer” so that students can focus on asking questions and interpreting the plots.

Materials:

1. Computers
2. Projector
3. *Food Habits Check-In* handout (LMR_U1_L12_A)
4. *Exploring Our Food Habits* handout (LMR_U1_L12_B)

Essential Concepts: Once Participatory Sensing data has been collected, the Dashboard and PlotApp perform the analysis step of the Data Cycle, though humans need to tell the computer which plots we wish to examine.

Lesson:



In preparation for this lesson, watch these two videos:

- Navigating the Dashboard: <https://www.youtube.com/embed/d0BDaHKOqOg>
- Navigating the PlotApp: <https://www.youtube.com/embed/Jks39Gxi6dA>

1. Ask students to reflect about their experience so far with the Food Habits Participatory Sensing campaign by completing the *Food Habits Check-In* handout (LMR_U1_L12_A).

Name: _____ Date: _____

Check-In
of Your Food Habits

1. How well have you done collecting data for this project?
Circle one of the choices below and explain why you ranked it at that level.

(5)	(4)	(3)	(2)	(1)	(0)
Excellent	Very Well	Average	Below Average	Not as well as I wanted	Collected no data

2. What do you think your snack healthy levels are? Did you eat more healthy snacks or less healthy? Why?

LMR_U1_L12_A

2. Demonstrate how to access the **IDS Homepage** found at <https://portal.thinkdataed.org/>.
3. Explain to students that all of the IDS web tools can be accessed through this page.
4. For this lesson, students will need to observe the teacher using the **Campaign Manager**, the **Dashboard**, and the **PlotApp**.
5. Click on the Campaign Manager. Explain that selecting any of the web tools on the IDS page without logging in first will redirect them to the login prompt.
6. Demonstrate how to log in to access the IDS software suite. Inform students that they will use the same login information they have used to collect data on their mobile devices or the browser-based version.

Inform students that the Campaign Manager is the place where they will access, download, and export their campaign data in subsequent lessons. It also provides shortcuts to the Dashboard and PlotApp. The Campaign Manager allows them to view and learn about the campaigns in which they are participating, and to edit the campaigns they will be creating later in this course. Show them the drop-down menu on the right-hand side and explain that they will only be concerned with the **Responses** tab for this lesson.

Explain that the Responses tab allows them to view, delete, or share their data, and to view shared data contributed by other users of the campaign.

7. Distribute the *Exploring Our Food Habits* handout (LMR_U1_L12_B).

Name: _____ Date: _____

Exploring Our Food Habits

Using the Dashboard, answer the following statistical questions:

VARIABLE(S)	SKETCH OF PLOT (ANALYSIS)	STATISTICAL QUESTIONS AND INTERPRETATIONS
	Quickly sketch an image of the plot, name the type of plot, and <i>appropriately label</i> your plot.	Answer the statistical question based on what is shown in the plot.
1. Variable: When		When were the majority of snacks eaten?

LMR_U1_L12_B

8. Inform students that the teacher will be navigating through the IDS website as the students follow along in the *Exploring Our Food Habits* handout (LMR_U1_L12_B).



9. Once completed, students should turn in the handout for assessment.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lesson 13: RStudio Basics

Objective:

Students will learn RStudio/Posit Cloud's interface, as well as a few basic commands to discover the structure behind a dataset.

Materials:

1. Computer
2. Projector
3. RStudio: <https://portal.thinkdataed.org/>
4. Video showing how to log into RStudio/Posit Cloud for the first time: <https://www.youtube.com/watch?v=DSf2viLkSCc>

Vocabulary:

pane, preview, console, plot, environment

RStudio Commands:

`data( ), View( ), names( ), help( ), dim( ), tally( ), load_labs( )`

Essential Concepts: The computer has a syntax, and it can only understand if you speak its language.

Lesson:

Before
inviting
students to
your
RStudio/
Posit Cloud
Teacher
Space!

Before inviting students to your RStudio/Posit Cloud Teacher Space, ensure that:

- Students sign-in to RStudio/Posit Cloud using the “Log In With Google” option using their school email address.
- Each student watches the video showing how to log into RStudio/Posit Cloud for the first time found here: <https://www.youtube.com/watch?v=DSf2viLkSCc>
- You are familiar with managing your RStudio/Posit Teacher Space. See video here: <https://www.youtube.com/watch?v=KwXW4Dfjn50>
- In preparation for this lesson, watch this video: <https://www.youtube.com/embed/WkxCfaol3pE>

1. Inform students that the Dashboard and PlotApp are data visualization tools that are coded in R, the statistical programming software that academics and professional statisticians use. The Introduction to Data Science course will utilize RStudio, which also runs on R. They will learn the programming language of RStudio for data analysis.
2. Demonstrate how to access RStudio by projecting the URL: <https://portal.thinkdataed.org/> on a screen. Then, click on the RStudio (Posit Cloud) icon on the page.
3. Inform students that they will log into RStudio using the “Log In with Google” option. Note that it is not the same as their IDS App & IDS Homepage login.
4. Once logged in, show each **pane**, or rectangular area, of the RStudio interface:
 - a. **preview** (spreadsheet) – where they will be able to see the variables and observations (index); rows and columns of data
 - b. **console** – where they will be entering their code
 - c. **plot** – where their plots/graphs/visualizations will be generated
 - d. **environment** – where they will see values and objects
5. Inform students that they will be looking at a dataset from The Centers for Disease Control and Prevention (CDC), a government agency that collects data about teenagers on a variety of topics.

6. Demonstrate how to load and view the CDC data file to the workspace by typing the following command in the console:

```
>data(cdc)  
>View(cdc)
```
7. Examine the environment pane. Ask a student to describe how the data are displayed. *Answer: The data are displayed in rows and columns.*
8. Demonstrate how to list the variables found in the CDC dataset. Students may take notes and write down commands in their DS journals:
 - a.

```
>names(cdc)
```
 - b. Ask: What do you notice? What is one variable of this dataset? How many variables are there? How does this output compare to the information in the preview pane? *Answer: This command lists the names of each variable in the dataset. There are 32 variables; age, sex, grade, height, weight, etc. Answers will vary about how this output compares to the information in the preview pane.*
9. Demonstrate how to obtain more detailed information about the dataset by typing the following command in the console.
 - a.

```
>help(cdc)
```
 - b. Ask: What unit of measurement is height reported in? *Answer: Height was reported in meters.*
10. Demonstrate how to find the number of rows and columns in the dataset.
 - a.

```
>dim(cdc)
```
 - b. Ask: Which number do you think represents the rows? Which one represents the columns? How does this output compare to the information in the preview and environment panes? How many observations are there in the dataset? How many variables does this dataset contain? *Answer: The first number represents the rows and the second number represents the columns. There are 17,232 rows, or 17,232 observations; and there are 32 columns, or 32 variables. This information is also visible in the preview pane.*
11. Next, show students how to access the number of observations of a specific variable.
 - a.

```
>tally(~seat_belt, data = cdc)
```
 - b. Ask: What do you notice? Describe the output. *Answer: Notice that six categories are displayed. Each category shows the number of observations contained in it. E.g., "Never" has 265 observations, meaning 265 teens reported never wearing their seat belt as a passenger in a motor vehicle. <NA> = Not Available, represents teens that did not provide information about their seat belt habits.*
12. Change the variable to *height*.
 - a.

```
>tally(~height, data = cdc)
```
 - b. Ask: What do you notice? Describe the output. *Answer: The levels are missing. It happened because the variable 'height' contains numbers, not categories.*
-  13. Let's take a closer look at the variables *seat_belt* and *height*. Maximize the console. Ask teams to discuss the following question:
What is the difference between the data from the variables *seat_belt* and *height*? *Answer: The data from the 'seat_belt' variable is categorical, which means it consists of groupings. The data from the variable 'height' is numerical, which means it consists of numbers.*
14. Summarize: In data science, the variable *seat_belt* is what we call a **categorical variable**, and the variable *height* is what we call a **numerical variable**.

15. Let's look at the other variables in this dataset. In pairs, categorize each variable as categorical or numerical:
 - a. eat_fruit *Answer: categorical*
 - b. weight *Answer: numerical*
 - c. grade *Answer: categorical*
 - d. drive_text *Answer: categorical*
16. Inform students that they will be learning RStudio code to work with data. They will be completing RStudio labs throughout the course.
17. Demonstrate how to load the menu of labs by typing the following code:

```
>load_labs( )
```
18. The load labs command displays a list of available labs and a selection prompt. To select Lab 1A, type number 1 after the selection prompt.
19. Next, direct students' attention to the plot pane. Show them the location of Lab 1A's presentation.
20. Click on the arrows at the bottom right-hand side of the presentation to view each slide. Pause on a slide titled "R's most important syntax." There are 3 boxes, each containing a line of code.
21. Explain that every time they see a grey box with a line of code, they are to type the code in the console. The output will appear either on the console itself or on the plot pane.
22. Type in one of the lines of code. In this particular case, the output will be a plot. Show students the location of the plot and demonstrate how to toggle between the plots and presentation tabs.
23. Inform students that they will be completing the first lab, 1A, the next day.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next 3 Days

Students should continue to collect nutritional facts data using the *Food Habits* Participatory Sensing campaign on their smart devices or via web browser.

Lab 1A: Data, Code & RStudio

Lab 1B: Get the picture?

Lab 1C: Export, Upload, Import

Complete Labs 1A, 1B and 1C prior to Lesson 14.

Lab 1A: Data, Code & RStudio

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Welcome to the labs!

- Throughout the year, you'll be putting your data science skills to work by completing the labs.
- You'll learn how to program in the **R** programming language.
 - The programming language used by actual data scientists.
- Your code will be written in RStudio which is an easy to use interface for coding using **R**.

So let's get started!

- The data for our first few labs comes from the Centers for Disease Control (CDC).
 - The CDC is a federal institution that studies public health.
- **Type these two commands into your console:**

```
data(cdc)  
View(cdc)
```
- **Describe the data that appeared after running View(cdc):**
 - **(1) Who is the information about?**
 - **(2) What sorts of information about them was collected?**
- To find out more information about the **cdc** data, type the command below into your console.
 - To get back to the slides find and click on the *Viewer* tab.

```
?cdc
```

Data: Variables & Observations

- Data can be broken up into two parts.
 1. *Observations*
 2. *Variables*
 - *Observations* are the *who* or *what* we are collecting data from/about.
 - *Variables* are the measurements or characteristics about our *observations*.
- If need be, re-type the command you used to **View** your data. Then answer the following:
 - **(3) Based on the data, describe a few characteristics about the first observation.**
 - **(4) What does the first column tell us about our observations?**
- In order to describe the first observation, notice that you had to look at the first row of the spreadsheet. Each row, in this case, describes a person.
- The columns of the spreadsheet represent variables.

Uncovering our Data's Structure

- Now that we've looked at our data, let's look at how RStudio is organized.
- RStudio's main window is composed of four *panes*.
- **Find the pane that has a *tab* titled *Environment* and click on the *tab*.**
 - This pane contains a list of everything that's currently available for **R** to use.
 - Notice that **R** knows we have our **cdc** data loaded.

- (5) How many students are in our `cdc` dataset?
- (6) How many variables were measured for each student?

Some new functions

- Type the following commands into the console:

```
dim(cdc)
nrow(cdc)
ncol(cdc)
names(cdc)
```

- (7) Which of these functions tell us the number of observations in our data?
- (8) Which of these functions tell us the number of variables?

First Steps

- Typing commands into the console is your first step into the larger world of *programming* or *coding* (terms which are often used interchangeably).
- Coding is all about learning how to send instructions to your computer.
 - The way we *speak* to the computer, using a coding language, is *syntax*.
- R is one of many coding languages. Each coding language is slightly different, and these differences are reflected in the syntax.
- *Capitalization, spelling and punctuation* are REALLY important.

Syntax matters

- Run the following commands.

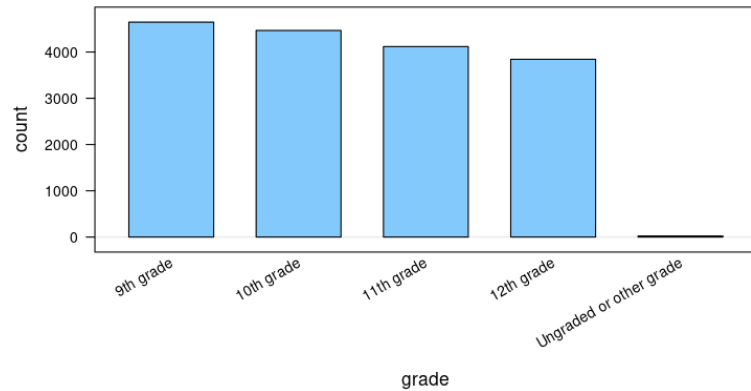
```
Names(cdc)
NAMES(cdc)
names(cdc)
names(CDC)
```

- (9) What happens after each command?
- (10) Which does R understand?

R's most important syntax

- Most of the commands you will be using follow the syntax below.
`function(y ~ x, data = ____)`
- To create graphs or plots you need to provide R with the following:
 - The name of the R function, often the plot's name, that tells the computer how to create your graph.
 - The variable(s) containing the information we want the function to use.
 - The dataset containing the variables.
- Notice that when we analyze a single variable the value for y is left blank.

```
bargraph(~grade, data = cdc)
```



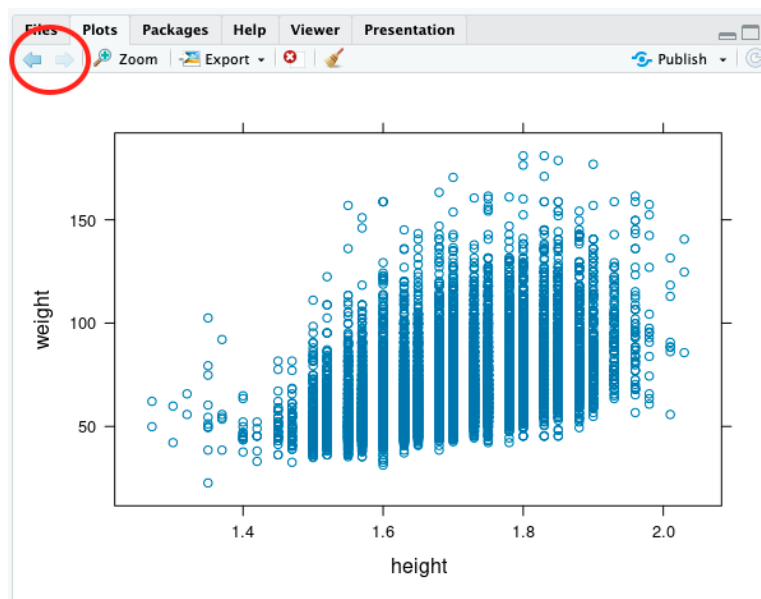
- Later on, we'll see we can use this syntax to do more than create graphs.

Syntax in action

- Search through the different panes. Find and then click on the *Plots* tab.
 - To get back to the slides, find and then click on the *Viewer* tab.
- (11) Would a histogram, bargraph, or scatterplot be useful for answering the question: *Is it unusual for students in the CDC dataset to be taller than 1.8 meters?*
- Run the three commands below then answer the question that follows.

```
histogram(~height, data = cdc)
bargraph(~drive_text, data = cdc)
xyplot(weight~height, data = cdc)
```

- (12) Do you think it's unusual for students in the *cdc* dataset to be taller than 1.8 meters? Why or why not?
- Hint: Use the arrow keys on the *Plots* tab to toggle between the plots.



Hint

On your own:

- After completing the lab, answer the following questions:
 - (13) What is *public health* and do we collect data about it?
 - (14) How do you think our data was collected? Does it include every high school aged student in the US?
 - (15) How might the CDC use this data? Who else could benefit from using this data?
 - (16) Write and run the code to visualize the distribution of weights of the students in the CDC data with a *histogram*. What is the *typical* weight?
 - (17) Write and run the code to create a *bargraph* to visualize the distribution of how often students ate fruit. About how many students did not eat fruit over the previous 7 days?

Lab 1B: Get the picture?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Where'd we leave off ...

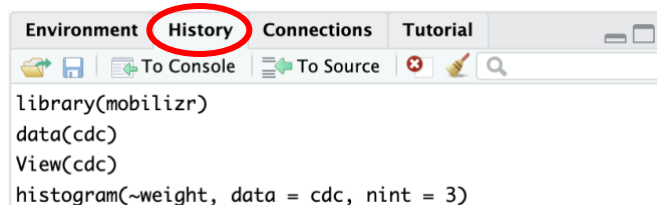
- In the previous lab, we started to get acquainted with the layout of RStudio and some of the commands.
- In this lab, we'll learn about different *types* of variables.
 - Such as those that are measured by numbers and others that have values that are categories.
- We'll also look at ways to visualize these different types of data using *plots* (a word data scientists use interchangeably with the word *graph*).
- **Find the *History* tab in RStudio and click on it. Figure out how to use the information to reload the *cdc* data.**

Variable Types

- Numerical variables have values that are measured in units.
- Categorical variables have values that describe or categorize our observations.
- **View your *cdc* data and find the columns for *height* and *sex* (use the *History* pane again if you need help to View your data).**
 - **(1) Is *height* a numerical or categorical variable? Why?**
 - **(2) Is *sex* a numerical or categorical variable? Why?**
 - **(3) List either the different categories or what you think the measured units are for *height* and *sex*.**

Which is which?

- **Run the code you used in the previous lab to display the names of your *cdc* data's variables (use the code displayed in the *History* pane to resubmit previously typed commands).**



- **Use the code's output to help you complete the following:**
 - **(4) Write down 3 variables that you think are *categorical* variables and why.**
 - **(5) Write down 3 variables that you think are *numerical* variables and why.**

Data Structures

- One way to get a good summary of your data is to look at the data's *structure*.
 - One way to view this info would be to click on the little blue arrow next to *cdc* in the *Environment* pane.
 - Another way would be to run the following in the console:

```
str(cdc)
```

- (6) What information does the `str` function output?
- (7) Were you able to correctly guess which variables were categorical and numeric? Which ones did you mislabel?

Visualizing data

- Visualizing data is a really helpful way to learn about our variables.
- (8) Choose one numeric variable. Write and run the codes to create a `bargraph` and a `histogram`.
- (9) Choose one categorical variable. Write and run the codes to create a `bargraph` and a `histogram`.
- (10) Which function, either `bargraph` or `histogram` is better at visualizing categorical variables? Which is better at visualizing numerical variables?

We have options

- (11) Write and run the code to make a graph that shows the distribution of people's `weight`.
 - (12) Describe the distribution of `weight`. Make sure to describe the *shape*, *center* and *spread* of the distribution.
- Options can be added to plotting functions to change their appearance. The code below includes the `nint` option which controls the number of *intervals* in a numerical plot.
 - Options, also known as *arguments*, are additional pieces of information you provide to a function, and are separated by commas.
 - Type the command below on your console and then answer the questions that follow:

```
histogram(~weight, data = cdc, nint = 3)
```

- (13) How did including the option `nint = 3` change the `histogram`?
- (14) Does setting `nint = 3` impact how you would describe the shape, center and spread?
- (15) Try other values for `nint`. What value produced the best graph? Why?

How often do people text & drive?

- (16) Write and run the code to make a graph that shows how often people in our data texted while driving.
 - (17) What does the y-axis represent?
 - (18) What does the x-axis tell us?
 - (19) Would you say that *most* people *never* texted while driving? What does the word *most* mean?
 - (20) Approximately what percent of the people texted while driving for 20 or more days? (Hint: There are 17232 students in our data.)

Does texting and driving differ by sex?

- (21) Write and run the code to make a side-by-side `bargraph` that could answer this question: *Does texting and driving differ by sex?* Use the following fill-in-the-blank code as a hint.
`bargraph(~____, data = _____, groups = _____)`
- (22) Write a sentence explaining how boys and girls differ when it comes to texting while driving.

- (23) Would you say that most girls never text and drive? Would you say that most boys never text and drive?
- (24) How did including the `groups` argument in your code change the graph?

Do males and females have similar heights?

- To answer this, what we'd like to do is visualize the distributions of heights, separately, for males and females.
 - This way, we can easily compare them.
- (25) Write and run the code to create a `histogram` for the `height` of males and females using the `groups` argument.
 - (26) Can you use this graphic to answer the question at the top of the slide? Why or why not?
 - (27) Is grouping numeric values, such as heights, as helpful as grouping categorical variables, such as texting & driving?

Do males and females have similar heights?, continued

- Why does this work for bargraphs but not histograms?
 - The `groups` argument uses color to differentiate between groups.
 - With bargraphs, each group is split with bars next to each other on the x-axis.
 - With histograms, the x-axis is a continuous set of numbers so the bars overlap making it difficult to compare center and spread.
- (28) Write and run the code to create a `split histogram` and answer the questions below:
`histogram(~___ | ___, data = ___ )`
- (29) Do you think males & females have similar heights? Use the plot you create to justify your answer.
- (30) Just like we did for the `histogram`, is it possible to create a `split bargraph`? Write and run the code to create a `bargraph` of `drive_text` that is split by `sex` to find out.

On your own:

- In this lab, we looked at the texting & driving habits of boys and girls.
- (31) What other factors do you think might affect how often people text and drive?
- (32) Choose one variable from the `cdc` data, make a graph, and use the graph to describe how `drive_text` use differs with this variable.

Lab 1C: Export, Upload, Import

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Whose data? Our data.

- Throughout the previous labs, we've been using data that was already loaded in RStudio.
 - But what if we want to analyze our own data?
- This lab is all about learning how to load our own Participatory Sensing data into RStudio.

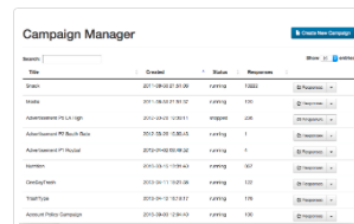
Export, upload, import

- Before we can perform any analysis, we have to load data into R.
- When we want to get our Participatory Sensing data into RStudio, we:
 - Export** the data from your class's campaign page.
 - Upload** data to RStudio server.
 - Import** the data into R's working memory.
- NOTE: You can watch the following video for a step-by-step walk-through of the process:*
<https://www.youtube.com/watch?v=4mChtv5qy1g>

Exporting

- To begin, go to the **IDS Tools** page.
 - Click on the **Campaign Manager**.
 - Fill in your username and password and click "Sign in".

Campaign Manager



The screenshot shows the Campaign Manager interface with a table of campaigns. The table has columns for Title, Created, Status, and Resources. There are 8 campaigns listed, each with a 'Download' button and a 'Resources' link.

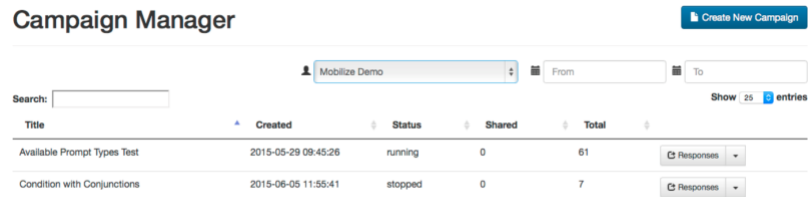
Title	Created	Status	Resources
Urban	2017-08-02 17:30:00	running	1000
Arts	2017-08-02 17:30:00	running	100
Advertisement for LA Light	2017-08-02 17:30:00	stopped	100
Advertisement for Beach Race	2017-08-02 17:30:00	running	1
Advertisement for Beach Race	2017-08-02 17:30:00	running	1
Arts	2017-08-02 17:30:00	running	100
Advertisement	2017-08-02 17:30:00	running	100
Advertisement	2017-08-02 17:30:00	running	100

Manage and create campaigns

Campaign Manager

If you forget your username or password, ask your teacher to remind you.

Campaign Manager



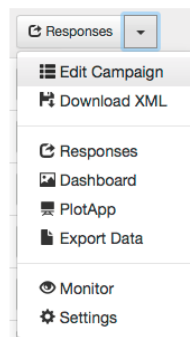
Title	Created	Status	Shared	Total	
Available Prompt Types Test	2015-05-29 09:45:26	running	0	61	Responses
Condition with Conjunctions	2015-06-05 11:55:41	stopped	0	7	Responses

Campaign Manager

- After logging in, your screen should look similar to this.
- Click on the dropdown arrow for the campaign you are interested in downloading.
 - At this point in the course, it will most likely be the Food Habits campaign.

Dropdown Arrow

- The options for the dropdown menu will look like this:



Campaign Tab

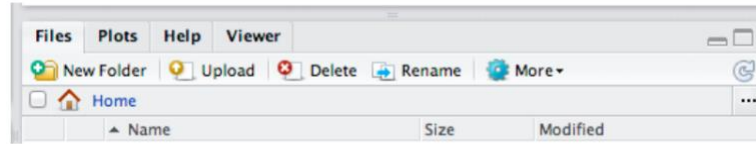
- Click on the option labeled **Export Data**.
 - Remember where you save your file!

Exporting

- When you clicked the **Export** link, a .csv file was saved on your computer.
- Now that the file is on your computer, we need to **upload** it into RStudio.

Uploading

- Look at the four different *panes* in RStudio.
 - Find the *pane* with a **Files** tab.
 - Click it!



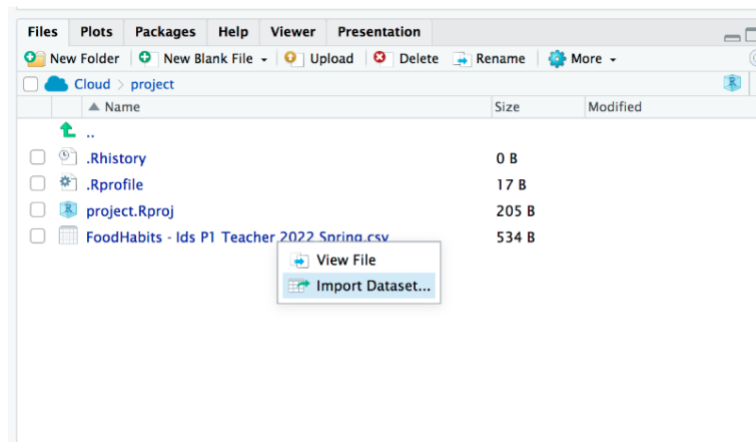
Upload Button

- Click the button on the Files pane that says “Upload”.
 - Click on “Choose File” and find the *SurveyResponses.csv* file you saved to your computer.
 - Hit the OK button.
- Voila!
 - If you look in the Files pane, you should be able to find your data!

Upload vs. Import

- By **Uploading** your data into RStudio you’ve really only given yourself access to it.
 - Don’t believe me? Look at the **Environment** pane ... where’s your data?
- To actually use the data we need to **Import** it into your computer’s memory.
- To compute more quickly and efficiently, R will only keep a few datasets stored in its memory at a time.
 - By importing data, you are telling R that this is a dataset that is important to store it in its memory so you can use it.

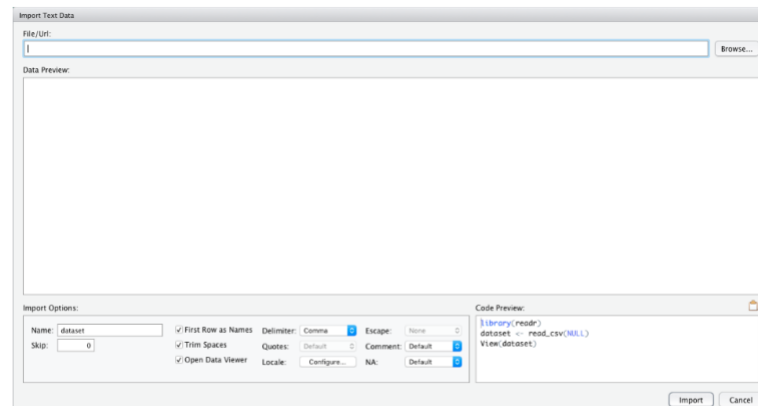
Importing



Import Data

- On the Files pane, find the data you want to **Import**.
- Click on the name of the file and choose the option “Import Dataset...”.

Data Preview

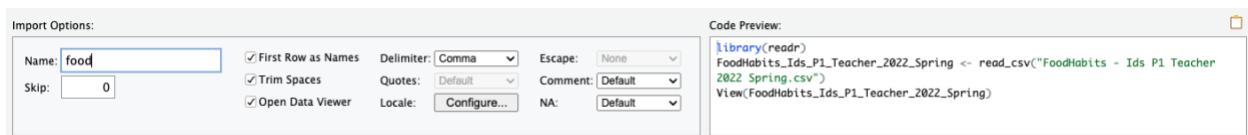


Data Preview

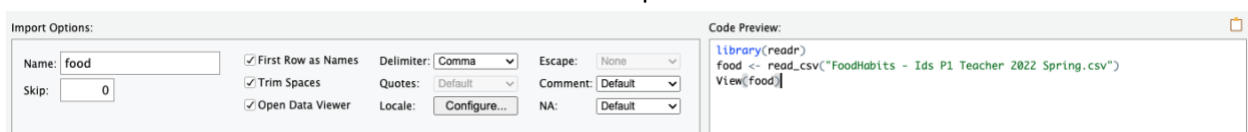
- You can give your data a name using the **Name** field in the lower left corner.

What's in a name?

- The name you give your data is what you will use when you write code to analyze your data.
 - Good names are short and descriptive.
 - For your Food Habits campaign, some good names to use would be “foodhabits” or even just “food”.
- When you're ready, click the *Import* button.
- Make sure the name under Import Options matches the name in the Code Preview before you click Import.



nonexample



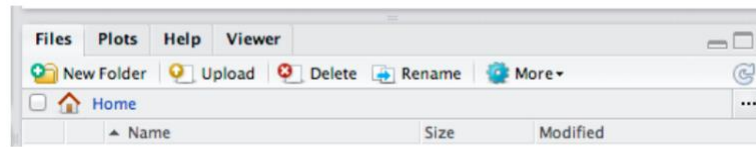
example

read.csv()

- After you click *Import* you might notice something appeared in your console.

```
data.file <- read_csv("~/SurveyResponse.csv")
View(data.file)
```
- This is the actual code **RStudio** uses to read your data when you clicked the *Import* button.
 - So instead of using the **RStudio** buttons, we can actually **Import** by writing code similar to what was output into the console!
 - This will come in handy later in the course.

A word on staying organized...



Rename Button

- The **Files** tab has a few other features to help keep you organized.
 - *CampaignName* – *Ids Teacher Year Semester* probably isn't the best name for your data. Click **Rename** to give it a clearer name.
 - Often, it's helpful to give your data file the same name as when you import your data.
 - So in this case, we could name our data file *food.csv*.

Analysis Time

- After you **Export**, **Upload**, and **Import** your data, you're ready to analyze.
- **(1) View your data and select a variable. Write and run a code to create an appropriate plot for that variable.**
- **(2) What variable did you choose? What does the plot tell us about that variable?**
- **Summarize the process:**
 - **(3) What might your future self forget about this process? Summarize the steps you took in this lab to use new data in RStudio.**

Lesson 14: Variables, Variables, Variables

Objective:

Students will learn how to read and interpret multiple variable plots: bivariate scatterplots, multiple variable scatterplots, stacked bar plots, and side-by-side bar plots. They will summarize their learning about multiple variable plots using a four-fold graphic organizer.

Materials:

1. *Scatterplot of Heights & Weights* handout (LMR_U1_L14_A)
2. *Scatterplot of Heights & Weights, Split by Sex* handout (LMR_U1_L14_B)
3. *Side-by-Side Bar Chart* handout (LMR_U1_L14_C)
4. *Faceted Histogram of Height by Sex* handout (LMR_U1_L14_D)
5. *Summarizing Multi-Variable Plots* graphic organizer (LMR_U1_L14_E)

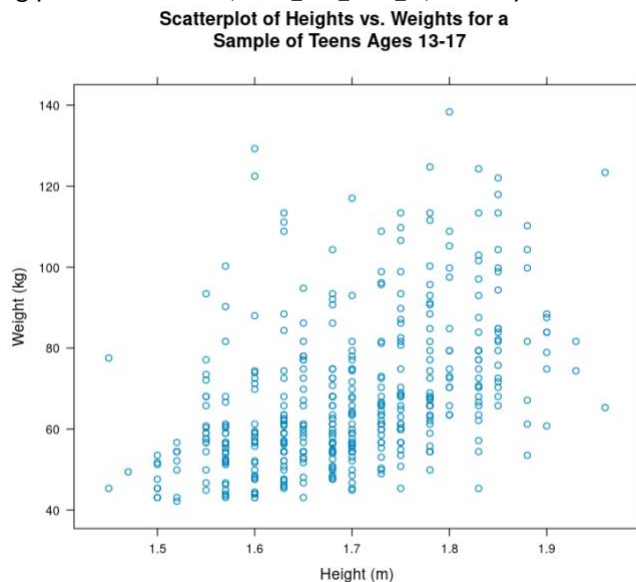
Vocabulary:

scatterplot, grouping, side-by-side bar plot

Essential Concepts: To examine whether two (or more) variables are related, we can plot their distributions on the same graph.

Lesson:

1. Begin by informing students that they will learn how to make visual displays using more than one variable, and by asking them to ponder the following questions:
 - a. What do you think is the relation between people's heights and weights?
 - b. Are taller people heavier? Always? Or is this just a tendency?Let's look at some data.
2. Display the following plot to the class (LMR_U1_L14_A) so they can see some actual data:



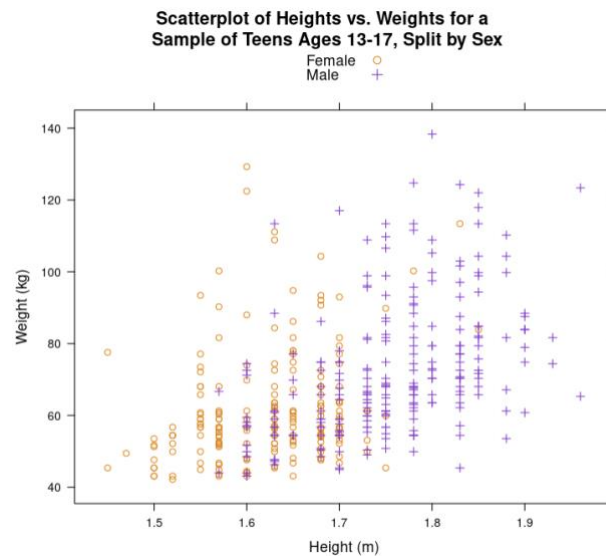
LMR_U1_L14_A

3. Ask students to individually answer the following questions about the plot on the handout (LMR_U1_L14_A):
 - a. What kind of plot is this and how will you remember its features? *Answer: Scatterplot. Answers will vary about how students will remember its features.*

- b. How many variables are displayed in this plot? Name the variable(s) and identify the type of variable(s). *Answer: Two variables. Weight in kilograms and height in meters. They are both numerical variables.*
- c. What do the axes show? *Answer: The x-axis shows the height of teens in meters, and the y-axis shows the weight of teens in kilograms.*
- d. Do taller people weigh more? *Answer: Not necessarily, but there is a tendency for this to be true.*



4. Discuss this plot with the class by eliciting students' responses to the questions. Students actively listen to the discussion by confirming, correcting, or adding to their own responses.
5. Close the discussion by asking students: What questions might you have about this plot? What additional information would be helpful? *Answers will vary.*
6. Now, suppose we could see which of these dots represented girls and which represented boys. Where do you think most of the girls' dots would be relative to the boys?
7. Display the following plot to the class (LMR_U1_L14_B):



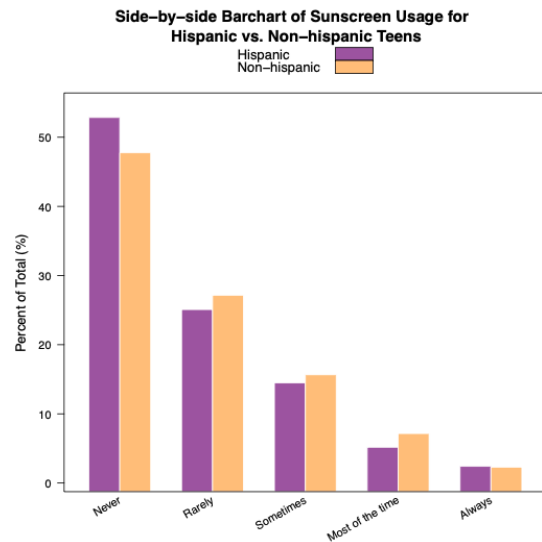
LMR_U1_L14_B

8. Ask students to individually answer the following questions about the plot on the handout (LMR_U1_L14_B):
 - a. What kind of plot is this and how will you remember its features? *Answer: Scatterplot. Answers will vary about how students will remember its features.*
 - b. How many variables are displayed in this plot? Name the variable(s) and identify the type of variable(s). *Answer: Three variables. Weight in kilograms and height in meters are numerical variables. Sex is categorical.*
 - c. What do the axes show? *Answer: The x-axis shows the height of teens in meters, and the y-axis shows the weight of teens in kilograms.*
 - d. What questions can we ask that this graph might answer? *Possible questions are: Who is taller, boys or girls? Who weighs more? Is the association between height and weight the same for boys as it is for girls?*



9. Discuss this plot with the class by eliciting students' responses to the questions. Students actively listen to the discussion by confirming, correcting, or adding to their own responses.
10. Follow-up discussion: when the data are split into categories, it is called **grouping**.
11. Close the discussion by asking students: What questions might you have about this plot? What additional information would be helpful? *Answers will vary.*

12. Display the following plot to the class (LMR_U1_L14_C):



LMR_U1_L14_C

13. Ask students to individually answer the following questions about the plot on the handout (LMR_U1_L14_C):

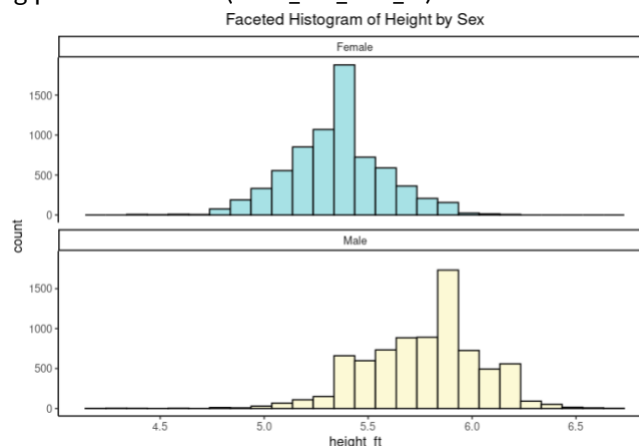
- What kind of plot is this and how will you remember its features? *Answer: Side-by-side bar chart. Answers will vary about how students will remember its features.*
- How many variables are displayed in this plot? Name the variable(s). *Answer: Two variables: whether or not someone is Hispanic, and how often they wear sunscreen.*
- What are the x-axis and y-axis telling us? *Answer: The x-axis shows how often a student wears sunscreen, and the y-axis shows the percentage of the total that fall into that category (broken into two bars, one for Hispanic and one for non-Hispanic).*
- What statistical questions can you answer with this graph? *Possible questions are: Do Hispanics and non-Hispanics have different approaches to sunscreen? What percent of Hispanics always/never wear sunscreen? How does that compare to non-Hispanics?*



14. Discuss this plot with the class by eliciting students' responses to the questions. Students actively listen to the discussion by confirming, correcting, or adding to their own responses.

15. Close the discussion by asking students: What questions might you have about this plot? What additional information would be helpful? *Answers will vary.*

16. Display the following plot to the class (LMR_U1_L14_D).



LMR_U1_L14_D

17. Ask students to individually answer the following questions about the plot on the handout (LMR_U1_L14_D):
- What kind of plot is this and how will you remember its features? *Answer: Split or faceted histogram. Answers will vary about how students will remember its features.*
 - How many variables are displayed in this plot? Name the variable(s). *Answer: Two variables; height and sex.*
 - What are the x-axis and y-axis telling us? *Answer: The x-axis shows height in feet, and the y-axis shows the total that fall into a certain range of heights (broken into two histograms, one for males and one for females).*
 - What statistical questions can you answer with this graph? *Possible questions are: Do males and females differ in height? What is the typical female height? What is the typical male height?*
18. Discuss this plot with the class by eliciting students' responses to the questions. Students actively listen to the discussion by confirming, correcting, or adding to their own responses.
19. Using the notes and sketches in their DS journals, students will summarize their learning of how to read and interpret basic multiple variable plots by completing the *Multiple Variable Plots* four-fold graphic organizer (LMR_U1_L14_E):



Name: _____ Date: _____

Summarizing Multiple Variable Plots

Numerical-Numerical (Scatterplots)	Numerical-Numerical-Categorical (Scatterplots with Grouping)
Categorical-Categorical (Side by Side Bar Graphs)	Numerical-Categorical (Faceted Histogram)

LMR_U1_L14_E

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Next 2 Days

Lab 1D: Zooming Through Data

Lab 1E: What's the Relationship?

Complete Lab 1D and 1E prior to the Practicum.

Lab 1D: Zooming Through Data

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Data with Clarity

- Previously, we've looked at graphs of entire variables (by looking at all of their values).
 - Doing this is helpful to get a *big picture* idea of our data.
- In this lab, we'll learn how to *zoom in* on our data by learning how to subset.
 - We'll also learn a few ways to manipulate the plots we've been making to make them easier to use for analyses.
- Import the data from your class's *Food Habits* campaign and name it **food**.

Another plotting function

- A **dotPlot** is another plot that can be used to analyze a numerical variable.
 - Dotplots are better suited for smaller datasets. If datasets are too large, the dots become too small to see.
 - Similarly, distributions with a large spread might impact the readability of the plot.
- **(1) Write and run code for creating a dotPlot of the amount of sugar in our food data.**
 - The code to create a **dotPlot** is exactly like you'd use to make a **histogram**.
 - Make sure to use a capital P in **dotPlot**.

More Options

- While a **dotPlot** should conserve the exact value of each data point, sometimes it behaves like a **histogram** in that it lumps values together.
- **(2) Write and run the code for a more accurate dotPlot by using the nint option.**
 - Set **nint** equal to the maximum value for **sugar** minus the minimum value for **sugar** plus one.
 - **On your food data spreadsheet, click on the sugar header to sort in ascending order (to obtain minimum).**
 - **Click on the sugar header again to sort in descending order (to obtain maximum).**
 - **Use your History pane to see how we included the option nint with the histogram function.**
- Pro-tip: If the **dotPlot** comes out looking wonky, try changing the value of the *character expansion option*, **cex**.
 - The default value is **1**. Try a few values between **0** and **1** and a few more values larger than **1**.

Splitting datasets

- In Lab 1B, we learned that we can *facet* (or split) our data based on a categorical variable.
- **(3) Write and run code splitting the dotPlot displaying the distribution of grams of sugar in two, by faceting on our observations' salty_sweet variable.**
 - **(4) Describe how R decides which observations go into the left or right plot.**
 - **(5) What does each dot in the plot represent?**

Altering the layout

- It would be much easier to compare the `sugar` levels of salty and sweet snacks if the dotPlots were stacked on top of one another.
- We can change the *layout* of our separated plots by including the `layout` option in our `dotPlot` function.
 - Add the following option to the code you used to create the dotPlot split by `salty_sweet`.

```
layout = c(1,2)
```

- Hint: Use a similar syntax used with the `nint` option to add the `layout` option to the `dotPlot` function.

Subsetting

- Subsetting is a term we use to describe the process of looking at only the data that conforms to some set of rules:
 - Geologists may subset earthquake data by looking at only large earthquakes.
 - Stock market traders may subset their trading data by looking only at the previous day's trades.
- There are *many* ways to subset data using RStudio, we'll focus on learning the most common methods.

The filter function

- Creating a plot for `Salty` and a plot for `Sweet` is useful for comparing `Salty` and `Sweet`. What if we want to examine one group by itself?
- The line of code below creates a subset of the `food` dataset containing only `Salty` snacks. We will break it down piece by piece in the next few slides.
- Run the line of code below.

```
food_salty <- filter(food, salty_sweet == "Salty")
```

- (6) View `food_salty` and write down the number of observations in it.

So what's really going on?

- Coding in R is really just about supplying directions in a way that R understands.
 - We'll start by focusing on everything to the right of the `<-` symbol

```
food_salty <- filter(food, salty_sweet == "Salty")
```

- `filter()` tells R that we're going to look at only the values in our data that follow a *rule*.
- The first blank should be the data we're going to filter down into a smaller set (based on our rule).
- `salty_sweet == "Salty"` is the rule to follow.

3 parts of defining rules

- We can decompose our rule, `salty_sweet == "Salty"`, into 3 parts:
 1. `salty_sweet`, is the particular *variable* we want to use to select our subset.
 2. `"Salty"` is the *value* of the variable that we want to select. We only want to see data with the value `"Salty"` for the variable `salty_sweet`.
 3. `==` describes how we want to relate our variable (`salty_sweet`) to our value (`"Salty"`). In this case, we want values of `salty_sweet` that are *exactly equal* to `"Salty"`.
- Notice: *Values* (that are also words) have quotation marks around them. *Variables* do not.

More on ==

- We can use the `head()` function to help us see what's happening when we write `salty_sweet == "Salty"`.
 - `head()` returns the values of the first 6 observations.
 - The `tail()` function returns the last 6 observations.
- Run the following code and answer the question below:

```
head(~salty_sweet == "Salty", data = food)
```
- (7) What do the values TRUE and FALSE tell us about how our rule applies to the first six snacks in our data? Which of the first six observations were Salty?

Saving values

- To use our subset data we need to save it first.
 - When we save something in R what we are really doing is giving a value, or set of values, a specific name for us to use later.
- The arrow `<-` is called the “assignment” operator. It assigns names (on the left) to values (on the right).
 - We now focus on everything to the left of, and including, the `<-` symbol.

```
food_salty <- filter(food, salty_sweet == "Salty")
```

Saving our subset

```
food_salty <- filter(food, salty_sweet == "Salty")
```

- This code then:
 - takes our subset data, (everything to the right of `<-`) ...
 - and assigns the subset data, by using the arrow `<-` ...
 - the name `food_salty`.
- We can now use `food_salty` to do anything we could do with the regular `food` data...
 - but only including those snacks who reported being `Salty`.
- As a result of assigning the subset data to `food_salty`, `food_salty` now appears in the *Environment* pane. Whenever data is assigned to a variable name, that variable name will appear in the *Environment* pane.
- (8) Write and run code using `food_salty` to make a dotPlot of the sodium in our Salty snacks.

Including more filters

- We often want to filter our data based on multiple rules.
 - For instance, we might want to filter our `food` data based on the food being salty AND having less than 200 calories.
- We can include multiple filters to our subsets by separating each rule with a comma like so:

```
my_sub <- filter(food, salty_sweet == "Salty", calories < 200)
```
- View the `my_sub` data we filtered in the above line of code and verify that it only includes salty snacks that have less than 200 calories.

Put it all together

- (9) Create a `dotPlot` and answer the question: About how much sugar does the typical sweet snack have?
- (10) Create a `dotPlot` and answer the question: How does the typical amount of sugar compare when `healthy_level < 3` and when `healthy_level > 3`?
- Because you are now working with subsets of data, it is important to be able to label our plots and make this distinction.
 - We can use the `main` option to add a title to our plots.
 - Add the following option to the code you used to create the `dotPlot` of the sugar in Sweet snacks.

```
main = "Distribution of sugar in sweet snacks"
```

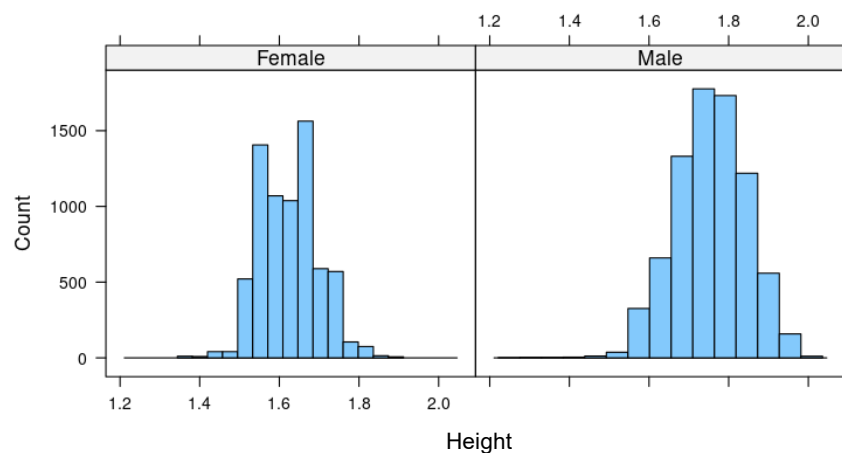
Lab 1E: What's the Relationship?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Finding patterns in data.

- To discover (*really*) interesting observations or relationships in data, we need to find them!
 - Which is difficult if we only look at the raw data.
- The best tool for finding patterns is often ... your own eyes.
 - Plots are an excellent way to help your eye search for patterns.
- In this lab, we'll learn how to include more variables in our plots to make them more informative.
- Import the data from your class's *Food Habits* campaign and name it **food**.

Where are the variables?



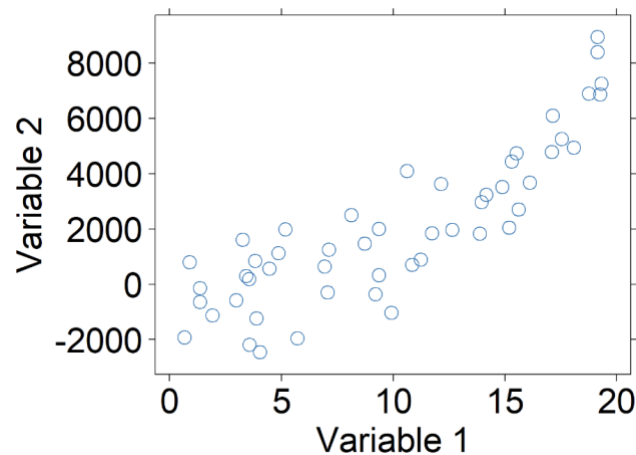
- **(1) How many variables were used to create this plot? Which variables were used and how were they used?**

Multiple variable plots

- The previous graph is an example of a *multiple variable plot*, which means that more than a single variable was used. In this case:
 - Variable 1: **height**
 - Variable 2: **sex**
- Multiple variable plots are tools for finding *relationships* between data.
- Let's take our **food** data and make some new multiple variable plots you haven't created before!

Scatterplots

- Scatterplots are useful for viewing how one *numerical* variable relates to another *numerical* variable.



Creating scatterplots

- (2) Fill in the blanks to create a scatterplot with `sodium` on the y-axis and `sugar` on the x-axis.

```
xyplot(____ ~ ____, data = food)
```

Scatterplots in action

- Use a scatterplot to answer the following questions:
 - (3) Do snacks that have more `protein` also have more `calories`? Why do you think that?
 - (4) What happens if you swap the `protein` and `calories` variables in your code? Does the relationship between the variables change?
 - (5) Does the relationship between `protein` and `calories` change when the snack is either `Salty` or `Sweet`? Write down the code you used to answer this question.

4-variable scatterplots

- When we make scatterplots, we can include:
 - 1 numerical variable on the x-axis
 - 1 numerical variable on the y-axis
 - Use 1 categorical variable to facet our scatterplot
 - Change the color of the points based on another categorical variable
- To change the color of our points, we can include the `groups` argument much like we did for bargraphs (use the `search` feature in the *History* pane if you need help).
- (6) Write and run code creating a scatterplot that uses these 4 variables: `sodium`, `sugar`, `cost`, `salty_sweet`.

Multiple facets

- It can sometimes be helpful to facet on more than 1 variable.
 - Splitting the data using 2 facets can give us additional insights that might otherwise be hidden.
- (7) Write and run code creating a dotPlot or histogram of the calories variable, but facet the data using:

```
healthy_level + salty_sweet
```
- (8) How does the healthy_level of a Salty or Sweet snack impact the number of calories in the snack?
- Although we are treating healthy_level as a categorical variable, R recognizes it as a numerical variable.
 - Verify this using the str function.
 - Notice that the faceted histograms or dotPlots do not have labels, but rather tick-marks.
 - You will have the opportunity to convert the healthy_level variable into a factor later on.
- Faceting your data on a numerical variable is NOT recommended.
 - Numerical variables often have so many different values that they overwhelm the plot and make it hard to read.

On your own

- Answer the following questions by creating an appropriate graph or graphs.
 - (9) Do healthier snacks have more or less ingredients than less healthy snacks?
 - (10) What other variables seem to be related to the number of ingredients of a snack? Describe their relationships.

Practicum: The Data Cycle & My Food Habits



Objective:

Students will apply what they have learned by engaging in the Data Cycle using the data they collected from the *Food Habits* campaign. Students will present their findings to the class.

Materials:

1. *The Data Cycle Practicum* handout (LMR_U1_Practicum_Data_Cycle)
2. Poster paper
3. Markers

Practicum The Data Cycle & My Food Habits

Instructions:

With a partner, you will engage in the Data Cycle to address the Research Topic:

What do our snacking habits reveal about us?

Task:

1. Create a Data Cycle poster.
2. The poster should illustrate how the Data Cycle is used to address the Research Topic.
3. Use RStudio to create at least one statistical graphic. The graphic **MUST** be included on the poster.
4. You and your partner will present your findings with appropriate evidence from the data.

Awards:

Your teacher will select the top posters in the following categories:

- Best Statistical Question
- Most Interesting Statistical Graphic
- Best Illustration of the Data Cycle

Scoring Guide

Below you will find some parameters to assist you in scoring. They are meant only as a guide.

4-point response:

- The poster correctly illustrates how the Data Cycle is used to address the big question.
- A histogram, bar chart, scatterplot, or other graphical representation was correctly created.
- An answer and a justification for the answer to the statistical question are presented.
- A justification includes mention of statistics concepts learned thus far. For example, "The variables are..." **AND** it includes acknowledgment of variability. For example, "There are between ____ and ____."

3-point response:

- The poster correctly illustrates how the Data Cycle is used to address the big question.
- A histogram, bar chart, scatterplot, or other graphical representation was correctly created.
- An answer and a justification for the answer to the statistical question are presented.
- A justification includes mention of statistics concepts learned thus far. For example, "The variables are..." **OR** it includes acknowledgment of variability. For example, "There are between ____ and ____."

2-point response:

- The poster partially illustrates how the Data Cycle is used to address the big question.
- A histogram, bar chart, scatterplot, or other graphical representation was created.
- An answer and a justification for the answer to the statistical question are presented.

1-point response:

- The poster incorrectly illustrates how the Data Cycle addresses the big question.
- An answer to the statistical question is presented but a justification is missing.
- A histogram, bar chart, scatterplot, or other graphical representation was correctly created.

0-point response:

- The Data Cycle is missing OR does not show how it addresses the big question.
- A histogram, a dot plot, or other graphical representation was incorrectly created OR is missing.
- No answer AND/OR no justification for the answer to the statistical question is presented.

Would You Look at the Time

Instructional Days: 9

Enduring Understandings

Data are useful for evaluating claims and reports. Summaries of categorical and numerical data show important features and patterns in the data. Data summaries provide evidence to make claims.

Engagement

The Bureau of Labor Statistics (BLS) collects data about daily time-use of Americans. Students will explore the *New York Times* multimedia graphic titled *How Different Groups Spend Their Time* to spark their curiosity about how they spend their own time. The graphic can be found at

http://www.nytimes.com/interactive/2009/07/31/business/20080801-metrics-graphic.html?_r=0.

Learning Objectives

Statistical/Mathematical:

S-ID 5: Summarize categorical data for two categories in two-way frequency tables. Interpret relative frequencies in the context of data (including joint, marginal, and conditional relative frequencies). Recognize possible associations and trends in the data.

S-IC 6: Evaluate reports based on data.

Data Science:

Understand that data are collected and stored in particular formats. Before data can be analyzed, it must be cleaned so it can be read.

Applied Computational Thinking Using RStudio:

- Create tabular displays of categorical data and summaries of numerical data.
- Create two-way frequency (and relative frequency) tables.
- Use RStudio to calculate joint, marginal, and conditional relative frequencies.
- Subset data frames and create new categorical variables from numerical variables.
- Clean and polish data to make it readable.

Real-World Connections:

Make claims that are based on data and begin to evaluate reports that make claims based on data.

Language Objectives

1. Students will describe associations between categorical variables.
2. Students will connect their time use activities with that of Americans.
3. Students will read informative texts to evaluate claims based on data.

Data File or Data Collection Method

Data Files:

1. Students' *Time-Use* campaign data
2. American Time-Use Survey (ATUS): data (atus)

Data Collection Method:

Time-Use Participatory Campaign: Students will monitor the amount of time they devote to activities such as sleeping, studying, eating, and partaking in media.

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Lesson 15: Americans' Time on Task

Objective:

Introduction to *Time Use Campaign*. Students will explore a multimedia graphic that incorporates data from the American Time Use Survey to spark their interest about how they spend their time. They will begin to learn how to evaluate reports that make claims based on data by reading the Pew Research Center's article *Teens and social media: Key findings from Pew Research Center surveys*.

Materials:

1. Computers
2. Data Collection Devices
3. Interactive multimedia graphic: *The New York Times' How Different Groups Spend Their Time* found at: <https://archive.nytimes.com/www.nytimes.com/interactive/2009/07/31/business/20080801-metrics-graphic.html>
4. Article: *Pew Research Center's Teens and social media: Key findings from Pew Research Center surveys* found at: <https://www.pewresearch.org/short-reads/2023/04/24/teens-and-social-media-key-findings-from-pew-research-center-surveys/>
5. K-L-W Graphic Organizer (LMR_TR_K-L-W Chart)
6. *Time Use Campaign* guidelines (LMR_U1_Campaign_Time_Use)

Vocabulary:

evaluate, claim

Essential Concepts: Learning to examine other analyses is an important part of statistical thinking.

Lesson:

1. Become familiar with the *Time Use Campaign Guidelines* (shown at the end of this lesson), particularly the big questions, to help guide students during the campaign (see Campaign Guidelines in Teacher Resources).
2. In pairs, ask students to make predictions based on the big questions in the *Time Use Campaign Guidelines*.
3. Next, inform students that *The Bureau of Labor Statistics* (BLS) collects data about Americans' daily time use and that they will be exploring time use through an interactive graphic.
4. Ask students to go to the multimedia graphic at the following URL:
<https://archive.nytimes.com/www.nytimes.com/interactive/2009/07/31/business/20080801-metrics-graphic.html>
5. Students will spend 10 minutes exploring the interactive graphic. Their task is to answer the following questions (display questions to students):
 - a. What variables are represented in this graphic? *Answer: The variables represented are activities that Americans spend their time doing. These include sleeping, eating, traveling, socializing, etc.*
 - b. Explain what the graphic is telling you. *Answer: The graphic shows how much time Americans over the age of 15 are spending doing these activities. This information is broken down by different categories of Americans (e.g., sex, ethnicity) and the percentage of Americans doing particular activity at a particular time (e.g., 5% of Americans are working at 6:00 am). The average time spent on a particular activity is also shown (e.g., average time spent at work for all Americans is 3 hours and 25 minutes).*
 - c. Where did the data come from? *Answer: The data come from thousands of Americans over the age of 15 who took a survey recalling every minute of a day in 2008.*
 - d. What are some interesting findings? Be prepared to share. *Answers will vary.*



6. Ask students to share their findings in pairs. Each pair will agree on and select one finding to share with the class. In a *Whip Around*, ask each pair to share their finding.
7. Inform students that they will continue to investigate Americans' daily time use. Using the K LW graphic organizer, read out loud the title of the *Pew Research Center* article: *Teens and social media: Key findings from Pew Research Center surveys*. Ask them to write what they know about the topic in the Know column.

Note to Teacher: If this is the first time using K LW, please take time to provide an overview of the graphic organizer.



8. Next, ask students to read the article individually: <https://www.pewresearch.org/short-reads/2023/04/24/teens-and-social-media-key-findings-from-pew-research-center-surveys/>
9. As they read, students may complete the Learn column of the K LW graphic organizer.
10. Ask students to complete the Want to Learn column when they finish reading the article.
11. When reading a newspaper, magazine, or blog that includes statistical analysis, it is important to **evaluate**, or think carefully, about **claims** that these articles state as fact.
12. Ask students to work in teams to evaluate the article based on the questions below:
 - a. Who was observed and what were the variables observed? *Answer: 1,316 U.S. teens (13 to 17-year-olds), and one parent from each household, were surveyed. The variables had to do with experiences with social media – time spent on YouTube, TikTok, Snapchat, Instagram, Facebook, and the effects it had on them.*
 - b. What statistical questions were they trying to answer? *Possible statistical question: How much time per day does today's typical 13 to 17-year-old spend on TikTok?*
 - c. Who collected the data? *Answer: The Pew Research Center conducted the study, but the survey was conducted online by Ipsos.*
 - d. How was the data collected? *Answer: Online by Ipsos via a survey from April 14 to May 4, 2022.*
 - e. What claim(s) did the article make? *Answer: There were 11 claims made in the article. The first claim: "Majorities of teens report ever using YouTube, TikTok, Instagram, and Snapchat."*
 - f. What are some statistics that the article used to make the claim(s)? *Examples include: Two-thirds of teens report using TikTok, followed by roughly six-in-ten who say they use Instagram (62%) and Snapchat (59%). Six-in-ten teens say they think they have little (40%) or no control (20%) over the personal information that social media companies collect about them. Some 22% of teens think their parents are extremely or very worried about them using social media.*



13. Select a whole group share-out/discussion strategy from the Instructional Strategies Teacher Resource to discuss the answers to the evaluation questions.
14. Inform students that they will engage in the Time Use Participatory Sensing campaign (LMR_U1_Campaign_Time_Use) and will begin to collect data about their own time use.
Reminder: Once logged into the app or the browser-based version, students may go to **Campaigns** to see the campaigns in which they are participating. They can then add the campaign by tapping the name of the campaign. If no campaigns are visible, ask them to click the campaign refresh option.
15. Emphasize that this data will be tracked throughout the day via a log of some sort – it might be helpful to split the log into three intervals where students pause and think about what they did before school, after school and in the evening. Once the log is complete/accounts for all 1,440 minutes of their day, students should then submit the survey corresponding to that day. They will keep a log for at least 5 days (of which 2 days include a weekend) but no more than 10 days.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Campaign Guidelines – Time Use

1. The Issue:

There have been many reports lately about people spending a large amount of time interacting with technology and the Internet. This raises some questions about time use:

- 1) How do I spend my time?
- 2) Are there groups that spend more time on certain activities than other groups?
- 3) How is my time use similar or different to other Americans?

2. Objectives:

Upon completing this campaign, students will have compared themselves to the U.S. population to find how they are similar to and/or different from other people in terms of time use. They will use single and multivariable plots, summary statistics, and frequency tables to find similarities and differences between groups of students, and between students and other residents of the United States.

3. Survey Questions: (Students will enter data for the activities in which they participate.)

Consider Data: The activity categories students will use to categorize how they spend their time are similar to those employed in the American Time Use Survey (ATUS). The ATU Survey offers nationally representative insights into how Americans allocate their time. Before students begin collecting data, it is important to discuss different activities in their day and how they might be classified. A class consensus of the meaning of the variables must be reached so that proper analysis and interpretations can be made.

Note: Students cannot double dip their time. For example, if they read during class, then those minutes spent reading do not count towards “read” but instead toward “school”.

Students will respond to the following questions:

Prompt	Variable	Data Type
For which day are you collecting data?	day	ordinal category (integers 1-10)
What activities did you participate in?	activities	n/a
a. How many MINUTES did you sleep?	sleep	number
b. How many MINUTES did you spend eating/drinking?	meals	number
c. How many MINUTES did you spend in classes at school?	school	number
d. How many MINUTES did you spend doing homework?	homework	number
e. How many MINUTES did you spend working at a job?	work	number
f. How many MINUTES did you spend grooming yourself?	grooming	number
g. How many MINUTES did you spend traveling/commuting?	travel	number
h. How many MINUTES did you spend doing household chores?	chores	number
i. How many MINUTES did you spend watching television (includes streaming)?	television	number
j. How many MINUTES did you spend playing video games?	videogames	number
k. How many MINUTES did you spend participating in sports/exercise/physical activity?	sports	number
l. How many MINUTES did you spend reading (not for class)?	read	number
m. How many MINUTES did you spend communicating (includes texting, emails, video and voice calls)?	communicate	number

n. How many MINUTES did you spend socializing (outside of class, in person)?	socialize	number
o. How many MINUTES did you spend on a spiritual activity?	spiritual	number
p. How many MINUTES did you spend purchasing items online or in a store?	purchases	number
q. How many MINUTES did you spend on hobbies/volunteering/leisure/extra-curricular activities (excluding sports and physical activity)?	extra	number
r. How many MINUTES did you spend on social media?	social_media	number
AUTOMATIC	location	lat, long
AUTOMATIC	time	time
AUTOMATIC	date	date

When should you take the survey? It is recommended that students keep a log of their time and submit one survey at the end of each day, accounting for every minute of each day of the campaign. It might be helpful to split the log into three intervals where students pause and think about what they did before school, after school and in the evening. Once the log is complete and accounts for all 1,440 minutes of their day, students should then submit the survey corresponding to that day.

How long? At least five days (maximum of ten days). Ideally, two of these days would include a weekend.

4. Motivation:

Use the

<https://archive.nytimes.com/www.nytimes.com/interactive/2009/07/31/business/20080801-metrics-graphic.html> Interactive Time Use graphic to explore how Americans spend their time.

After the first day, monitor the data collection and ensure that each student has submitted a survey for Day 1.

Discuss data collection issues. What makes it hard? Does this affect the quality of data?

5. Technical Analysis:

RStudio and American time use interactive graphic.

Single/Multivariable plots: histograms, bargraphs, scatterplots, etc.

Numerical summaries: mean, median, MAD, standard deviation.

Frequency tables: One and two-way tables.

6. Guiding Questions:

- 1) On average, how long do students think they spend on homework?
- 2) Are certain activities that take up most of the time in our day?
- 3) Are there groups of students who spend their time similarly to one another?

7. Report:

Students will complete a practicum in which they answer a statistical question based on the Time Use data collected.

Homework & Next Day

For the next 5 days, students will collect data using the Time Use campaign on their smart devices or via web browser.

Lab 1F: A Diamond in the Rough and

Data Collection Monitoring

1. **Data Collection Monitoring:** Display the IDS Campaign Monitoring Tool, found at <https://portal.thinkdataed.org/>. Click on **Campaign Monitor** and sign in.
 - a. See *User List* and sort it by *Total*. Ask: Who has collected the most data so far?
 - b. Click on the pie chart. Ask: How many active users are there? How many inactive users are there?
 - c. See *Total Responses*. Ask: How many responses have been submitted?
 - d. Using TPS, ask students to think about what they can do to increase their data collection.
2. Inform students that you will conduct another data collection check with the whole class in a couple of days, and that they will understand the private vs. shared data after they have completed the campaign collection.

Complete Lab 1F prior to Lesson 16.

Lab 1F: A Diamond in the Rough

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Messy data? Get used to it

- Since lab 1, the data we've been using has been pretty *clean*.
- Why do we call it *clean*?
 - Variables were named so we could understand what they were about.
 - There didn't seem to be any *typos* in the values.
 - Numerical variables were considered numbers.
 - Categorical variables were composed of categories.
- Unfortunately, more often than not, data is *messy* until YOU clean it.
- In this lab, we'll learn a few essentials for cleaning *dirty* data.

Messy data?

- What do we mean by messy data?
- Variables might have *non-descriptive names*
 - *Var01*, *V2*, *a*, ...
- Categorical variables might have *misspelled categories*
 - *"blue"*, *"Blue"*, *"blu"*, ...
- Numerical variables might have been *input incorrectly*. For example, if we're talking about people's height in inches:
 - *64.7*, *6.86*, *676*, ...
- Numerical variables might be *incorrectly coded* as categorical variables (or vice-versa).
 - *"64.7"*, *"68.6"*, *"67.6"*

The American Time Use Survey

- To show you what *dirty* data looks like, we'll check out the *American Time Use Survey*, or *ATU* survey.
- What is ATU survey?
 - It's a survey conducted by the US government (specifically the Bureau of Labor Statistics).
 - They survey thousands of people to find out exactly what activities they do throughout a single day.
 - These thousands of people combined together give an idea about how much time the typical person living in the US spends doing various activities.

Load and go:

- **Type the following commands into your console:**

```
data(atu_dirty)
View(atu_dirty)
```

- **(1) Just by viewing the data, what parts of our ATU data do you think need cleaning?**

Description of ATU Variables

- The description of the actual variables:
 - `caseid`: Anonymous ID of survey taker.
 - `V1`: The age of the respondent.
 - `V2`: The sex of the respondent.
 - `V3`: Whether the person is employed full-time or part-time.
 - `V4`: Whether the person has a physical difficulty.
 - `V5`: How long the person sleeps, in minutes.
 - `V6`: How long the survey taker spent on homework, in minutes.
 - `V7`: How long the respondent spent socializing, in minutes.

New name, same old data

- To fix the variable names, we need to *assign* a new set of names in place of the old ones.
 - Below is an example of the `rename` function:

```
atu_cleaner <- rename(atu_dirty, age = V1,  
                      sex = V2)
```

- Use the example code and the variable information on the previous slide to rename the rest of the variables in `atu_dirty`.
- (2) Write down the new names you chose for the rest of the variables in `atu_dirty`.
 - Names should be short, contain no spaces and describe what the variable is related to. So use abbreviations to your heart's content.

Next up: Strings

- In programming, a *string* is sort of like a *word*.
 - It's a value made up of *characters* (i.e. letters)
- The following are examples of strings. Notice that each *string* has quotes before and after.

```
"string"  
"A1B2c3"  
"Hot Cocoa"  
"0015"
```

Numbers are words? (Sometimes)

- In some cases, `R` will treat values that look like *numbers* as if they were *strings*.
- Sometimes we do this on purpose.
 - For example, we can code `Yes/No` variables as `"1"/"0"`.
- Sometimes we don't mean for this to happen.
 - The *number of siblings* a person has should not be a string.
- Look at the structure of your data and the variable descriptions from a few slides back:
 - (3) Write down the variables that should be *numeric* but are improperly coded as *strings* or *characters*.

Changing strings into numbers

- To fix this problem, we need to tell R to think of our “*numeric*” variables as numeric variables.
- We can do this with the `as.numeric` function.
 - An example using this function is below:

```
as.numeric("3.14")  
## [1] 3.14
```

- Notice: We started with a string, "3.14", but `as.numeric` was able to turn it back into a number.

Mutating in action

- (4) Look at the variables you thought should be *numeric* and select one. Then fill in the blanks below to see how we can correctly code it as a number:

```
atu_cleaner <- mutate(atu_cleaner,  
  age = as.numeric(age),  
  ____ = as.numeric(____))
```

- Once you have this code working, use a similar line of code to correctly code the other *numeric* variables as numbers.

Deciphering Categorical Variables

- We mentioned earlier that we sometimes code categorical variables as numbers.
 - For example, our `sex` variable uses "01" and "02" for "Male" and "Female", respectively.
- It's often much easier to analyze and interpret when we use more descriptive categories, such as "Male" and "Female".

Factors and Levels

- R has a special name for *categorical* variables, called *factors*.
- R also has a special name for the different *categories* of a *categorical* variable.
 - The individual categories are called *levels*.
- To see the levels of `sex` and their counts, type:

```
tally(~sex, data = atu_cleaner)
```

- (5) Use similar code as we used above to write down the levels for the three factors in our data.

A level by any other name...

- If we know that '01' means 'Male' and '02' means 'Female' then we can use the following code to recode the *levels* of `sex`.
- Type the following command into your console:

```
atu_cleaner <- mutate(atu_cleaner,  
  sex = recode(sex,  
    "01" = "Male",  
    "02" = "Female"))
```

- This code is definitely a bit of a mouthful. Let's break it down.

Allow me to explain

```
atu_cleaner <- mutate(atu_cleaner,  
                      sex = recode(sex,  
                                   "01" = "Male",  
                                   "02" = "Female"))
```

- This code is saying:
 - Replace my current version of `atu_cleaner`...
 - with a mutated one where ...
 - the `sex` variable's levels ...
 - have been recoded...
 - where "01" will now be "Male"...
 - and "02" will now be "Female".

Finish it off!

- **Recode the categorical variable about whether the person surveyed had a physical challenge or not. The coding is currently:**
 - "01": Person surveyed *did not* have a physical challenge.
 - "02": Person surveyed *did* have a physical challenge.
- **Write a script that:**
 1. Loads the `atu_dirty` dataset
 2. Cleans the data as we have in this lab
 3. Saves a copy of the cleaned data (see next slide).
- **NOTE: You can watch this video to learn about RScripts:**
<https://www.youtube.com/watch?v=OPqjL9AzmkE>

The final lines

- The last few lines of your script are extremely important because they will save all your work.
- Be sure to **View** your data and check its `structure` to make sure it looks clean and tidy before saving.
- **Run the code below:**

```
atu_clean <- atu_cleaner
```

- This code will create a new data frame in your *Environment* called `atu_clean` which is a final copy of `atu_cleaner`.
 - If `atu_clean` is swept from your *Environment* all of the changes you made will NOT be saved.
 - You would need to re-run the script to clean the data again.
- To permanently save your changes you need to save the file as an R data file or `.Rda`.
- **Run the code below:**

```
save(atu_clean, file = "atu_clean.Rda")
```

- Look in your *Files* pane for the `atu_clean.Rda` file.
 - This is a permanent copy of your clean `atu` data.
 - To load the data onto your *Environment* click on the file.
 - A pop-up window confirming the upload will appear.

Flex your skills

- Now that you have learned some cleaning data basics, it's time to revisit the `food` data.
 - Import your food data onto the *Environment* pane.
- Run the code below:

```
histogram(~calories | healthy_level, data = food)
```
- (6) Write and run code using the `as.factor()` function to convert `healthy_level` into a categorical variable and re-run the `histogram` function.
 - Notice that the `healthy_level` categories are now numbers as opposed to tick-marks. This is an improvement but an even better solution would be to `recode` the categories.
- (7) Write and run code to `recode` the `healthy_level` categories and re-run the `histogram` function.
 - “1” = “Very Unhealthy”
 - “2” = “Unhealthy”
 - “3” = “Neutral”
 - “4” = “Healthy”
 - “5” = “Very Healthy”
- If your `food` data is cleared from your *Environment*, the changes that you made to the `healthy_level` variable will not be saved.
- To save your changes permanently, save your `food` file as an R data file.

Lesson 16: Categorical Associations

Objective:

Students will learn to construct, interpret, and calculate the joint relative frequencies of two-way frequency tables.

Vocabulary:

two-way frequency table, relative frequency

Essential Concepts: A two-way table is a summary of the association/relationship between two categorical variables. Joint relative frequencies answer questions of the form “What proportion of the people/objects had *this* value on the first variable and *this* value on the second?”

Lesson:

1. Launch the lesson by displaying the following scenario:

Rosa has a theory that musicians are more likely to own a cat. To find out, she decided to collect data that would help her understand the relationship between instrument playing and cat ownership among the students in her art class. She conducted a survey and found that out of the 35 students in her art class, 16 owned a cat and out of those that owned a cat, 7 played an instrument. She also discovered that 9 owned a cat but did not play an instrument. There were also 9 students who neither owned a cat nor played an instrument.

2. Inform students that Rosa asked two questions that provided the data for her two-way frequency table. What could those two questions be?

Possible Answer: Question 1 – Do you play an instrument?

Question 2 – Do you own a cat?

3. What variables did Rosa collect? What were the values of those variables?

Answer: Variable 1: Owns Cat – yes/no.

Variable 2: Plays Instrument – yes/no.

4. In pairs, write out on paper what the original data must have looked like.

Possible answer:

<i>Owns Cat</i>	<i>Plays Instrument</i>	
<i>Yes</i>	<i>Yes</i>	<i>(There are 7 of these)</i>
<i>Yes</i>	<i>No</i>	<i>(9 of these)</i>
<i>No</i>	<i>Yes</i>	<i>(10 of these)</i>
<i>No</i>	<i>No</i>	<i>(9 of these)</i>

5. Inform students that today they will be looking at associations in categorical variables.

Note: If necessary, review the difference between questions for categorical and numerical variables.

6. Explain that their task is to summarize Rosa’s findings in one table that shows totals. Remind students to use their knowledge of data structures from Lesson 2, especially organizing in rows and columns.
7. Allow time for teams to wrestle with how to organize their data in one table. As teams work, walk around monitoring their data tables.
8. Select a few data tables to display and share with the entire class.
9. Explain the *Anonymous Author* strategy to students (see Instructional Strategies in Teacher Resources).
10. Display the data tables and ask students to engage in the *Anonymous Author* strategy. You may want to start the discussion by asking about the total number of students Rosa surveyed.

11. Make sure the last data table you display shows a correctly structured **two-way frequency table**. A two-way frequency table displays the data that pertains to two categories from one group. One category is represented in rows and the other is represented in columns. In this exercise, the group is a class of art students.

Cat Ownership and Instruments

	Plays an instrument	Does not play instrument	Total
Owens a cat	7	9	16
No cats	10	9	19
Total	17	18	35



12. Based on the Cat Ownership and Instruments table, ask student teams to generate questions that can be asked and answered by the data.
13. In a *Whip Around*, ask student teams to share one of their questions.
14. Explain that a two-way frequency table can show relative frequencies. A **relative frequency** is how often something occurs in relation to the total number of occurrences and is expressed as a proportion or percentage of a total. For example, what is the relative frequency of those who own a cat and play an instrument? *Answer: 7/35 or 0.2 or 20%.*
- Note:** Review how to write a proportion and how to express a proportion as a percent.
15. Ask students to calculate the relative frequencies for the entire table. They may check their calculations with a partner.

Cat Ownership and Instruments

Relative Frequencies

	Plays an instrument	Does not play instrument	Total
Owens a cat	$7/35 = 0.20$	$9/35 \approx 0.26$	$16/35 \approx 0.46$
No cats	$10/35 \approx 0.29$	$9/35 \approx 0.26$	$19/35 \approx 0.54$
Total	$17/35 \approx 0.49$	$18/35 \approx 0.51$	$35/35 = 1.00$

16. In teams, students will generate 2 questions about 2 categorical variables. Allow teams 3-5 minutes to generate their questions. They should choose two categorical variables that they predict might be associated.
17. The team will create a two-way table that corresponds to their categorical variables. They will be surveying their fellow classmates in the next lesson in order to fill out their table.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework



Rosa posed this statistical question:

What proportion of students did not play an instrument and did not own a cat? Answer: 0.26

Use what you know about two-way tables to answer her question.

Lesson 17: Interpreting Two-Way Tables

Objective:

Students will calculate conditional, marginal, and joint frequencies and explain what they mean in the context of the data.

Materials:

1. Poster paper
2. Markers
3. *Analyzing Categorical Variables* (LMR_U1_L17_A)
Advance preparation required: Needs to be cut into sets (see step 23 in lesson)
4. *Interpreting Categorical Variables* (LMR_U1_L17_B)

Vocabulary:

marginal (relative) frequency, joint (relative) frequency, conditional relative frequency

Essential Concepts: Marginal (relative) frequencies tell us about the distribution of a single variable. Conditional relative frequencies tell us about the distribution of one variable when “subsetting” the other.

Lesson:

1. **Time Use Campaign Data Collection Monitoring:**
 - a. Display the IDS Campaign Monitoring Tool, found at <https://portal.thinkdataed.org/>. Click on **Campaign Monitor** and sign in.
 - b. Inform students that you will be monitoring their data collection again today.
 - i. See *User List* and sort it by *Total*. Ask: Who has collected the most data so far?
 - ii. Click on the pie chart. Ask: How many active users are there? How many inactive users are there?
 - iii. See *Total Responses*. Ask: How many responses have been submitted?
 - iv. Using TPS, ask students to think about what they can do to increase their data collection.
 - c. Remind students that this is the last day to collect data.
2. Ask student teams to take out the 2 questions and the two-way table that they created in the previous day's lesson.
3. Before teams ask the class their questions, ask them to strategize about how they will collect and record their data, because they can only ask the 2 questions.
4. Students in the class will respond to each question by raising their hands.
5. In a *Whip Around*, have each team ask their 2 questions. Pause briefly between teams so that the asking team has time to collect and record their data.
6. Students will use their frequency tables before the end of the lesson.
7. Recall that in the previous lesson, students learned to calculate relative frequencies.
8. Display the *Cat Ownership and Instruments* table:

Cat Ownership and Instruments

	Plays an instrument	Does not play instrument	Total
Owens a cat	7	9	16
No cats	10	9	19
Total	17	18	35

9. Suppose that we want to know the following information (display questions):
- How many students own a cat and play an instrument? *Answer: 7*
 - What is the proportion of students that own a cat and play an instrument? *Answer: $7/35 = 0.2$*
 - What is the proportion of students who do not own a cat and do not play an instrument? *Answer: $9/35 \approx 0.257$*
10. In teams, discuss where on the table you would find this information and how you would calculate it. The questions should look familiar because they are the same type of questions asked in the previous lesson. The specific answers are not important; but what is important is to know how to obtain the information. *Possible answer: You would find the answers in the cells that make up the body of the table.*
11. After a few minutes, ask a team to volunteer a response. Mark up the cells in the body of the table to show that the cells with the initial counts are called **joint frequencies**, like in 9a. When you take the value of each cell over the total number of observations (the proportion), these are the **joint relative frequencies**, like in 9b and c. *See example:*

Cat Ownership and Instruments

	Plays an instrument	Does not play instrument	Total
Owens a cat	7	9	16
No cats	10	9	19
Total	17	18	35

12. There are multiple types of relative frequencies. Let's look at other ways of understanding a two-way frequency table using another type of relative frequency.
13. Suppose that we want to know the following information (display questions):
- How many students own a cat? *Answer: 16*
 - What is the proportion of students who own a cat? *Answer: $16/35 \approx 0.46$*
 - What is the proportion of students who do not play an instrument? *Answer: $18/35 \approx 0.51$*



14. In teams, discuss where on the table you would find this information and how you would calculate it. The specific answers are not important; but what is important is to know how to obtain the information. *Possible answer: You would find the proportion of students who do not play an instrument by dividing the number in the "Total" row that is in the "Does not play instrument" column by the total number of students (35).*
15. After a few minutes, ask a team to volunteer a response. Mark up the margins on the table to show that the cells with the initial total counts are called **marginal frequencies**, like in 13a. When you take the total of each row (or column) over the total number of observations, these are the **marginal relative frequencies**, like in 13b and c. *See example:*

Cat Ownership and Instruments

	Plays an instrument	Does not play instrument	Total
Owens a cat	7	9	16
No cats	10	9	19
Total	17	18	35

16. Suppose that we wanted to answer the question: Of the students who play an instrument, what proportion own a cat?

17. In teams, discuss where on the table you would find this information and how you would calculate it. The specific answers are not important; but what is important is to know how to obtain the information. *Possible answer: You would first find the “Total” number of students who play an instrument (17). Then, to find the proportion of those students who own a cat, you would look at how many students own a cat within that column (7) and divide it by the “Total” who “Plays an instrument” (7/17).*



18. After a few minutes, ask a team to volunteer a response. Encourage students to agree or disagree with the explanations provided. Lead students to see that the total for the “Plays an instrument” column is important because we are only concerned with that subset of the group. Mark up the “Plays an instrument” column on the table to show that we have conditioned, or are bound, by this variable. This is a **conditional relative frequency by column**.

Cat Ownership and Instruments
Conditional Relative Frequencies by Column

	Plays an instrument	Does not play instrument	Total
Owens a cat	$7/17 \approx 0.41$	$9/18 = 0.50$	$16/35 \approx 0.46$
No cats	$10/17 \approx 0.59$	$9/18 = 0.50$	$19/35 \approx 0.54$
Total	$17/17 = 1.0$	$18/18 = 1.0$	35

19. Finally, suppose that we wanted to answer the question: Do a greater proportion of students in Rosa’s art class who do not play an instrument own a cat than those who do play an instrument?
20. In teams, discuss where on the table you would find this information and how you would calculate it. The specific answers are not important; but what is important is to know how to obtain the information. *Possible answer: You would need to find an additional proportion. You would first find the “Total” number of students who do not play an instrument (18). Then, to find the proportion of those students who own a cat, you would look at how many students own a cat within that column (9) and divide it by the “Total” who “Does not play an instrument” (9/18).*



21. After a few minutes, ask a team to volunteer a response. Encourage students to agree or disagree with the explanations provided. Lead students to see that the total for the “Does not play instrument” column is important because we are now also concerned with that subset of the group. Mark up the “Does not play instrument” column on the table to show that we have conditioned, or are bound, by this variable. Compare the values that show the **conditional relative frequency** for each column, “Does not play instrument” and “Plays an instrument”. A higher proportion of students who do not play an instrument own a cat than students who do play an instrument.

Cat Ownership and Instruments
Conditional Relative Frequencies by Column

	Plays an instrument	Does not play instrument	Total
Owens a cat	$7/17 \approx 0.41$	$9/18 = 0.50$	$16/35 \approx 0.46$
No cats	$10/17 \approx 0.59$	$9/18 = 0.50$	$19/35 \approx 0.54$
Total	$17/17 = 1.0$	$18/18 = 1.0$	35

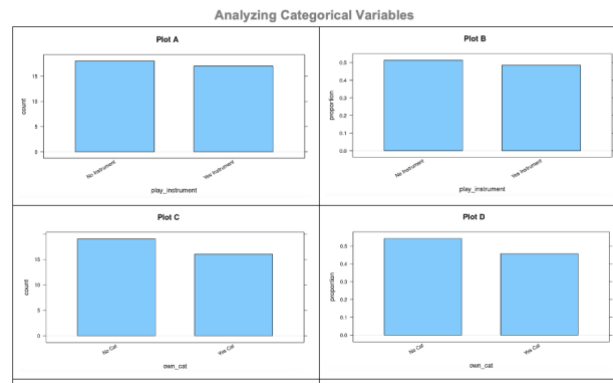
Note: These are conditional relative frequencies by column. We can also calculate conditional relative frequencies by row if we were interested in knowing the difference in playing an instrument for those who own cats versus those who do not.

22. Was Rosa right? *Answer: Looking back at our scenario introduced in the last lesson, “Rosa has a theory that musicians are more likely to own a cat,” our analysis in steps 19-21 above does not provide evidence to support her theory. We found that a higher proportion of students who do not play an instrument own a cat than students who do play an instrument.*

23. Distribute one full set of cards from the *Analyzing Categorical Variables* file (LMR_U1_L17_A) to each student team.

Advance preparation required:

Print the *Analyzing Categorical Variables* file (LMR_U1_L17_A). The handouts can then be cut into a total of 20 cards (12 visuals, 8 numerical summaries). You will need enough sets of the cards for each student team to share one full set. For example, if there are 5 student teams in a class, then 5 copies of the file will need to be printed so that each team gets all 20 cards.



LMR_U1_L17_A

24. Distribute the *Interpreting Categorical Variables* handout (LMR_U1_L17_B) to each student team.

Interpreting Categorical Variables					
Instructions: Decide which visualization(s) and numerical summaries can be used to answer each statistical question. Then answer each statistical question, citing a numerical summary as evidence.					
Statistical Question	What is the proportion of students who do not play an instrument?	How many students neither own a cat nor play an instrument?	Do a greater proportion of students in Rosa's art class who own cats prefer to play an instrument than those who do not own cats?	Is there a difference in cat preference for those who play instruments versus those who don't?	
Visualization					
Numerical Summaries					

LMR_U1_L17_B

25. Each student team will work together and decide which visualization(s) and numerical summaries can be used to answer each statistical question. They will then answer each statistical question, citing a numerical summary as evidence.

Note: Student teams may tape or glue visuals and numerical summaries onto LMR_U1_L17_B, or they can simply write the plot letter and table number in the appropriate box. The blank column is for student teams to write a statistical question than can be answered with a visual and a numerical summary that was not used.



26. After student teams have been allotted ample time to complete LMR_U1_L17_B, lead a class discussion to go over the answers. It is extremely important to have students justify their answers by referring to their visuals and tables. For example, the statistical question “How many students neither own a cat or play an instrument?” can be answered with Plot E, Plot G, Plot I, Plot K, and with Tables 1 and 7.



27. Ask students to refer back to the two-way frequency tables they created earlier. Have each team create one poster that shows their two-way frequency table. Then, ask each team to ask 5 questions about the data in their table that must be answered by a:

- marginal frequency
- marginal relative frequency
- joint frequency
- joint relative frequency
- conditional relative frequency (either by row or column)

28. If time permits, pair teams up and ask them to present their findings to each other.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next Day



Using the data below, generate 2 questions: one must be answered with a marginal relative frequency and the other must be answered by a conditional relative frequency.

Work Preference and the Color Red

Which emotion do you most relate with the color red?

	Love	Anger	Fear	Total
Morning shift	7	11	5	23
Evening shift	12	15	10	37
Total	19	26	15	60

Lab 1G: What's the FREQ?

Complete Lab 1G prior to the Practicum.

Lab 1G: What's the FREQ?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Clean it up!

- In Lab 1F, we saw how we could *clean* data to make it easier to use and analyze.
 - You cleaned a small set of variables from the American Time Use (ATU) survey.
 - The process of cleaning and then analyzing data is very common in Data Science.
- In this lab, we'll learn how we can create frequency tables to detect relationships between categorical variables.
 - **For the sake of consistency, rather than using the data you cleaned, you will use the pre-loaded ATU data.**
 - **(1) Write and run code using the `data()` function to load the `atu_clean` data file to use in this lab.**

How do we summarize categorical variables?

- When we're dealing with categorical variables, we can't just calculate an *average* to describe a *typical* value.
 - (Honestly, what's the average of categories *orange*, *apple* and *banana*, for instance?)
- When trying to describe categorical variables with numbers, we calculate *frequency tables*.

Frequency tables?

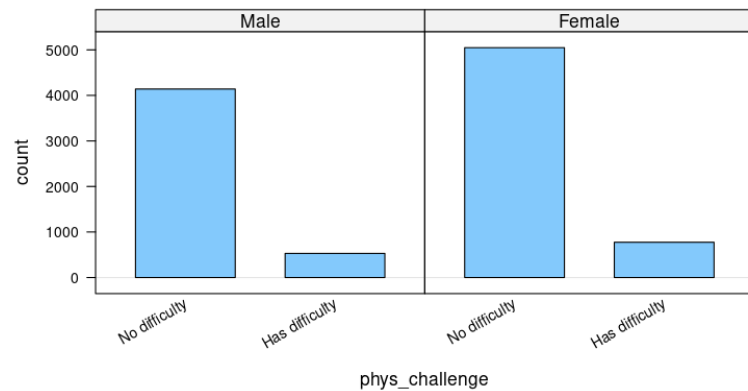
- When it comes to categories, about all you can do is *count* or *tally* how often each category comes up in the data.
- **(2) Fill in the blanks below to answer the following: How many more *females* than *males* are there in our ATU data?**

```
tally(~ ____, data = ____)
```

2-way Frequency Tables

- Counting the categories of a single variable is nice, but often times we want to make comparisons.
- For example, what if we wanted to answer the question:
 - Does one **sex** seem to have a higher occurrence of physical challenges than the other?
- We could use the following plot to try and answer the question:

```
bargraph(~phys_challenge | sex, data = atu_clean)
```



- The split `bargraph` helps us get an idea of the answer to the question, but we need to provide precise values.
- (3) Write and run a line of code, that's similar to how we facet plots, to obtain a tally of the number of people with physical challenges and their sex.
 - (4) Write down the resulting table.

Interpreting 2-way frequency tables

- Recall that there were 1153 more women than men in our dataset.
 - If there are more women, then we might expect women to have more physical challenges (compared to men).
- Instead of using *counts* we use *percentages*.
- Include: `format = "percent"` as an option to the code you used to make your 2-way frequency table.
 - (5) Does one sex seem to have a higher occurrence of physical challenges than the other? If so, which one and explain your reasoning.
- It's often helpful to display totals in our 2-way frequency tables.
 - To include them, include `margin = TRUE` as an option in the `tally` function.

Conditional Relative Frequencies

- There is a difference between `sex | phys_challenge` and `phys_challenge | sex!`

```
tally(~sex | phys_challenge, data = atu_clean, margin = TRUE)
```

```
##           phys_challenge
##    sex      No difficulty  Has difficulty
##    Male           4140           530
##    Female          5048           775
##    Total           9188          1305
```

```
tally(~phys_challenge | sex, data = atu_clean, margin = TRUE)
```

```
##           sex
##    phys_challenge  Male  Female
##    No difficulty   4140   5048
##    Has difficulty   530   775
##    Total           4670   5823
```

Conditional Relative Frequencies, continued

- At first glance, the two-way frequency tables might look similar (especially when the `margin` option is excluded). Notice, however, that the totals are different.
- The totals are telling us that `R` calculates conditional frequencies by column!
- What does this mean?
 - In the first two-way frequency table the groups being compared are the people with `No difficulty` and `Has difficulty` on the distribution of `sex`.
 - In the second two-way frequency table the groups being compared are `Male` and `Female` on the distribution of physical challenges.
- (6) Add the option `format = "percent"` to the first `tally` function. How were the percents calculated? Interpret what they mean.

On your own

- (7) Describe what happens if you create a 2-way frequency table with a numerical variable and a categorical variable.
- (8) How are the types of statistical investigative questions that 2-way frequency tables can answer different than 1-way frequency tables?
- (9) Which `sex` has a higher rate of *part time employment*?

Practicum: Teen Depression



Objective:

Using the CDC dataset, students will apply their learning of statistical concepts to determine possible factors that might be associated with depression in teens. They will create graphical representations to analyze and interpret the data. Students will present their findings to their teams the following day.

Materials:

1. *Teen Depression Practicum* (LMR_U1_Practicum_Depression)
2. *NIMH Mental Health Sheet* (LMR_U1_Practicum_NIMH_Sheet)
3. Poster paper
4. Markers

Practicum Teen Depression

Background:

The Centers for Disease Control and Prevention (CDC) collect data about teenagers on a variety of topics. One of these topics is depression. According to the *Children and Mental Health* sheet published by the National Institute for Mental Health, depression is a real problem among teens.

Instructions:

With a partner, you will read the *Children and Mental Health* sheet and then use the CDC data to address the Research Topic.

Research Topic:

What factors are associated with depression in teens?

Task:

1. Create a poster that addresses the Research Topic.
2. Generate a statistical question that might address the Research Topic.
3. Use RStudio to create at least one statistical graphic. The graphic **MUST** be included on the poster.
4. You and your partner will present your findings with appropriate evidence from the data.

Awards:

Your teacher will select the top posters in the following categories:

- Best Statistical Question
- Most Interesting Statistical Graphic

Scoring Guide:

4-point response:

- The poster identifies possible factors that are in the dataset that might be associated with depression in teens.
- A graphical representation that shows an association was created.
- An answer and a justification for the answer to the statistical question are presented.
- A justification includes mention of statistics concepts learned thus far. For example, "The variables are..."; **AND** it includes acknowledgment of variability. For example, "There are between ____ and ____."

3-point response:

- The poster identifies possible factors in the dataset that might be associated with depression in teens.
- A graphical representation that shows an association was created.
- An answer and a justification for the answer to the statistical question are presented.
- A justification includes mention of statistics concepts learned thus far. For example, “The variables are...”; **OR** it includes acknowledgment of variability. For example, “There are between ____ and ____.”

2-point response:

- The poster identifies possible factors that might be associated with depression in teens.
- A graphical representation that shows an association was created.
- An answer and a justification for the answer to the statistical question are presented.

1-point response:

- The poster identifies possible factors that might be associated with depression in teens.
- An answer to the statistical question is presented but a justification is missing.
- A histogram, bar chart, scatterplot, or other graphical representation was correctly created.

0-point response:

- The poster does not identify possible factors that might be associated with depression in teens.
- A histogram, a dot plot or other graphical representation was incorrectly created **OR** is missing.
- No answer **AND/OR** no justification for the answer to the statistical question is presented.

Next Day

Lab 1H: Our time.

Complete Lab 1H prior to the End of Unit Project.

Lab 1H: Our time.

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

We've come a long way

- The labs until now have covered a huge range of topics:
 - We've learned how to make plots for different types of variables.
 - We know how to subset our data to get a more refined view of our data.
 - We've covered cleaning data and making two-way frequency tables.
- In this lab, we're going to combine all of these ideas and topics together to find out how we spend our time.

First steps first.

- **Export, Upload, Import the data from your class's Time Use campaign.**
 - The data, as-is, is very messy and hard to interpret/analyze.
 - **(1) Fill in the blank with the name of your imported data to format it:**
- ```
timeuse <- timeuse_format(____)
```
- This function formats/cleans the data so that each row represents a typical day for each student in the class.
  - **Hint: Search your History tab for the code to save your formatted timeuse data as an R data file (.Rda).**

### timeuse\_format specifics

- In case you're wondering, the `timeuse_format` function:
  - Takes each student's daily data and adds up all of the time spent doing each activity for each day.
  - The time spent on each activity for each day is then averaged together to create a *typical day* in the life of each student.

### Exploring your data

- Start by getting familiar with your `timeuse` data:
  - **(2) How many observations and variables are there?**
  - **(3) What are the names of the variables?**
  - **(4) Which row represents YOUR typical day?**

### How do we spend our time?

- We would like to investigate the *research question*: "How did our class spend our time?"
  - To do this, we'll perform a statistical investigation.
- **(5) State and answer two statistical investigative questions based on our research question.**
- **(6) State one way in which your personal data is typical and one way that it differs from the rest of the class.**
- **(7) Justify your answers by using appropriate statistical graphics and summary tables.**
  - **(8) If you subset your data, explain why and how it benefited your analysis.**

## **End of Unit 1 Project: Evaluate Claims from the Media**

### **Objective:**

 Students will apply their learning of the first unit in the curriculum by completing an end of unit project.

### **Materials:**

1. *End of Unit 1 Project: Evaluating Claims from the Media* (LMR\_U1\_End\_of\_Unit\_Project)

### **End of Unit 1 Project Evaluating Claims from the Media**

Congratulations! You are on your way to becoming a Data Scientist. You have now learned some basic statistics concepts - along with RStudio skills - to help you analyze and interpret data. It is time to apply what you have learned so far.

You will apply what you have learned by engaging in the following:

1. Use an article from the list provided below, or find an article, report, blog post, etc., in a magazine, newspaper, or other media related to the topic of nutrition or time use that makes a claim. Use an article we have not used in class.
  - a. *Child Nutrition in America Today: A New Look at a Crisis of Our Own Making*  
[https://link.gale.com/apps/doc/A724157086/GPS?u=lnoca\\_brecksv&sid=bookmark-GPS&xid=f72e0a60](https://link.gale.com/apps/doc/A724157086/GPS?u=lnoca_brecksv&sid=bookmark-GPS&xid=f72e0a60)
  - b. *Americans Are Eating More Ultra-Processed Foods*  
<https://www.nyu.edu/about/news-publications/news/2021/october/ultra-processed-foods.html>
  - c. *2022 Food and Health Survey Spotlight: Snacking*  
<https://foodinsight.org/survey-spotlight-snacking/>
  - d. *Screen Time Statistics: Average Screen Time by Country*  
<https://www.comparitech.com/tv-streaming/screen-time-statistics/>
  - e. *Gallup: Teens Spend More Time On Social Media Than On Homework*  
<https://www.forbes.com/sites/bradadgate/2023/10/18/gallup-teens-spend-more-time-on-social-media-than-on-homework/>
2. Analyze the article or report based on the following questions:
  - a. What claim(s) did the article make?
  - b. What statistical questions were they trying to answer?
  - c. Does the article cite data? If so:
    - i. Who was observed and what were the variables observed?
    - ii. Who collected the data?
    - iii. How was the data collected?
    - iv. What are some statistics that the article used to make the claim(s)?
  - d. If there was no data, how did the article justify its claim?
3. Determine whether the class's Food Habits or Time Use campaign data supports, refutes, or is inconclusive of the claim(s) made in the article.
4. Use RStudio to do your analysis using either the Food Habits or Time Use campaign data and create graphics/plots that support your reasoning.
5. Generate other statistical questions that you would like to investigate further after you reach your conclusion.
6. Write a summary of your analysis that is no more than 4 pages long. Include graphics/plots/tables that provide evidence to support your reasoning. Be sure to include everything in items 1-5.
7. Prepare a 2-minute presentation of your report. Make sure you refer to your graphics/plots/tables during your presentation.



# Introduction to Data Science

## Unit 2

## Introduction to Data Science

### Daily Overview: Unit 2

| Theme                                      | Day   | Lessons and Labs                                           | Campaign                                                  | Topics                                                   | Page |
|--------------------------------------------|-------|------------------------------------------------------------|-----------------------------------------------------------|----------------------------------------------------------|------|
| What Is Your True Color?<br>(10 days)      | 1^+   | Lesson 1: What Is Your True Color?                         | Personality Color - data                                  | Subsets, relative frequency                              | 121  |
|                                            | 2     | Lesson 2: What Does Mean Mean?                             | Personality Color                                         | Measures of center – mean                                | 124  |
|                                            | 3     | Lesson 3: Median In the Middle                             | Personality Color                                         | Measures of center – median                              | 128  |
|                                            | 4     | Lesson 4: How Far Is It from Typical?                      | Personality Color                                         | Measures of spread – MAD                                 | 131  |
|                                            | 5     | Lab 2A: All About Distributions                            | Personality Color                                         | Measures of center & spread – mean, median, MAD          | 135  |
|                                            | 6     | Lesson 5: Human Boxplots                                   |                                                           | Boxplots, IQR                                            | 137  |
|                                            | 7     | Lesson 6: Face Off                                         |                                                           | Comparing distributions                                  | 140  |
|                                            | 8     | Lesson 7: Plot Match                                       |                                                           | Comparing distributions                                  | 143  |
|                                            | 9     | Lab 2B: Oh, the Summaries...                               | Personality Color                                         | Boxplots, IQR, numerical summaries, custom functions     | 145  |
|                                            | 10    | Practicum: The Summaries                                   | Food Habits or Time Use                                   | Data cycle, comparing distributions                      | 148  |
| How Likely Is It?<br>(7 days)              | 11    | Lesson 8: How Likely Is It?                                |                                                           | Probability, simulations                                 | 151  |
|                                            | 12    | Lesson 9: Dice Detective                                   |                                                           | Simulations to detect unfairness                         | 154  |
|                                            | 13    | Lesson 10: Marbles, Marbles                                |                                                           | Probability, with replacement                            | 158  |
|                                            | 14    | Lab 2C: Which Song Plays Next?                             |                                                           | Probability of simple events, do loops, set.seed()       | 160  |
|                                            | 15    | Lesson 11: This AND/OR That                                |                                                           | Compound probabilities                                   | 163  |
|                                            | 16    | Lab 2D: Queue It Up!                                       |                                                           | Probability with & without replacement, sample()         | 167  |
|                                            | 17    | Practicum: Win, Win, Win                                   |                                                           | Probability estimation through repeated simulations      | 170  |
| Are You Stressing or Chilling?<br>(8 Days) | 18^   | Lesson 12: Don't Take My Stress Away                       | Stress/Chill – data                                       | Introduction to campaign                                 | 173  |
|                                            | 19    | Lesson 13: The Horror Movie Shuffle                        | Stress/Chill – data                                       | Chance differences – categorical variables               | 177  |
|                                            | 20    | Lab 2E: The Horror Movie Shuffle                           | Stress/Chill – data                                       | Inference for categorical variable, do loops, shuffle()  | 181  |
|                                            | 21    | Lesson 14: The Titanic Shuffle                             | Stress/Chill – data                                       | Chance differences – numerical variables                 | 184  |
|                                            | 22    | Lab 2F: The Titanic Shuffle                                | Stress/Chill – data                                       | Inference for numerical variable, do loops, shuffle()    | 188  |
|                                            | 23+   | Lesson 15: Tangible Data Merging                           | Stress/Chill – data                                       | Merging datasets                                         | 190  |
|                                            | 24    | Lab 2G: Getting It Together                                | Stress/Chill & Personality Color                          | Merging datasets, stacking vs. joining                   | 192  |
|                                            | 25    | Practicum: What Stresses Us?                               | Stress/Chill & Personality Color                          | Analyzing merged data                                    | 194  |
| What's Normal?<br>(5 Days)                 | 26    | Lesson 16: What Is Normal?                                 |                                                           | Introduction to normal curve                             | 196  |
|                                            | 27    | Lesson 17: Normal Measure of Spread                        |                                                           | Measures of spread – SD                                  | 199  |
|                                            | 28    | Lesson 18: What's Your Z-Score?                            |                                                           | z-scores, shuffling                                      | 202  |
|                                            | 29    | Lab 2H: Eyeballing Normal                                  |                                                           | Normal curves overlaid on distributions & simulated data | 206  |
|                                            | 30    | Lab 2I: R's Normal Distribution Alphabet                   |                                                           | Normal probability, rnorm(), pnorm(), quantiles, qnorm() | 208  |
| End of Unit Project<br>(5 Days)            | 31-35 | End of Unit 2 Project: Comparing Groups Using Our Own Data | Stress/Chill, Personality Color, Food Habits, or Time Use | Synthesis of above                                       | 210  |

^=Data collection window begins.

+=Data collection window ends.

## IDS Unit 2: Essential Concepts

### Lesson 1: What Is Your True Color?

Students will understand that the 'typical' value is a value that can represent the entire group, even though we know that not all members of the group share the same value.

### Lesson 2: What Does Mean Mean?

The center of a distribution is the 'typical' value. One way of measuring the center is with the mean, which finds the balancing point of the distribution. The mean gives us the typical value but does not tell the whole story. We need a way to measure the variability to understand how observations might differ from the typical value.

### Lesson 3: Median In the Middle

Another measure of center is the median, which can also be used to represent the typical value of a distribution. The median is preferred for skewed distributions or when there are outliers because it better matches what we think of as 'typical'.

### Lesson 4: How Far Is It from Typical?

Mean Absolute Deviation (MAD) measures the variability in a sample of data – the larger the value, the greater the variability. More precisely, the MAD is the typical distance of observations from the mean. There are other measures of spread as well, notably the standard deviation and the interquartile range (IQR).

### Lesson 5: Human Boxplots

A common statistical question is "How does this group compare to that group?". This is a hard question to answer when the groups have a lot of variability. One approach is to compare the centers, spreads, and shapes of the distributions. Boxplots are a useful way of comparing distributions from different groups when all of the distributions are unimodal (one hump).

### Lesson 6: Face Off

Writing (and saying) precise comparisons between groups in which variability is present based on the (a) center, (b) spread, (c) shape, and (d) unusual outcomes help to make statements in context of the data. Actual comparison statements should use terms such as "less than," "about the same as," etc.

### Lesson 7: Plot Match

Boxplots are an alternative visualization of histograms or dotplots. They capture most, but not all, of the features we can see in a dotplot or histogram.

### Lesson 8: How Likely Is It?

Probability is an area that we humans have poor intuition about. Probability measures a long-run proportion: 50% chance means the event happens 50% of the time *if you repeated it forever*. When we don't repeat forever, we see variability.

### Lesson 9: Dice Detective

In the short-term, actual outcomes of chance experiments vary from what is 'ideal'. An ideal die has equally likely outcomes. But that does not mean we will see exactly the same number of one-dots, two-dots, etc.

### Lesson 10: Marbles, Marbles...

There are two ways of sampling data that model real-life sampling situations: with and without replacement. Larger samples tend to be closer to the "true" probability.

### **Lesson 11: This AND/OR That**

What does “A or B” mean versus “A and B” mean? These are compound events and two-way tables can be used to calculate probabilities for them.

### **Lesson 12: Don’t Take My Stress Away!**

Generating statistical questions is the first step in a Participatory Sensing campaign. Research and observations help create applicable campaign questions.

### **Lesson 13: The Horror Movie Shuffle**

We can “shuffle” data based on categorical variables. The statistic we use is the difference in proportions. The distribution we form by shuffling represents what happens if chance were the only factor at play. If the actual observed difference in proportions is near the center of this shuffling distribution, then we would conclude that chance is a good explanation for the difference. But if it is extreme (in the tails or off the charts), then we should conclude that chance is NOT to blame. Sometimes, the apparent difference between groups is caused by chance.

### **Lesson 14: The Titanic Shuffle**

We can also “shuffle” data based on numerical variables. The statistic we use is the difference in means. The distribution we form by this form of shuffling still represents what happens if chance were the only factor at play. When differences are small, we suspect that they might be due to chance. When differences are big, we suspect they might be ‘real’.

### **Lesson 15: Tangible Data Merging**

We can enhance the context of a statistical problem by merging related datasets together. To merge data, each dataset must have a “unique identifier” that tells us how to match up the lines of the data.

### **Lesson 16: What Is Normal?**

The Normal curve, also called the Gaussian distribution and the “bell curve”, is a model that describes many real-life distributions and is usually called the Normal Model.

### **Lesson 17: A Normal Measure of Spread**

The standard deviation is another measure of spread. This is commonly used by statisticians because of its role in common models and distributions, such as the Normal Model.

### **Lesson 18: Shuffling with Normal**

Z-scores allow us a way to measure how extreme a value is, regardless of the units of measurement. Typically, z-scores will range between -3 and +3, and so values that are at or more extreme than -3 or +3 standard deviations are considered extremely rare.

# What is Your True Color?

Instructional Days: 10

## Enduring Understandings

Statistics enable us to make sense of large amounts of data. Numerical summaries capture important elements of a distribution. Measures of center, also known as measures of central tendency, show the tendency of quantitative data to gather around a central value. Measures of spread, also known as measures of variability, show how much the quantitative data is spread out. Measurements of the propensity for the data to cluster on a central location and the range of variability within the data can provide insightful indicators about the data.

## Engagement

Students will complete the *True Colors Personality Test* to discover the qualities and characteristics of their personality styles. Students will use the results from the personality color test to learn about subsetting data and finding measures of center and spread. The data from their personality test will be collected in a survey using the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org/>

## Learning Objectives

### Statistical/Mathematical:

A-SSE A: Interpret the structure of expressions.

S-ID 2: Use statistics appropriate to the shape of the data distribution to compare center (median, mean) and spread (interquartile range, standard deviation) of two or more different data sets.

S-ID 3: Interpret differences in shape, center, and spread in the context of the data sets, accounting for possible effects of extreme data points (outliers).

S-IC 6: Evaluate reports based on data.

### Focus Standards for Mathematical Practice for All of Unit 2:

SMP-4: Model with mathematics.

SMP-5: Use appropriate tools strategically.

### Data Science:

Understand the information that numerical summaries provide about the data. Understand that a boxplot is a graphical representation of a numerical summary.

### Applied Computational Thinking Using RStudio:

- Calculate numerical summaries (mean, median, Mean of Absolute Deviations (MAD), and Interquartile Range (IQR)).
- Create graphical representations to compare two or more datasets, including boxplots.

### Real-World Connections:

We must be able to synthesize vast amounts of data into coherent, comprehensible measures. Today's media is continuously publishing articles that include statistical references. Critical consumerism requires that we understand the information provided in summaries of data.

## Language Objectives

1. Students will consider the appropriateness of using the mean or median as the measure of center for different types of distributions.
2. Students will determine and explain an appropriate measure of spread.
3. Students will engage in an "Active Debate" activity to compare shape, center, and spread of distributions.
4. Students will engage in partner and whole group discussions and presentations to express their understanding of statistical investigations.

## Data File or Data Collection Method

### Data Files:

1. Students' *Personality Color* survey: data(colors)

### Data Collection Method:

**True Colors Personality Test:** Students will complete the *Personality Color* survey that will collect their data about their personality styles.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 1: What is Your True Color?

### Objective:

Students will collect data that might tell them about their personality type and will understand how to subset their data.

### Materials:

1. True Colors Personality Test (LMR\_U2\_L1)
2. Posted signs for each Personality Color: Blue, Gold, Green, and Orange  
**Advance preparation required:** Posters need to be labeled and posted (see step 5 in lesson)
3. Poster paper
4. Markers
5. Data collection devices

### Vocabulary:

subsets

**Essential Concepts:** Students will understand that the 'typical' value is a value that can represent the entire group, even though we know that not all members of the group share the same value.

### Lesson:

1. Ask students to consider the following questions (they do not need to record any responses):
  - a. How well do you know yourself?
  - b. How well do you know your classmates?
2. There are things students know and don't know about themselves. The *True Colors Personality Test* (LMR\_U2\_L1) claims to identify personality types (later, students can gather more evidence to test these claims). Students will use these data to explore fundamental statistical concepts.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Discovering Our Personality through TRUE COLORS**  
(Adapted from Head Start of Greater Dallas -- <http://thead.org>)

**Instructions:** Compare all 4 boxes in each row. Do NOT analyze each word; just get a general sense of each box. **Score each of the 4 boxes in each row from most to least as it describes you:**  
4 = most, 3 = a lot, 2 = somewhat, 1 = least.

|       | A                                                                                  | B                                                                                           | C                                                                            | D                                                                                    |
|-------|------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|
| Row 1 | Active<br>Variety<br>Sports<br>Opportunities<br>Spontaneous<br>Flexible            | Organized<br>Planned<br>Neat<br>Parental<br>Traditional<br>Responsible                      | Warm<br>Helpful<br>Friends<br>Authentic<br>Harmonious<br>Compassionate       | Learning<br>Science<br>Quiet<br>Versatile<br>Inventive<br>Competent                  |
|       | Score <input type="text"/>                                                         | Score <input type="text"/>                                                                  | Score <input type="text"/>                                                   | Score <input type="text"/>                                                           |
| Row 2 | Curious<br>Ideas<br>Questions<br>Conceptual<br>Knowledge<br>Problem Solver         | Caring<br>People Oriented<br>Feelings<br>Unique<br>Empathetic<br>Communicative              | Orderly<br>On-time<br>Honest<br>Stable<br>Sensible<br>Dependable             | Action<br>Challenges<br>Competitive<br>Impetuous<br>Impactful                        |
|       | Score <input type="text"/>                                                         | Score <input type="text"/>                                                                  | Score <input type="text"/>                                                   | Score <input type="text"/>                                                           |
| Row 3 | Helpful<br>Trustworthy<br>Dependable<br>Loyal<br>Conservative<br>Organized         | Kind<br>Understanding<br>Giving<br>Devoted<br>Warm<br>Poetic                                | Playful<br>Quick<br>Adventurous<br>Confrontive<br>Open Minded<br>Independent | Independent<br>Exploring<br>Competent<br>Theoretical<br>Why Questions<br>Ingenuous   |
|       | Score <input type="text"/>                                                         | Score <input type="text"/>                                                                  | Score <input type="text"/>                                                   | Score <input type="text"/>                                                           |
| Row 4 | Follow<br>Rules<br>Useful<br>Save Money<br>Concerned<br>Procedural<br>Cooperative  | Active<br>Free<br>Winning<br>Daring<br>Impulsive<br>Risk Taker                              | Sharing<br>Getting Along<br>Feelings<br>Tender<br>Inspirational<br>Dramatic  | Thinking<br>Solving Problems<br>Perfectionistic<br>Determined<br>Complex<br>Composed |
|       | Score <input type="text"/>                                                         | Score <input type="text"/>                                                                  | Score <input type="text"/>                                                   | Score <input type="text"/>                                                           |
| Row 5 | Puzzles<br>Seeking Info<br>Making Sense<br>Philosophical<br>Principled<br>Rational | Social Causes<br>Easy Going<br>Happy Endings<br>Approachable<br>Affectionate<br>Sympathetic | Exciting<br>Lively<br>Hands On<br>Courageous<br>Skillful<br>On Stage         | Pride<br>Tradition<br>Do Things Right<br>Orderly<br>Conventional<br>Careful          |
|       | Score <input type="text"/>                                                         | Score <input type="text"/>                                                                  | Score <input type="text"/>                                                   | Score <input type="text"/>                                                           |

|                                            |                                          |                                          |                                           |
|--------------------------------------------|------------------------------------------|------------------------------------------|-------------------------------------------|
| <b>Total Orange Score</b><br>A, F, R, N, S | <b>Total Gold Score</b><br>B, G, I, M, T | <b>Total Blue Score</b><br>C, E, O, D, R | <b>Total Green Score</b><br>Q, V, L, P, U |
| <input type="text"/>                       | <input type="text"/>                     | <input type="text"/>                     | <input type="text"/>                      |

If any of the scores in the colored boxes are less than 5 or greater than 20, you have made an error. Please go back and read the instructions.

LMR\_U2\_L1

- Distribute the first 2 pages of the *True Colors Personality Test* (LMR\_U2\_L1). DO NOT include the final page, which contains the descriptions of each color. Instruct students on how to complete the test and allow time for them to complete it (see page 2 in handout).

**Note:** When adding scores for each color at the bottom of the test, make sure that students have **NOT** added straight down each column.

- Students should have a score for each of the 4 colors. Ask students to record each color and its respective score in their IDS journal. Inform them that the color with the *highest* score describes their personality. We can refer to this as their predominant color. They should record their predominant personality color in their DS journal as well. Tell students that you will show them what each color means at the end of the lesson.
- Post a sign for each personality color on different walls of the classroom. For example, Blue on the north wall, Gold on the east wall, etc.



- Ask students to gather by the wall corresponding to their predominant personality color. The students should record answers to the following questions in their DS journals.
  - How many students are in your color group?
  - How many students are in each of the other color groups?
  - What is the predominant personality color in your class?



- Then ask students to determine some common characteristics of the people in their group. Questions to help steer the discussions are included below. Each team should come up with a consensus to describe their team and will share their descriptions with the whole class. The goal is to get the students to think about “What is typical?” within their groups.
  - What are your likes and dislikes?
  - What things do you have in common?
  - What are your favorite activities?
  - What’s your favorite color?
  - Do you prefer mornings or nights?
  - What’s your favorite type of music?
- As the groups are presenting, record some dominant characteristics on the board for each color. The students will be able to compare their perceived traits with the actual descriptions from the activity at the end of the class.
- Next ask students in each color group to gather into two **subsets**: introvert and extrovert. Inform them that subsetting is another way to organize collected data. Create a two-way frequency table like the one below on the board to record the results.

|                         |           | Color  |      |      |       |       |
|-------------------------|-----------|--------|------|------|-------|-------|
| Introvert/<br>Extrovert |           | Orange | Gold | Blue | Green | Total |
|                         | Introvert |        |      |      |       |       |
|                         | Extrovert |        |      |      |       |       |
|                         | Total     |        |      |      |       |       |

- Distribute poster paper and markers to each team.
- Inform the students that they will be creating visuals for this data by comparing subsets. The Orange and Gold groups should create visualizations that subset the color variable by introverts and extroverts, and the Blue and Green groups should create visualizations that subset introverts and extroverts by color.
- If students are confused or stuck, have them recall the topic of two-way tables and relative frequencies from Unit 1 (Lessons 16 & 17). The Orange and Gold groups will be looking at the columns and comparing introverts/extroverts, while the Blue and Green groups will be looking at the rows and comparing colors.

13. Once all groups have completed their visuals, the Orange and Gold teams should choose one of their 2 posters to display to the class. The Blue and Green groups should do the same and select one of their visuals.



14. Display both visuals on the board and discuss their similarities and differences. Ask students to analyze and interpret the visualizations by discussing the following questions for each of the visualizations:
- What type of plot is this and how many variables are present? *Answers will vary by class.*
  - What information about this subset can I gather from this visualization? *Answers will vary by class.*
  - What do I see the most/least of? *Answers will vary by class.*
  - What is the typical personality color for this subset? Or, what is the typical group (introverts/extroverts) for this subset? *Answers will vary by class.*
15. Ask students to summarize their impressions of the class's personality color data by writing this summary in their DS journals.
16. Distribute the description of each personality color to students (page 3 of LMR\_U2\_L1). Remind them that the highest score is considered their predominant color and the second highest score is considered their secondary color. If there is a tie for their predominant or secondary colors, ask students to choose the color that describes their personality better.
17. Compare the given descriptions on the handout to the characteristics listed on the board for each group during step 7 above. Do the descriptions match what the students originally thought? How accurate are the descriptions? If time allows, ask a couple of students to share their comparisons.
18. Students will now record their data by completing the *Personality Color* campaign on the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org>.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework

If not finished in class, students should complete the *Personality Color* survey either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org/>

## Lesson 2: What Does Mean Mean?

### Objective:

Students will learn that values that gather around the center of a distribution show the typical value. This value is also referred to as the mean, or average.

### Materials:

1. *Pennies on a Ruler* handout (LMR\_U2\_L2\_A)  
**Digital Option:** IDS Balancing Point app  
*Balancing Point* handout (LMR\_U2\_L2\_B)
2. Markers (1 for each table)
3. Rulers (1 for each table)
4. Pennies (6 for each table group)
5. Tape
6. Exported, printed, and reproduced class's *Personality Color* survey data  
**Advance preparation required:** The teacher must share students' data on the IDS Home page (<https://portal.thinkdataed.org>) before it can be exported and printed. Students will keep for use in subsequent lessons (see step 9 of lesson)
7. *Mr. Jones Mile Run Times* handout (LMR\_U2\_L2\_C)

### Vocabulary:

measures of central tendency (or center), typical, measures of variability (or spread), mean, average, balancing point

**Essential Concepts:** The center of a distribution is the 'typical' value. One way of measuring the center is with the mean, which finds the balancing point of a distribution. The mean gives us the typical value but does not tell the whole story. We need a way to measure the variability to understand how observations might differ from the typical value.

### Lesson:



1. In student pairs, ask students to discuss what they think the following terms mean:
  - a. Measures of central tendency. *Answer: A value that shows the tendency of quantitative data to gather around a central, or typical, value. Also known as measures of center. Students will learn about two such measures: the mean and the median.*
  - b. Measures of variability. *Answer: Values that show how much the quantitative data varies. Also known as measures of spread.*  
**Note:** Measures of variability are not taught during this lesson but will be addressed as part of Lesson 4.
2. Ask a pair to share what they think these two terms mean. Pairs who are listening must decide whether they agree or disagree with the pair that shared. Lead a discussion based on their statements of agreement or disagreement.
3. Communicate to the class that they will be learning more about these measures and what they tell us about data as we progress through this unit.
4. By a show of hands, ask students how many are familiar with finding the **mean**, or **average**.
5. Select a student to share his/her process for finding the mean. *Possible answer: Add up all of the numbers. Then divide by how many numbers there are.*
6. Another way to find the mean is to find the **balancing point** of a distribution. They will learn about the balancing point via the activity in steps 7 and 8 below.
7. Distribute the *Pennies on a Ruler* handout (LMR\_U2\_L2\_A) along with a marker, ruler, tape, and 6 pennies to each table group. If you prefer to not print the document, you can project it on the board instead.

**Note:** Alternatively, you can substitute the digital version of this activity using the *Balancing Point* handout (LMR\_U2\_L2\_B).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Pennies on a Ruler

The **balancing point** of a dataset is the point on a number line where the data distribution is balanced. Use the instructions below to find the **balancing point** of the following set of numbers: 2, 3, 6, 8, 9, 11.

Instructions:

- Estimate the balancing point:
  - Tape a marker securely to your desk.



- Model the dataset by centering pennies on the 2-inch, 3-inch, 6-inch, 8-inch, 9-inch, and 11-inch marks on a 12-inch ruler.



- Carefully place the ruler on top of the marker. Make sure that the coins do not move from their original positions. If necessary, you can tape the pennies to the ruler. Try to balance the ruler on the marker. To the nearest half inch, at what value on the ruler is the data balanced?



LMR\_U2\_L2\_A

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Balancing Point

The **balancing point** of a dataset is the point on a number line where the data distribution is balanced.

- Use the instructions below to find the **balancing point** of the following set of numbers: 2, 3, 5, 6, 9.

Instructions:

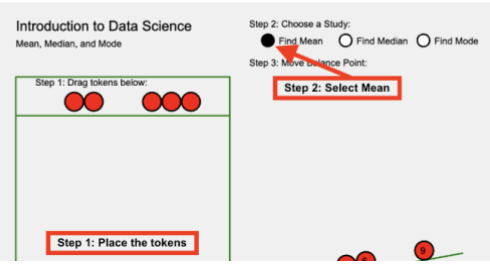
**Step 1:** Drag a token and place it above each of the numbers in the set (2, 3, 5, 6, 9)

**Step 2:** Make sure "Find Mean" is selected.

**Step 3:** Click and drag the yellow triangle (fulcrum) until the green line is balanced (horizontal).

**Step 4:** Click on "Show Calculation" to have the computer calculate the mean.

What do you notice?



LMR\_U2\_L2\_B

- Guide the students through the handout and have them share their findings throughout the activity. Be sure to emphasize the idea that the mean of a distribution can be identified by finding its balancing point.
- Next, distribute the class's *Personality Color* survey data to the students.
- Have student pairs find the variable **Blue** (whether or not that was their predominant color) in the class's printed data.
- As a class, make a dot plot on the board to show the distribution of **Blue** values. Each student should come to the board and draw a dot to indicate where their value is in the distribution. Ask the students:
  - What do you think the typical **Blue** score is? *Answers will vary by class. They should be driven to an answer in the center of the distribution.*
  - Are the data roughly symmetric? Where is the balancing point of this distribution? *Answers will vary by class. Once a value is chosen, indicate the location on the dot plots.*
- As a class, compute the mean **Blue** score for the entire class on the board and compare this value to the class's prediction of the balancing point. Students may not remember exactly how to compute the mean, so you can remind them of the general algorithm or refer them back to their responses from step 5 above.

13. Show the students the formula for calculating the mean:

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$$

14. Now that they have calculated the mean for the **Blue** score, ask them to identify each symbol in the formula with a step in their algorithm for finding the mean, and discuss the meaning of the symbols in the formula as a class. *Answer:  $x_i$  represents each individual data point and  $n$  represents the total number of observations.*
15. Indicate the location of the calculated mean on the dot plots by drawing a vertical line at the value on the x-axis. Ask student pairs to engage in a conversation about how close the mean value is to their predicted balancing point and why their prediction was made that way. Select a pair to share their discussion with the whole class.
16. Using the *Personality Color* survey data from step 9 above, ask student pairs to compute the mean score for each of the other three personality colors.
17. Inform the students that, during the next lesson, they will learn about another method that can be used for measuring the center of a distribution.
18. Now, you can inform the class about an even easier method of calculating the mean – using RStudio! Explain that the command RStudio uses to calculate the mean incorporates the algorithm of summing up all the data and dividing by the total number of observations. Students will be able to use this command for quick calculations now.
- Note:** If you have already “Exported, Uploaded, Imported” the class’s *Personality Color* campaign data, you can simply use the exact command below to calculate the mean **Blue** score:
- ```
> mean(~blue, data = colors)
```
- In general, the function can be denoted as follows:
- ```
> mean(~variable, data = datafile)
```
- So, for our specific example, **blue** is the **variable** we want to find the mean value of, and **colors** is the **datafile**.
19. Have the students *Think-Pair-Share* to discuss how the mean value of a group of data could be used to easily describe complicated things. For example, instead of giving someone the entire class’s **Blue** scores, we could just tell him/her the mean score and he/she would have a general idea about the class.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students should complete the *Mr. Jones Mile Rule Times* handout (LMR\_U2\_L2\_C) for homework. They can practice finding the mean of distributions by determining a balancing point for the data. Answers to the handout are below.

**Note:** The mean values in part (3) do NOT need to be exact.

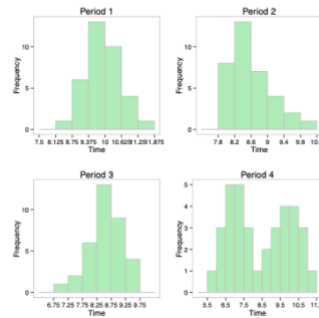
Name: \_\_\_\_\_

Date: \_\_\_\_\_

Mr. Jones  
Mile Run Times

## Background:

Mr. Jones is a physical education teacher at a high school. Every year, his students must run a distance of one mile in a specified amount of time. At the beginning of the year, he gives the students a practice one-mile run and records their times. He hopes to compare these run times to others throughout the year to see if there is any improvement in the students' running paces.



LMR\_U2\_L2\_C

- What kind of plots did Mr. Jones create for his classes? *Answer: Histograms.*
- Where does each distribution balance? Find and label the balancing point of each distribution.  
*Answer: The balancing point for all of these distributions is at the mean.*
- Based on the balancing points you found, what would you say the mean mile run time is for each class?
  - Period 1: 9.91
  - Period 2: 8.48
  - Period 3: 8.45
  - Period 4: 8.17

### Lesson 3: Median in the Middle

#### Objective:

Students will learn that the median is another way to measure the center, or typical-ness, of a distribution, and will understand how medians compare and contrast with the mean.

#### Materials:

1. Sticky notes (one per student)  
**Advance preparation required:** Have one sticky note for each student (see step 6 in lesson)
2. Poster paper
3. Graphics from *Medians – Dotplots or Histograms?* (LMR\_U2\_L3\_A)
4. *Where is the Middle?* handout (LMR\_U2\_L3\_B)
5. Exported, printed, and reproduced class's *Personality Color* survey data

#### Vocabulary:

median

**Essential Concepts:** Another measure of center is the median, which can also be used to represent the typical value of a distribution. The median is preferred for skewed distributions or when there are outliers because it better matches what we think of as 'typical'.

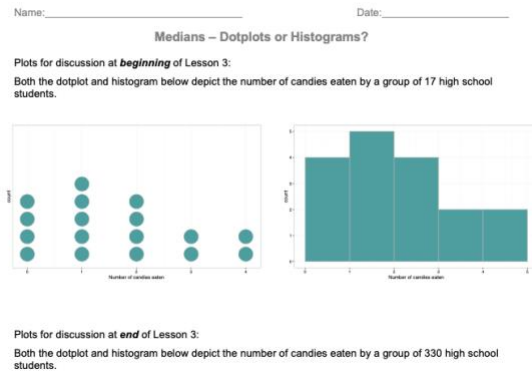
#### Lesson:

1. Remind students that, during the previous lesson, they learned about the mean as the balancing point of a distribution and as a measure of center. In statistics, there are a few values that can be considered as measures of center – the mean is one, and another is the **median**. The median is the middle value in a group of ordered observations.
2. As a simple example, write or display the following group of numbers on the board:  
$$8, 2, 6, 3, 7, 4, 9, 5, 5$$
3. Since there are 9 numbers in the list above, we should use the 5<sup>th</sup> number as the median because it is directly in the middle and there are 4 numbers above it, and 4 numbers below it.
4. However, students should realize that they cannot simply pick the middle number of the list as it is currently written (this would give a median value of 7). Instead, they must first arrange the numbers in numerical order (from lowest to highest).  
$$2, 3, 4, 5, 5, 6, 7, 8, 9$$
5. Now they can identify that the true median value of this list of numbers is 5.
6. Next, randomly distribute one sticky note to each student.

**Advance preparation required:** There should be one card for every student in the class. All of the cards, except one, need to have the value 0 written on them. One card should have the value 1,000,000 written on it.

- 
7. Place poster paper on the board and have the students create a dotplot by placing their sticky notes at the corresponding values on the axis. Then, ask and record answers to the following questions:
    - a. What is the typical value of these data? *Answer: 0 – all sticky notes but one have a value of 0.*
    - b. Using the formula we learned in class, calculate the mean, or average, value of this distribution. *Answers will vary by class/class size. Example: for a class with 28 students enrolled, there would be 27 values of 0 and 1 value of 1,000,000. Therefore, the mean value would be  $(0 \cdot 27 + 1,000,000)/28 \approx 35,714.3$ .*
    - c. Does the mean you calculated match your understanding of “typical”? Why is the mean not capturing our notion of “typical”? *Answer: The 1,000,000 value is heavily skewing the calculation of the mean. It is pulling the mean to a higher value than what we consider to be typical for these data.*

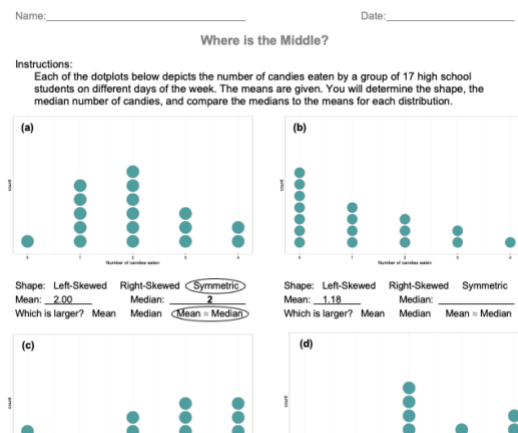
8. Since we introduced the idea of the median as a measure of center at the beginning of class, have the students find the median value of the data on their sticky notes. If time permits, have them place the sticky notes in a line across the board in order (from least to greatest) and have them find the middle number. The median value will be 0.
9. Ask students why there is such a large difference between the mean and median values even though they are both measures of center? Is there a specific reason why the mean is larger than the median for this particular set of data? *Answer: In this case, there was an outlier value that skewed the distribution and forced the balancing point to move to the right.*
10. Display the first 2 plots in the *Medians – Dotplots or Histograms?* file (LMR\_U2\_L3\_A). They are labeled as plots for discussion for the *beginning* of class. Both the dotplot and histogram depict the number of candies eaten by a group of 17 high school students.



LMR\_U2\_L3\_A



11. For the first 2 plots, ask students:
  - a. Which plot makes it easier to find the median number of candies eaten – the dotplot or the histogram? Why? *Answer: The dotplot is easier because we can simply find the middle dot and record the value. It is harder on the histogram, because we would have to add up the amount in each bar to find the middle person.*
  - b. What is the median value? *Answer: The median number of candies eaten is 1 candy.*
12. Inform the students that they will practice finding medians of distributions using the *Where is the Middle?* handout (LMR\_U2\_L3\_B). They will be determining medians when distributions have different shapes (e.g., symmetric, left-skewed, right-skewed).
13. Distribute the *Where is the Middle?* handout (LMR\_U2\_L3\_B). Students should complete the handout individually first, then compare answers with their team members. Once each team has agreed upon their answers, discuss the handout as a class.



LMR\_U2\_L3\_B



14. Ask the following questions to elicit a team discussion about the relationship between means and medians:

- What did you notice about the relationship between the mean and median values for the symmetric distributions? *Answer: The mean and median values in the symmetric distributions – plots (a) and (d) – are fairly similar. For plot (a), the mean and median are exactly equal. For plot (d), the mean is actually larger than the median, but not by much ( $2.29 > 2$ ).*
- What did you notice about the relationship between the mean and median values for the left-skewed distributions? *Answer: The mean value was smaller than the median value in both of the left-skewed distributions – plots (c) and (f). Both plots had the same values for the mean (2.53) and the median (3.00) – clearly, the mean is much smaller than the median ( $2.53 < 3$ ).*
- What did you notice about the relationship between the mean and median values for the right-skewed distributions? *Answer: The mean value was larger than the median value in both of the right-skewed distributions – plots (b) and (e). For plot (b), the mean was only slightly higher than the median ( $1.18 > 1$ ). For plot (e), the mean was a decent amount higher than the median ( $0.47 > 0$ ).*



15. Steer the discussion towards the relationship between the shape of a distribution and its corresponding mean and median values.

- Is there a pattern that emerges between the mean and median values for differently shaped distributions? *Answer: Yes! It seems that symmetric distributions will produce similar mean and median values, left-skewed distributions will produce smaller means and higher medians, and right-skewed distributions will produce higher means and smaller medians.*
- For each of the plots in the *Where is the Middle?* handout (LMR\_U2\_L3\_B), which value better matches your idea of “typical” for that specific distribution? *Answer: For plot (a), both the mean and median agree and appear to be the balancing point of the distribution – both match what we think is typical. For plot (b), the median seems to be more typical, but the values are very close. For plot (c), the median appears to be a more typical value. For plot (d), both the mean and median appear to be capturing our idea of typical. For plot (e), the median is a better match to typical. For plot (f), the median is also a better match.*

16. Steer the discussion so that students recognize that the better measures of center for skewed distributions are typically medians, and the better measures for center for symmetric distributions are typically means.

17. Display the last 2 plots in the *Medians – Dotplots or Histograms* file (LMR\_U2\_L3\_A). They are labeled as plots for discussion for the *end* of class. Both the dotplot and histogram depict the number of candies eaten by a group of 330 high school students.



18. For the last 2 plots, ask students:

- Which plot makes it easier to find the median number of candies eaten – the dotplot or the histogram? Why? *Answer: The histogram is easier because we can estimate based on the distribution’s shape. There are too many dots in the dotplot to find the exact middle person.*
- What is the median value? *Answer: The median number of candies eaten is 7 candies.*

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students should calculate the median values for each of their personality color scores. They should compare the median values to the mean values (calculated in Lesson 2) and make a decision about the possible shape of the distribution if we were to create a dotplot of the scores.

## Lesson 4: How Far is it from Typical?

### Objective:

Students will understand that the mean of the absolute deviations (MAD) is a way to assess the degree of variation in the data from the mean and adjusts for differences in the number of points in the dataset ( $n$ ). The MAD measures the total distance between all the data values from the mean and divides it by the number of observations in the dataset.

### Materials:

1. Masking tape (or painter's tape) – approximately 4-5 feet long – one for each student team
2. *How Far Apart?* handout (LMR\_U2\_L4) – will be used again in Lesson 17
3. Exported, printed, and reproduced class's *Personality Color* survey data

### Vocabulary:

measures of variability (spread), deviation, mean of absolute deviations (MAD)

**Essential Concepts:** Mean Absolute Deviation (MAD) measures the variability in a sample of data – the larger the value, the greater the variability. More precisely, the MAD is the typical distance of observations from the mean. There are other measures of spread as well, notably the standard deviation and the interquartile range (IQR).

### Lesson:

1. Remind students that they learned about 2 different measures of center during the previous 2 lessons: the mean and the median. Have the students recall when it is appropriate to use each value based on the shape of the distribution.
  - a. Mean – use with symmetric distributions.
  - b. Median – use with skewed distributions or when there are outliers.
2. Inform the students that, during today's lesson, they will learn about **measures of variability** – also known as measures of **spread**. These values show us how much the quantitative data varies from the center of a distribution. Similar to measures of center, we will use two different measures of spread: (1) the mean of absolute deviations (MAD), and (2) the interquartile range (IQR).

**Note:** IQR will be discussed in detail during Lesson 5.



3. Introduce the term **deviation**. Using *Think, Pair, Share*, ask students what they think this word means and how it could relate to variability. *Answer: A deviation is the act of departing from an established course or accepted standard. Common synonyms include departure, detour, difference, digression, divergence, fluctuation, inconsistency, shift, etc.*
4. On the classroom floor next to each student team, place a 4-5-foot-long piece of masking tape (or painter's tape). Then, propose the following scenario:

Your team has been invited to guest star at the circus! You have been asked to perform as part of the tightrope act – a routine that requires tremendous focus and balance to walk across a tightly pulled rope that is suspended high in the air. In order to practice your balancing skills, the circus has provided your team with a line of tape that will represent the tightrope.
5. Have the students consider the piece of tape (aka the rope) to be the “typical” path they must take to finish the circus act. Since they do not want to fall from the suspended tightrope while performing at the actual circus, they will need to practice walking directly on the middle of the line at all times. If they *deviate* from the line, they will no longer be walking the “typical” path and will likely fall.
6. Each team should select one student to be their starting performer.
7. In teams of 4, one student is the performer, two are measuring the distance of the deviation (one on each side of the tape), and one is the recorder.

8. Place a ruler perpendicular to the “rope” and measure the distance, in centimeters, from the path to the center of the back of the heel of the student who walks while attempting to balance across the “rope”.
9. The performer will walk the tightrope by looking straight up to the sky – first they look to place a foot on the line, then walk naturally while looking up to the sky, and repeating one step at a time for 4 steps, measuring after each step. Any time the performer missteps, this is considered a variation from the typical value. *You can have students take turns so everyone gets a chance to balance, walk, and to measure, depending on time in your class.*
10. Now that the students have an idea about what it means to deviate from something they consider typical, they can start looking at distributions to see how data points vary from their typical value.
11. Inform students that they were observing deviations from typical while calculating actual differences between the rope and the performer’s steps. When data are quantified with numbers, we can then calculate how far away each value is from the center.
12. One such calculation that is popular among data scientists is the mean of absolute deviations (MAD). Ask students to consider the components of the MAD in math terms, and brainstorm what the MAD value might represent. **Answer:**  
*mean – an average*  
*absolute – in mathematics, we talk about absolute value, the positive difference between 2 numerical values*  
*deviation – as discussed earlier in the lesson, deviation represents how much things vary*
13. Using the 3 components in step 12 from above, explain that the MAD measures the absolute distance of each data point from the mean, and then finds the average of all those distances.
14. Display the formula for the MAD for the whole class to see.

$$MAD = \frac{\sum_{i=1}^n |x_i - \bar{x}|}{n}$$

15. Discuss what each symbol in the formula means and how we use it to perform the calculation.  
*Answer:  $x_i$  represents each individual data point,  $\bar{x}$  represents the mean value, and  $n$  represents the total number of observations. The  $\Sigma$  symbol represents the summation – this tells us to add up all the absolute distances from each point to the mean.*
16. To practice using this formula with actual data, students will calculate and compare the MAD values for 2 distributions.
17. Distribute the *How Far Apart?* handout (LMR\_U2\_L4), which contains 2 of the dotplots, plots (a) and (c), from the *Where is the Middle?* handout (LMR\_U2\_L3\_B) used in Lesson 3. As before, the dotplots depict the number of candies eaten by a group of 17 high school students on different days of the week. The means are also given.



Name: \_\_\_\_\_ Date: \_\_\_\_\_

How Far Apart?

Instructions:  
Each of the dotplots below depicts the number of candies eaten by a group of 17 high school students on different days of the week. The means are given.  
**Note:** the plots are labeled (a) and (c) to correspond with the plots on the *Where is the Middle?* handout (LMR\_U2\_L2\_B).  
Answer questions (i) – (iii) below.

(a) Mean = 2.00

Shape: Left-Skewed   Right-Skewed   Symmetric

(c) Mean = 2.53

Shape: Left-Skewed   Right-Skewed   Symmetric

i. Determine the shape of each distribution by circling the corresponding option below the dotplot.

ii. Without doing any calculations – just by looking at the distributions – which one do you think will have a larger MAD value? Why?

\_\_\_\_\_

LMR\_U2\_L4

The calculations for each plot are shown below for the teacher's reference.

#### MAD for plot (a)

$$\begin{aligned}
 MAD &= \frac{1|0 - 2| + 5|1 - 2| + 6|2 - 2| + 3|3 - 2| + 2|4 - 2|}{17} \\
 &= \frac{1(2) + 5(1) + 6(0) + 3(1) + 2(2)}{17} \\
 &= \frac{2 + 5 + 0 + 3 + 4}{17} \\
 &= \frac{14}{17} \\
 &\approx 0.8235
 \end{aligned}$$

#### MAD for plot (c)

$$\begin{aligned}
 MAD &= \frac{3|0 - 2.53| + 0|1 - 2.53| + 4|2 - 2.53| + 5|3 - 2.53| + 5|4 - 2.53|}{17} \\
 &= \frac{3(2.53) + 0(1.53) + 4(0.53) + 5(0.47) + 5(1.47)}{17} \\
 &= \frac{7.59 + 0 + 2.12 + 2.35 + 7.35}{17} \\
 &= \frac{19.41}{17} \\
 &\approx 1.1418
 \end{aligned}$$



18. Students may work in pairs to complete the handout. After all student pairs have come to an agreement on their answers, pose the following questions to the class as a whole:
- Which MAD value did you think would be larger based only on the look/shape of the distributions? Why? *Answer: Since plot (c) is skewed to the left, it probably has a larger MAD because more points will be further away from the mean than in plot (a).*
  - Which MAD value was actually larger when you calculated it? *Answer: The MAD value for plot (c) was larger (1.1418 > 0.8253).*
  - Did your prediction match the actual calculated values, or were you surprised by the results? *Answer: Yes. The distribution with the wider spread (more variability) had the larger MAD value.*

19. To continue exploring with the class's Personality Color survey data, student teams should calculate the MAD value for their *Blue* scores. Does the MAD value seem reasonable based on the dotplot they created during Lesson 2? *Answers will vary by class.*

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

**Homework & Next Day**



Students should calculate the MAD values for each of the other 3 personality color scores and compare the values of the 4 color scores.



Declutter your Environment Pane

- Unit 2 utilizes new datasets so it's a good idea to declutter your Environment Pane.
- Demonstrate it on your own Environment or refer students to this video: <https://www.youtube.com/embed/N5KpS0MFk7Y>

**Lab 2A: All About Distributions**

Complete Lab 2A prior to Lesson 5.

## Lab 2A: All About Distributions

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### In the beginning...

- Most of the labs thus far have covered how to visualize, summarize, and manipulate data.
  - We used visualizations to explore how your class spends their time.
  - We also learned how to clean data to prepare it for analyzing.
- Starting with this lab, we'll learn to use **R** to answer statistical investigative questions that can be answered by calculating the mean, median and MAD.

### How to talk about data

- When we make plots of our data, we usually want to know:
  - Where is the *bulk* of the data?
  - Where is the data more *sparse*, or *thin*?
  - What values are *typical*?
  - How much does the data *vary*?
- To answer these questions, we want to look at the *distribution* of our data.
  - We describe *distributions* by talking about where the *center* of the data are, how *spread* out the data are, and what sort of *shape* the data has.

### Let's begin!

- *Export, upload and import* your class's *Personality Color* data.
  - Name your data **colors** when you load it.
- Before analyzing a new dataset, it's often helpful to get familiar with it. So:
  - **(1) Write down the names of the 4 variables that contain the point-totals, or scores, for each personality color.**
  - **(2) Write down the names of the variables that tell us an observation's introvert/extrovert designation and whether they are involved in sports.**
  - **(3) How many variables are in the dataset?**
  - **(4) How many observations are in the dataset?**

### Estimating centers

- **(5) Write and run code creating a dotPlot of the scores for your predominant color.**
  - Pro-tip: If the **dotPlot** comes out looking wonky, include the **rint** and **cex** options.
- Based on your **dotPlot**:
  - **(6) Which values came up the most frequently? About how many people in your class had a score similar to yours?**
  - **(7) What, would you say, was a typical score for a person in your class for your predominant color? How does your own score for this color compare?**

### Means and medians

- *Means* and *medians* are usually good ways to describe the *typical* value of our data.

- (8) Fill in the blank to calculate the **mean** value of your predominant color score:  

```
mean(~____, data = colors)
```
- (9) Write and run a similar line of code to calculate the **median** value of your predominant color.
  - (10) Are the **mean** and **median** roughly the same? If not, use the **dotPlot** you made in the last slide to describe why.

## Estimating Spread

- Now that we know how to describe our data's *typical* value we might also like to describe how closely the rest of the data are to this *typical* value.
  - We often refer to this as the **variability** of the data.
  - Variability is seen in a **histogram** or **dotPlot** as the horizontal *spread*.
- (11) Write and run code re-creating a **dotPlot** of the scores for your *predominant color* and then write and run the code below filling in the blank with the name of your predominant color.  

```
add_line(vline = mean(~____, data = colors))
```
- (12) Look at the spread of the scores from the mean score then complete the sentence below:
  - Data points in my plot will usually fall within \_\_\_\_ units of the center.

## Mean Absolute Deviation

- The **mean absolute deviation** finds how far away, on average, the data are from the mean.
  - We often write *mean absolute deviation* as *MAD*.
- (13) Write and run code calculating the **MAD** of your *predominant color* by filling in the blank:  

```
MAD(~____, data = colors)
```
- (14) How close was your estimate of the spread for your predominant color (from the previous slide) to the actual value?

## Comparing introverts/extroverts

- Do introverts and extroverts differ in their typical scores for your predominant color?
  - Answer this investigative question using a **dotPlot** and numerical summaries.
- (15) Write and run code making a **dotPlot** of your *predominant color* again; but this time, facet the plot by the introvert/extrovert variable. Include the **layout** option to stack the plots as well as the **ntint** and **cex** options.
- (16) Describe the shape of the distribution of scores for the extroverts. Do the same for the introverts.
- (17) Using similar syntax to how you facet plots, write and run code *calculating* either the **mean** or **median** to describe the *center* of your predominant color for introverts and extroverts.
- (18) Do introverts and extroverts differ in their typical scores for your predominant color?
- (19) Based on the MAD, which group (introverts or extroverts) has more variability for your predominant color's scores?

## On your own

- Do introverts and extroverts in your class differ in their color scores?
  - Perform an analysis that produces *numerical summaries* and *graphs*.
  - (20) Then, write a few sentences that address this statistical investigative question and considers the *shape*, *center* and *spread* of the distributions of the graphs you create.

## Lesson 5: Human Boxplots

### Objective:

Students will learn how and when to use boxplots to compare groups of data. They will learn how to compute and interpret another measure of spread: the IQR.

### Materials:

1. Poster paper, 3-4 feet long  
**Advance preparation required:** Needs to be labeled and posted in class (see step 9 in lesson)
2. Tape
3. Poster paper
4. Markers
5. *Ages of Oscar Winners* handout (LMR\_U2\_L5)

### Vocabulary:

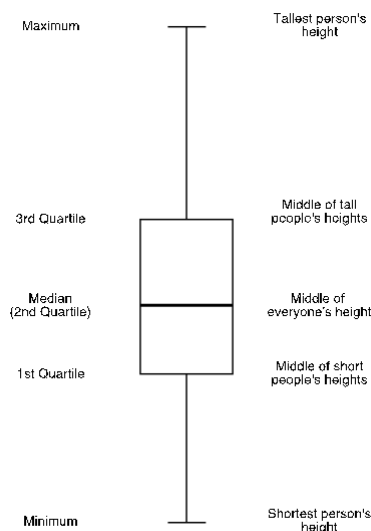
boxplot, quartiles, first quartile ( $Q_1$ ), third quartile ( $Q_3$ ), quantiles, minimum, maximum, five-number summary, range, interquartile range (IQR)

**Essential Concepts:** A common statistical question is “How does this group compare to that group?” This is a hard question to answer when the groups have a lot of variability. One approach is to compare the centers, spreads, and shapes of the distributions. Boxplots are a useful way of comparing distributions from different groups when all the distributions are unimodal (one hump).

### Lesson:

1. Remind students that we have been using the following numerical and graphical summaries to look at data:
  - a. Measures of center – mean, median
  - b. Measures of spread – MAD
  - c. Graphing – dotplots, histograms
2. Explain that all of these tools help us describe data to someone who may not actually be viewing it. Today, we will explore another way to summarize and describe data to others with the use of another type of statistical plot that involves breaking data up into distinct pieces: a **boxplot**.
3. For the next activity, students will need to carry their DS journals and a pen with them.
4. Instruct students to stand up and move their chairs away from the longest wall in the classroom. Ask them to line up against the wall (in no particular order).  
**Note:** If there isn't enough room for everyone to line up together inside the classroom, you may do this activity outside along a building wall.
5. Say, “I want to know which person represents the typical height of students in our class. Can I tell by looking at the line as it currently stands? How would I be able to tell?” Students should discuss with a partner.
6. Ask students to share their discussions. Call on students to contribute to what has been shared if needed. Guide students to see that organizing the data (in other words, themselves) can give you a visual for their heights. Then tell them to line up in height order from shortest to tallest along the wall.
7. Once students are arranged (and this may take a little time – allow students to develop their own algorithm for finding the ordering), ask them how they might be able to describe their distribution of heights. *Possible answers include: mean, median, MAD.*
8. Ask them to split themselves into two groups, one half that is taller and one half that is shorter and have them decide which student represents the class's median height.

9. Have the median student stand next to the wall directly in front of the poster paper.  
**Advance preparation required:** Before class begins, tape a piece of poster paper, approximately 3-4 feet long, vertically to a wall in the classroom. The students will be creating a plot using lines drawn at certain students' heights.
10. Draw a horizontal line on the poster paper to mark the location of the median by having the actual student stand in front of the poster paper so you can mark his/her exact height. Be sure to label this point as the median and include the student's actual height, in inches.
11. Next, ask the two halves to split again, so there are now four groups of students.
12. The breaks between each group are called **quartiles** because they break the data into four groups (*quartile* comes from the Latin word *quartus*, which is also the root of the Spanish word *cuatro*). The lower break represents the **first quartile** (because 25% of the class is shorter than this student's height), and the upper break represents the **third quartile** (because 75% of the class is shorter than this height). Another term that can be used in place of percentiles is **quantiles** because they represent the *quantity* of data that is lower than that value.
13. Using the student who represents the first quartile, draw another horizontal line on the poster paper marking his/her height. The student should stand in the same spot as the student who represented the median so that the line for this student is drawn underneath the median line. Be sure to label this point as the first quartile (or  $Q_1$ ) and include the student's actual height, in inches.
14. Using the student who represents the third quartile, draw another horizontal line on the poster paper marking his/her height. The student should stand in the same spot as the student who represented the median so that the line for this student is drawn above the median line. Be sure to label this point as the third quartile (or  $Q_3$ ) and include the student's actual height, in inches.
15. Finally, ask the tallest and shortest student to stand in front of the poster paper and draw horizontal lines at their heights. The shortest person represents the **minimum** height of the students in the class, and the tallest person represents the **maximum** height. Be sure to label the points as the minimum and maximum, and include each student's actual height, in inches.
16. When you finish, you should have five lines, which represent the **five-number summary**: minimum, first quartile ( $Q_1$ ), median, third quartile ( $Q_3$ ), and maximum. Draw a box using the first and third quartiles as the edges of the box. The median line will be contained within the box. Extend a line from the first quartile down to the minimum and extend a line from the third quartile up to the maximum. Your class's boxplot should look similar to the following:



17. Students should now be facing the newly created boxplot. Allow students time to sketch the boxplot in their DS journals, with the appropriate labels.



18. Ask students:

- What is the difference between the largest and smallest heights? Is there a large difference between the tallest and shortest person? *Answer: Students should calculate maximum - minimum.* Inform students that this difference is known as the **range** of the dataset.
- What is the difference between the quartiles  $Q_1$  and  $Q_3$ ? What percent of our class falls within these two values? *Answer: Students should calculate  $Q_3 - Q_1$ . 50% of the class falls between these two height values.* Inform students that this difference is known as the **interquartile range (or IQR)**.



19. Remind students that they learned about one measure of spread (the MAD) during the previous lesson and tell them that we now have another measure of spread – the IQR. Pose the following questions to the students:

- What does it mean when the IQR is small? *Answer: The middle 50% of heights are close to each other.*
- What does it mean when the IQR is large? *Answer: The middle 50% of heights are more spread out.*

20. Finally, subset the class into introverts and extroverts. Ask each group of students (the introverts and extroverts) to create a boxplot of their group's heights on a piece of poster paper using the techniques they just learned as a class.

21. Ask each group to share their boxplot with the class. Lead a discussion about the similarities and differences between the plots and be sure to include how they compare to the overall combined boxplot of heights they created earlier. In the discussion, have the students calculate the IQR for both plots and make a comparison by asking: What does the IQR tell us about each group?  
*Answers will vary by class.*

### Class Scribes:



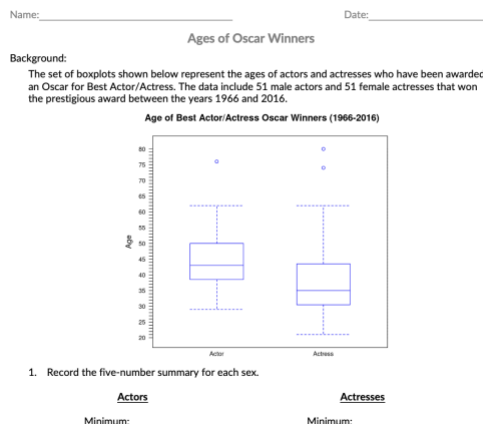
One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Homework



Students should complete the *Ages of Oscar Winners* handout (LMR\_U2\_L5) for homework using their newly acquired knowledge of boxplots.

**Note:** Page 2 of the handout is an answer key for teacher reference only.



LMR\_U2\_L5

## Lesson 6: Face Off

### Objective:

Students will informally compare two or more distributions using their knowledge of shape, center, and spread to answer statistical questions. They will learn how to find the difference between two means and two medians using a histogram or dotplot.

### Materials:

1. *Comparing Commute Times with Dotplots* handout (LMR\_U2\_L6\_A)
2. *Comparing Exam Scores with Histograms* handout (LMR\_U2\_L6\_B)
3. Timer
4. *Comparing Fuel Efficiency with Boxplots* handout (LMR\_U2\_L6\_C)

### Vocabulary:

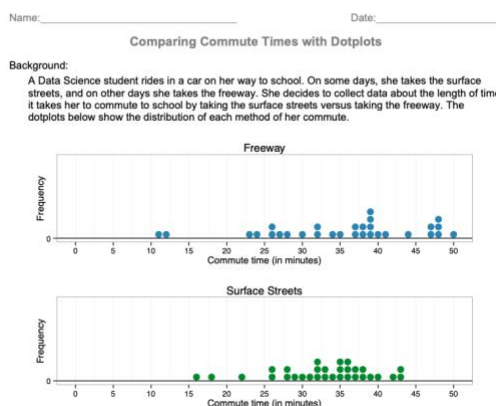
rebuttal

**Essential Concepts:** Writing (and saying) precise comparisons between groups in which variability is present based on the (a) center, (b) spread, (c) shape, and (d) unusual outcomes help to make statements in context of the data. Actual comparison statements should use terms such as “less than,” “about the same as,” etc.

### Lesson:

1. Poll students about the method of transportation they use for their daily school commute. How many of them walk, ride in a car, take the bus, ride a bike, etc.? Record their responses on the board. Ask them to estimate the typical amount of time it takes for them to get to school, in minutes.
2. Inform students that they have learned important features of distributions that will allow them to make decisions when working with data. More specifically, they will be able to use their knowledge of measures of center and measures of spread to compare 2 distributions in order to make a decision.
3. In teams, have students complete the *Comparing Commute Times with Dotplots* handout (LMR\_U2\_L6\_A). Allow students time to read the “Background” portion of the handout and then discuss what statistical question(s) the student in the scenario is trying to answer.

**Note:** Page 2 of the handout is an answer key for teacher reference only.

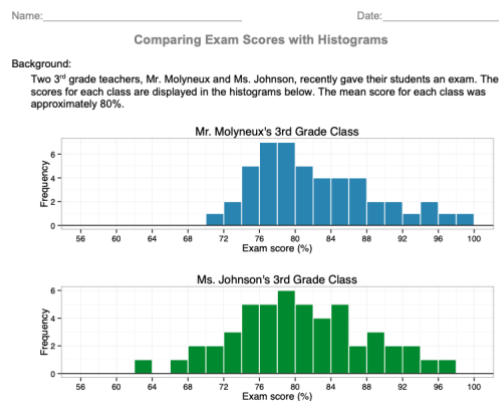


LMR\_U2\_L6\_A

4. Once teams decide on their recommendation, engage half of the class in an *Active Debate*. Half of the students will stand in a debate line and the other half will “fishbowl” the debate. Roles will reverse later in the lesson (see step 12 below).

5. Of those students standing on the debate line, half will argue the reasons why they recommend street travel and the other half will argue the reasons why they recommend freeway travel.
6. On the debate line, each student will stand face to face with a student who has the opposite recommendation. In other words, a student who recommends street travel will stand facing a student who recommends freeway travel.
7. Using a timer, allow one minute for students who recommend freeway travel to argue their point to the person they are facing. Then, repeat for students who recommend street travel. Students should not interrupt or respond; they should only listen to the other side.
8. Next, give debaters two minutes to prepare a **rebuttal** of the other person's argument. For example, if one student claimed that freeway travel is better, the other student may ask where the evidence is in the data or show that the data does not support the claim.
9. Allow each debater two minutes to present his/her rebuttal.
10. Finally, ask debaters if any of them changed their recommendations after engaging in the debate.
11. In teams, have students complete the *Comparing Exam Scores with Histograms* handout (LMR\_U2\_L6\_B). Allow students time to read the "Background" portion of the handout and then discuss what statistical question(s) the student in the scenario is trying to answer.

**Note:** Page 2 of the handout is an answer key for teacher reference only.



LMR\_U2\_L6\_B

12. Repeat debate process (steps 4-10 from above) with the other half of the class.
13. Summarize the lesson by conducting a class discussion about what to look for when comparing distributions. Students should be precise when estimating values of means, medians, MAD, and IQR. They should also be able to comment on when it is most appropriate to use each measure of center and spread. *Answer: If a distribution is symmetric, it is best to use the mean as a measure of center and the MAD as a measure of spread. If a distribution is skewed, or has outliers, it is best to use the median as a measure of center and the IQR as a measure of spread.*

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Similar to the activities they did during class today, for homework, students should complete the *Comparing Fuel Efficiency with Boxplots* handout (LMR\_U2\_L6\_C).

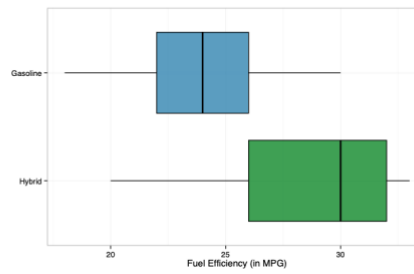
**Note:** Page 2 of the handout is an answer key for teacher reference only.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

### Comparing Fuel Efficiency with Boxplots

**Background:**

Data were collected on 50 cars – some were hybrids and others were regular gasoline vehicles. The fuel efficiency, in miles per gallon (MPG), is displayed in the boxplots below for each type of vehicle.



LMR\_U2\_L6\_C

## Lesson 7: Plot Match

### Objective:

Students will learn how to create a boxplot from an already-established dotplot.

### Materials:

1. From Dotplots to Boxplots handout (LMR\_U2\_L7\_A)
2. Sets of plots from Plot Match file (LMR\_U2\_L7\_B) – one for each team

**Advance preparation required:** Needs to be cut out prior to class (see step 8 in lesson)

### Vocabulary:

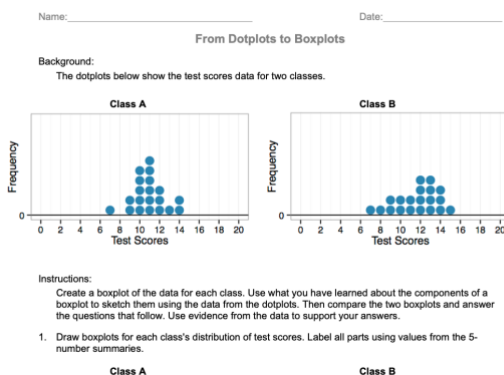
representation

**Essential Concepts:** Boxplots are an alternative visualization of histograms and dotplots. They capture most, but not all, of the features we can see in a dotplot or histogram.

### Lesson:

1. Ask students to complete an *Entrance Slip* by recalling the components of the five-number summary that make up a boxplot. **Answer:** Five-number summary: minimum, 1<sup>st</sup> quartile ( $Q_1$ ), median, 3<sup>rd</sup> quartile ( $Q_3$ ), maximum.
2. Randomly select students to share the components and briefly discuss what each means in a boxplot. If students are missing a component, ask them to add the component to their list.
3. Remind students that during Lesson 5, they created a boxplot from students' heights.
4. Explain that a boxplot is one **representation** of the distribution of a variable in a dataset. They have worked with other representations of distributions. Ask students:
  - a. What other representations of distributions have we seen? **Answers may include:** dotplots, bar charts, scatterplots, histograms, and tables.
5. Distribute the *From Dotplots to Boxplots* handout (LMR\_U2\_L7\_A). In teams, students will sketch boxplots from dotplots. They will need to determine the five-number summaries of each plot and should clearly label each value on their boxplots.

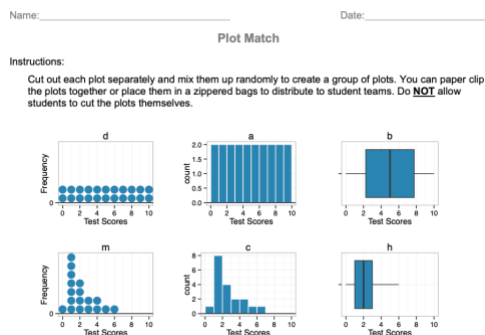
**Note:** Page 3 of the handout is an answer key for teacher reference only.



LMR\_U2\_L7\_A

6. Students should answer the 3 questions included in the handout. They can discuss their answers in pairs, and then have a class share out of the responses.
7. Once the discussion wraps up, inform the students that they will now attempt to find plots that represent the same data but are plotted differently.
8. Distribute one set of plots, from the *Plot Match* file (LMR\_U2\_L7\_B), to each student team.  
**Advance preparation required:** Each student team will receive a set of plots containing all 15 plots from the *Plot Match* file (LMR\_U2\_L7\_B). Copies will need to be cut and sorted prior to class. To keep the plots together, you can either paper clip them or place them in zippered bags.

**Note:** Do not distribute the handout for students to cut out the plots!



LMR\_U2\_L7\_B

9. Inform students that they are now going to gather in their teams and practice matching different representations of distributions. Each group will receive 15 plots (5 dotPlots, 5 histograms, and 5 boxplots). Their task is to determine which dotplots, histograms, and boxplots represent the same data.
10. Once each group has decided on their 5 groupings, engage the students in a class share out until all students agree. Then, have the students record their responses to the following statements and/or questions in their DS journals:
  - a. What types of data are best for using a histogram? *Answer: Histograms are useful for almost any type of data. They can easily show the shape of a distribution (including skewness and multiple peaks). They are usually best with larger datasets.*
  - b. What types of data are best for using a dotplot? *Answer: Dotplots can also easily show the shape of a distribution. They are preferred over histograms when there is a relatively small amount of data.*
  - c. What types of data are best for using a boxplot? *Answer: Boxplots are useful when the distribution has one mode (one peak). They are also useful to describe data that are heavily skewed or that contain outliers.*
  - d. Describe some characteristics of data that become hidden when a boxplot is used instead of a dotplot or histogram. *Answer: Dotplots and histograms can show the number of modes in a distribution, but a boxplot cannot. If a distribution is bimodal, we will not be able to tell in a boxplot. In general, we lose the ability to talk about the overall shape of the distribution.*
11. Display the uncut version of the *Plot Match* file (LMR\_U2\_L7\_B) so that students see the letters that correspond to each set of representations.

**Solution key:**

Set 1: Plots (d), (a), (b)

Set 4: Plots (e), (l), (i)

Set 2: Plots (m), (c), (h)

Set 5: Plots (n), (k), (g)

Set 3: Plots (f), (j), (o)

12. Have a few students share out their responses. For homework, students will record some pros and cons of using different types of graphical representations to display the same data.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework & Next Day



Students should reflect on today's class discussion and record their ideas of some pros and cons of using different types of graphical representations to display the same data.

### Lab 2B: Oh the Summaries...

Complete Lab 2B prior to the Practicum.

## Lab 2B: Oh the Summaries ...

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Just the beginning

- Means, medians, and MAD are just a few examples of *numerical summaries*.
- In this lab, we will learn how to calculate and interpret additional summaries of distributions such as: minimums, maximums, ranges, quartiles and IQRs.
  - We'll also learn how to write our first custom function!
- Start by loading your *Personality Color* data again and name it `colors`.

### Extreme values

- Besides looking at *typical* values, sometimes we want to see *extreme* values, like the smallest and largest values.
  - To find these values, we can use the `min`, `max` or `range` functions. These functions use a similar syntax as the `mean` function.
- **(1) Find and write down the `min` value and `max` value for your predominant color.**
- **(2) Apply the `range` function to your predominant color and describe the output.**
  - The *range* of a variable is the difference between a variable's smallest and largest value.
  - Notice, however, that our `range` function calculates the maximum and minimum values for a variable, but not the difference between them.
  - Later in this lab you will create a custom `Range` function that will calculate the difference.

### Quartiles (Q1 & Q3)

- The **median** of our data is the value that splits our data in half.
  - Half of our data is smaller than the *median*, half is larger.
- Q1 and Q3 are similar.
  - 25% of our data is smaller than Q1, 75% are larger.
  - 75% of our data is smaller than Q3, 25% are larger.
- **(3) Fill in the blanks to compute the value of Q1 for your predominant color.**  
`quantile(~____, data = _____, p = 0.25)`
- **(4) Use a similar line of code to calculate and write down Q3, which is the value that's larger than 75% of our data.**

### The Inter-Quartile-Range (IQR)

- **(5) Write and run code making a `dotPlot` of your predominant color's scores. Make sure to include the `nint` option.**
- Visually (don't worry about being super-precise):
  - **Cut the distribution into quarters so the *number of data points* is equal for each piece. (Each piece should contain 25% of the data.)**
    - Hint: You might consider using the `add_line(vline = )` to add vertical lines to the quarter marks.
  - **(6) Write down the numbers that split the data up into these 4 pieces.**
  - **(7) How long is the interval of the middle two pieces?**
  - This length is the **IQR**.

## Calculating the IQR

- The `IQR` is another way to describe *spread*.
  - It describes how *wide* or *narrow* the middle 50% of our data are.
- Just like we used the `min` and `max` to compute the `range`, we can also use the 1st and 3rd quartiles to compute the `IQR`.
- (8) Use the values of Q1 and Q3 you calculated previously for your predominant color and find the `IQR` by hand.
  - (9) Then write and run code using the `iqr()` function to calculate it for you.
- (10) Which personality color score has the widest spread according to the `IQR`? Which is narrowest?

## Boxplots

- By using the medians, quartiles, and min/max, we can construct a new single variable plot called the **box and whisker** plot, often shortened to just **boxplot**.
- (11) By showing someone a `dotPlot`, how would you teach them to make a *boxplot*? Write out your explanation in a series of steps for the person to use.
  - (12) Use the steps you write to create a sketch of a *boxplot* for your predominant color's scores in your journal.
  - (13) Then write and run code using the `bwplot` function to create a *boxplot* using `R`.

## Our favorite summaries

- In the past two labs, we've learned how to calculate numerous *numerical summaries*.
  - Computing lots of different summaries can be tedious.
- (14) Fill in the blanks below to compute some of our *favorite summaries* for your predominant color all at once.

```
favstats(~____, data = colors)
```

## Calculating a range value

- We saw in the previous slide that the `range` function calculates the maximum and minimum values for a variable, but not the difference between them.
- We could calculate this difference in two steps:
  - Step 1: Use the `range` function to `assign` the max and min values of a variable the name `values`. This will store the output from the `range` function in the *Environment* pane.

```
values <- range(~____, data = ____)
```
  - Step 2: Use the `diff` function to calculate the difference of `values`. The input for the `diff` function needs to be a vector containing two numeric values.

```
diff(values)
```
- (15) Use these two steps to calculate and write down the *range* of your predominant color.

## Introducing custom functions

- Calculating the *range* of many variables can be tedious if we have to keep performing the same two steps over and over.
  - We can combine these two steps into one by writing our own custom `function`.

- Custom functions can be used to combine a task that would normally take many steps to compute and simplify them into one.
- The next slide shows an example of how we can create a custom function called `mm_diff` to calculate the absolute difference between the `mean` and `median` value of a `variable` in our `data`.

### Example function

```
mm_diff <- function(variable, data) {
 mean_val <- mean(variable, data = data)
 med_val <- median(variable, data = data)
 abs(mean_val - med_val)
}
```

- The function takes two *generic* arguments: `variable` and `data`.
- It then follows the steps between the curly braces `{ }`.
  - Each of the *generic* arguments is used inside the `mean` and `median` functions.
- Copy and paste the code above into an *R Script* and run it.
- The `mm_diff` function will appear in your *Environment* pane.

### Using mm\_diff()

- After running the code used to create the function, we can use it just like we would any other numerical summary.
  - (16) In the *console*, fill in the blanks below to calculate the absolute difference between the `mean` and `median` values of your predominant color:

```
____(~____, data = ____)
```

- (17) Which of the four colors has the largest absolute difference between the `mean` and `median` values?
  - (18) By examining a `dotPlot` for this personality color, make an argument why either the `mean` or `median` would be the better description of the center of the data.

### Our first function

- (19) Using the previous example as a guide, create a function called `Range` (note the capital 'R') that calculates the *range* of a variable by filling in the blanks below:

```
____ <- function(____, ____) {
 values <- range(____, data = ____)
 diff(____)
}
```

- (20) Use the `Range` function to find the personality color with the largest difference between the `max` and `min` values.

### On your own

- (21) Write a function called `myIQR` that uses the `quantile` function to compute the middle 30% of the data.

## Practicum: The Summaries

### Objective:

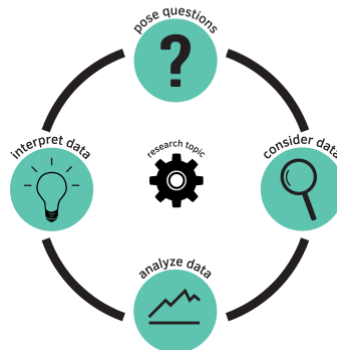
Students will engage in their first statistical investigation using the Data Cycle. They will pose a statistical question based on a dataset from a Participatory Sensing campaign, analyze, and interpret the data.

### Materials:

1. *The Summaries Practicum* (LMR\_U2\_Practicum\_The\_Summaries)

**Note to Teacher:** Before assigning the practicum to your students, engage the class in a discussion about the sample statistical questions below. Guide the discussion so that students identify not only the groups being compared in each question, but also what is being compared about the groups. Remind them of the Data Cycle from Unit 1.

### The Data Cycle



### Practicum

### The Summaries

Using the *Food Habits* campaign data or *Personality Color* survey data, develop a new statistical question that compares two or more groups. Some sample statistical questions (about other datasets) are below:

- Which sex shows a bigger range in age, male or female Oscar winners?  
**Grouping variable:** *sex (male, female)*  
**Variable:** *ages*
- Do children, teenagers, or adults spend more money on candy?  
**Grouping variable:** *age group (child, teenager, adult)*  
**Variable:** *the amount of money spent on candy*
- How does the median height of teenage males compare to that of females?  
**Grouping variable:** *sex (male, female)*  
**Variable:** *height*
- How do the average temperatures of Los Angeles, Las Vegas, and San Francisco compare?  
**Grouping variable:** *city (Los Angeles, Las Vegas, San Francisco)*  
**Variable:** *daily maximum temperature*

Remember, a statistical question is one that anticipates variability in the question and then addresses the variability in the answer:

Based on the data you chose (*Food Habits* or *Personality Color*), you need to:

1. Write down your question and think about ways you could answer it using RStudio.
2. Describe the data you are using to answer your question and explain why it is appropriate.
3. Analyze the data to provide evidence that supports the answer to your question. Include plot(s) and numerical summaries (mean, median, MAD, IQR, etc.) related to your plots.

4. Interpret the data to answer your statistical question. You should:
  - a. Provide the plot(s) and numerical summaries related to your plot(s).
  - b. Describe what the plot shows.
  - c. Explain why you chose to make that particular plot(s) and the related numerical summaries.
  - d. Explain how the plot and numerical summary answer your statistical question.
5. Write and submit a one-page report.

**Note:** You may use the scoring guide in Unit 1 to give you an idea of how to score the Practicum.

# How Likely Is It?

Instructional Days: 7

## Enduring Understandings

Probability measures the long run frequency of occurrence for chance outcomes. Probabilities can be approximated by designing and conducting simulations, and also via mathematical calculation.

## Engagement

Students will watch a scene from the movie *Rosencrantz and Guildenstern are Dead* and discuss the likelihood of getting “heads” when tossing a coin 78 times in a row. The scene can be found at: <https://www.youtube.com/watch?v=NblnZ5oJ0bc>

## Learning Objectives

### Statistical/Mathematical:

S-CP 2: Understand that two events A and B are independent if the probability of A and B occurring together is the product of their probabilities, and use this characterization to determine if they are independent.

S-CP 9: (+) Use permutations to perform [informal] inference.

\*This standard will be addressed in the context of data science.

S-IC 6: Evaluate reports based on data.

### Data Science:

Understand how algorithms are used to design simulations.

### Applied Computational Thinking using RStudio:

- Design and conduct simulations in RStudio.
- Compare actual data to simulated data using RStudio.
- Re-randomization of permuted data.
- Use estimated probabilities from samples to determine theoretical probabilities.

### Real-World Connections:

Learn to use simulations to determine expectations of events.

## Language Objectives

1. Students will construct summary statements about their understanding of probability.
2. Students will use a Venn Diagram to compare and contrast compound events and their probabilities.
3. Students will engage in group presentations to express their understanding of simulations to assess probabilities.

## Data File or Data Collection Method

### Data Files:

Simulated data in RStudio.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 8: How Likely Is It?

### Objective:

Students will understand the basic rules of probability. They will learn that previous outcomes do not give information about future outcomes if the events are independent.

### Materials:

1. Video: "Heads" from the movie *Rosencrantz and Guildenstern are Dead* found at: <https://www.youtube.com/watch?v=NblnZ5oJ0bc>
2. Projector for RStudio functions

### Vocabulary:

probability, simulation, model, sample proportion, chance, independence

**Essential Concepts:** Probability is an area that we humans have poor intuition about. Probability measures a long-run proportion: 50% chance means the event happens 50% of the time *if you repeated it forever*. When we don't repeat forever, we see variability.

### Lesson:



1. Ask students to consider possible synonyms to the word **chance**. If someone says, "That just happened by chance," what does that mean? *Synonyms: possibility, prospect, expectation, unintentional, unplanned. The actual definition of chance is "a possibility of something happening."*
2. Then ask them which game - chess or the board game, "Sorry" - is more based on chance? Why? **Note:** any game can be chosen. *"Sorry" is more based on chance because many outcomes are determined by dice rolls. In chess, there are certain strategies and movements that can be planned, so it is more a game of skill. With Sorry, the players roll a die (number cube), so the numbers they roll have an impact on how well they do in the game.*
3. Next ask students if they can think of situations where chance is the only force at play. *Possible responses: card games, slot machines, the lottery, coin flipping, and rock-paper-scissors.*



4. Play the "Heads" video from the movie *Rosencrantz and Guildenstern are Dead* found at: <https://www.youtube.com/watch?v=NblnZ5oJ0bc>.



5. In their DS journals, ask students to write down their initial reactions to the clip by responding to the following questions:
  - a. Is it *possible* to get 78 heads in a row when tossing a coin? *Answer: Yes, it is possible to get 78 heads in a row since one coin toss does not determine the next coin toss.*
  - b. Do you think it is *likely* to get 78 heads in a row? *Answer: No, although it is possible to get 78 heads in a row.*
  - c. How many times should we get heads when tossing a coin? *Answer: 1 out of 2 times or 50% of the time.*
  - d. On average, how many times out of the 78 tosses should the characters have gotten heads? *Answer: Roughly about 39 times.*
6. Ask students to discuss their findings with their team members and come to an agreement on their responses. Afterwards, conduct a *Whip Around* and ask each team to share its findings. Are there any differences between the teams? Any similarities?
7. As teams share their responses, students should add to or revise their individual findings in their DS journals.
8. Explain to students that, from the concept of chance, we can start learning about **probability**. Chance is simply the possibility that something will happen, and probability is a measurement for how often something happens in the "long run." Students may have ideas about how to calculate probabilities based on prior classes or knowledge but inform them that IDS will be taking a different approach by using simulations (see next step).

9. Since we don't want to actually flip a coin 78 times like the actors did in the video, we can have RStudio simulate them for us. A **simulation** is a way of creating random events that are close to real-life situations without actually doing them. It is a kind of **model**, which is a way of representing real world situations so that predictions can be made.
10. Explain to students that R has a function that does coin flipping for us and that it assumes an equal probability of heads and tails. Using a projector to display your computer screen to the whole class, demonstrate how to do one simulation of a coin flip in RStudio. Use the following function:

**> rflip(1)**

11. Explain that the value of 1 in the argument part of the function tells R to flip the coin 1 time. If we want to flip the coin 10 times, we could simply change the function to **rflip(10)**.



12. Run the function again using 10 as the number of times to flip the coin. Ask students:
- How many heads ("H"s) were there? *Answers will vary for each sample.*
  - How many Tails ("T"s) were there? *Answers will vary for each sample.*
  - In the output, what does **Flipping 10 coins [Prob(Heads) = 0.5]** mean? *Answer: This is RStudio telling us that we are tossing the coin 10 times and that the probability of getting heads should be 0.5 (it is flipping a fair coin).*
  - In the output, what does **Number of Heads: 3 [Proportion Heads: 0.3]** mean? *Answer: This is RStudio telling us that in our sample, we got heads 3 out of the 10 times we flipped the coin. The sample proportion is automatically calculated for us by dividing the number of heads by the total number of tosses (in this case, 3/10=0.3).*

**Note:** This is example of an output. Your sample may have a different value for the number of heads that appeared, and thus a different value for the proportion of heads.



13. To relate back to the video at the beginning of class, repeat the simulation once more, but use 78 as the number of coin flips **rflip(78)**. Ask students:
- How many heads ("H"s) were there? Since we know to expect about 39 heads if the coin is fair, does the value seem reasonable? *Answers will vary for each sample. Most likely, you will see values near 39.*
  - How many Tails ("T"s) were there? *Answers will vary for each sample.*
  - What proportion of the coin flips were heads? *Answers will vary for each sample.*

14. Using the **rflip(78)** command, run the simulation 3-5 more times and have students record the values for the number of heads and the proportion of heads.

*As an example, we ran the function 3 times and saw the following values:*

*Sample 1 - number of heads: 45*

*proportion of heads: 0.577*

*Sample 2 - number of heads: 33*

*proportion of heads: 0.423*

*Sample 3 - number of heads: 42*

*proportion of heads: 0.538*



15. Have students answer the questions listed below. The important thing to note is that the values can and (almost always) WILL change each time you run the simulation to create a new sample.
- How do the proportions of heads in the samples compare to each other? *Answers will vary.*
  - How do the proportions of heads compare to the true probability of heads (1/2 or 50%)? *Answers will vary, but students should notice that most of the probabilities are close to 50%.*
  - Why is there a 50% chance of getting heads during each coin flip? *Answer: Since there are two sides to a coin, both should be equally likely to come up. So there is a 1 out of 2 chance of getting heads and 1 out of 2 chance of getting tails.*

16. Ask students to engage in a discussion with their group about the statement below, then have a few group reporters share out.
  - a. If a coin was flipped 78 times, I would claim that the coin is unfair if I got less than # heads or more than # heads.
17. Inform students that you are going to perform 500 simulations. Each simulation represents a coin being flipped 78 times. For each simulation, the computer will record the number of heads in the 78 flips. A histogram will be created that represents the number of heads in each of the simulations. The histogram is a model that will display what typically happens when a fair coin is flipped 78 times.
18. Copy and paste the code below in an RScript and run each line of code, one at a time, for the students:
 

```
set.seed(11) #reproducibility
flips <- do(500)*rflip(78)
View(flips) # 4 variables
histogram(~heads, data = flips)
favstats(~heads, data = flips)
```
19. Engage the students in a discussion about the histogram:
  - a. What is the distribution telling us? *Answer: When flipping a fair coin 78 times, what typically happened was that it landed on heads between 36 and 42 times ( $36/78 = 0.46$  to  $42/78 = 0.54$ ). It was not uncommon for the coin to land on heads 28-35 ( $0.36$ - $0.45$ ) times or 43-51 ( $0.55$ - $0.65$ ) times. What was very uncommon, however, was landing on heads less than 28 times (less than 36%) or more than 51 times (more than 65%).*
  - b. Were the group's cut-offs (item #16) similar to what the chance model displayed? *Answers will vary. Some groups' intervals might be very wide and others very narrow.*
  - c. Use the chance model (histogram) which displays what typically happens when a fair coin is flipped 78 times, make a call for the scenarios below – fair or unfair?
    - i. You flip a coin 78 times and get 37 heads. *Answer: Fair. 37 was very common based on the histogram.*
    - ii. You flip a coin 78 times and get 46 heads. *Answer: Fair. 46 was less common but not too uncommon.*
    - iii. You flip a coin 78 times and get 20 heads. *Answer: Unfair. In the 500 simulations, not once did we see a FAIR coin land on heads 20 times.*



20. Next, pose the following question:
  - a. If you get a heads on the first toss of a coin, will you definitely get a heads on the next toss? Will you definitely get a tails on the next toss? *Answer: No. One coin toss should not affect another coin toss. Each time you flip the coin, the chances of getting heads versus tails remains the same.*
21. Introduce the concept of **independence**. Explain that, when tossing a fair coin, there is no relationship between each toss. The second toss does NOT depend on the first toss; therefore, the coin tosses are independent of each other.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students will create a Tweet (they do not have to post it online). Using 280 characters or fewer, write a Tweet about the meaning of probability.

## Lesson 9: Dice Detective

### Objective:

Students will learn how to use simulations to detect an unfair die.

### Materials:

1. Poster paper – with 6 columns labeled 1, 2, 3, 4, 5, 6
2. 2 dice (number cubes)

**Note:** You can use regular hard dice, or soft foam dice (can be found at dollar stores)

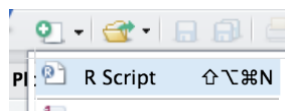
3. Projector for RStudio functions

**Essential Concepts:** In the short-term, actual outcomes of chance experiments vary from what is 'ideal'. An ideal die has equally likely outcomes. But that does not mean we will see exactly the same number of one-dots, two-dots, etc.

### Lesson:

1. In pairs, ask students to quickly share their Tweets from the previous lesson's homework. Collect the Tweets and select a few to share with the class. Of the Tweets shared, ask students which one is closest to the definition: probability measures how often something happens in the "long run."
2. Remind students that, during the previous lesson, they were introduced to simulations. The motivation for using simulations is that we can use the calculated sample proportions to estimate probabilities of real-life events.
3. During today's lesson, we will be continuing to learn about probability and simulations to determine if an event is not fair (one example: a coin is weighted and lands on heads more often than tails).
4. Ask students what they know about dice (number cubes). If they have never heard of them, show one to the class and explain how it works. *Answer: A die (number cube) is a 6-sided cube. Each face of the cube is labeled with dots to represent a number between 1 and 6. For example, if the face has 3 dots, then it represents the number 3. The cube itself is weighted so that there is an equal probability of rolling each of the 6 numbers.*
5. Have a discussion about what the students would expect the probability of rolling the number 1 should be if a die (number cube) were tossed into the air and allowed to fall back to the ground (or table). *Answer: Since there are 6 numbers on the die, each number should be equally likely to occur, so the probability of rolling a 1 is  $1/6$ .*
6. Display a piece of poster paper on the board with columns labeled 1, 2, 3, 4, 5, 6. Explain that each column represents the numbers on the die (number cube). We will be using the poster to tally the results of actual dice (number cube) rolls.
7. Select two students to be dice (number cube) rollers and give each student one die. As noted in the Materials section above, you can have the students use either regular hard dice, or softer foam ones (can be found at dollar stores).
8. Tell the class that each student roller will be rolling the dice 6 times (so there will be a total of 12 rolls for our sample). Ask:
  - a. If they are rolling the dice 6 times, how often do you think Student 1 will roll a 3? *Answer: Out of 6 rolls, we would expect to see each of the numbers one time, so we will most likely see about one 3 for Student 1.*
  - b. Would you expect the Student 2 to roll a 3 just as often? Why? *Answer: Yes, we should expect the same thing from Student 2 because we have independent events. There are actually two ways that independence plays a part here: (1) each student is independent from the other and has no effect on what the other will roll, (2) the 6 die rolls for Student 1 are all independent of each other because each face of the cube has an equal chance of happening on any given roll. So, if Student 1 gets a 3 during his/her first roll, that doesn't give us any information about what he/she will get on the second roll.*

- c. Since we will have 12 rolls (and therefore 12 samples), how many tally marks should we expect in each column on the chart? *Answer: We would expect to see 2 tally marks in each column (each number will probably be rolled twice).*
9. Have each student roller toss his/her die one time and share the outcomes with the rest of the class. As they do this, place a tally mark in the corresponding column on the chart. Repeat this process 5 more times so that each student has a total of 6 rolls.
10. As a class, observe the results in the chart and discuss the following:
  - a. Do the data from these 12 rolls match what we expected (see responses from step 8 above)? Is this surprising? *Answers will vary by class. Some values may have shown up more than we expected (example: the number 3 was rolled 3 times), and others may not have been rolled at all (example: the number 5 was never an outcome). We only have a small sample of data, so it's not surprising for our results to vary from the expected outcomes.*
  - b. If the data do not match our expectations, does this mean the dice are unfair in some way? *Answer: Even if they don't match our expectations, this does not mean the dice are unfair – we simply don't have enough data yet to know. We would need to roll the dice more.*
  - c. If we wanted to purposely create an unfair die, what are some ways we could achieve that? *Answers will vary by class. Some examples include: (1) We could add tape to one face of the die to give that side more weight. This would increase the chances of the number that is directly opposite of it appearing because the die will land on the heavier side more (and therefore the side facing up will be the number opposite). (2) We could chip the edge of one corner of the die. This would throw off the original balance and favor certain sides.*
11. Similar to the previous day's lesson with coin flipping, we can also simulate dice rolls in RStudio. The function required is called `roll_die()`. The arguments for this function are a bit different than the `rflip()` function from yesterday. We cannot simply put `roll_die(1)` for the computer to roll a die one time. Instead, the function was built with 2 possible dice to choose from – die A and die B.
12. Inform the students that one of the dice in the function is fair and the other is unfair. A die is unfair if it favors one outcome over another. They will attempt to determine which dice is the unfair one by doing multiple simulations.  
**Note to Teacher:** Many simulations require multiple functions, or code, to perform. This is where RScripts are helpful. An RScript can be used to test code, write notes, and let us easily execute multiple lines of code at one time. This would be a good place to introduce students to RScripts.



13. Using a projector to display your computer screen to the whole class, demonstrate how to open an RScript. Type the following function on your script and click Run. Run simply pastes the function onto the console.
 

```
roll_die("A", times = 1)
```
14. The output will show one number that represents what value on the die the computer rolled. Go back to your script and modify the function to roll die A 12 times.
 

```
roll_die("A", times = 12)
```
15. Compare the results of these 12 simulated rolls to the results of the 12 actual rolls completed by the two students during step 9 above. If there is space available on the tally chart, you can add the computer results to it for an easy comparison.

16. Ask students how we could record data from these simulations if we wanted to roll the die 100 times. Would they want to hand count the number of times each value occurred? Is there a function in RStudio that will count them for us? *Answer: It would be difficult to count every individual value in the output on the screen. However, we can use the `tally()` function to find out how many times each die value appeared.*
17. To make using the `tally()` function easier, we should assign a name to each simulation so we don't have to type the entire function multiple times. We can also have it calculate the proportions for each value. Add the functions below to your RScript and run them one at a time.

```
sample1 <- roll_die("A", times = 12)
tally(sample1)
tally(sample1, format = "proportion")
```

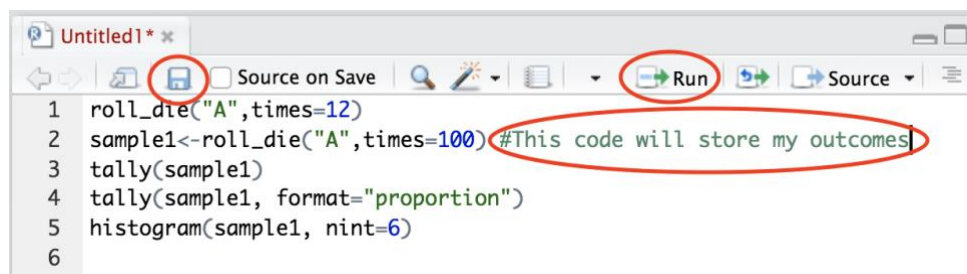
18. Remind the students that if the die is fair, then each side of the die should appear roughly the same number of times. Therefore, the proportions should be fairly similar to each other and to the true probability of 1/6.

19. Add the function below to your script, but before running it, ask the students what they think a histogram of the simulated data might look like and then run the command on your screen.

**Note:** Be sure to include the argument `nint = 6` so that the resulting histogram has six bars. *Answer: If the die is fair, each of the bars in the histogram should be roughly the same height.*

```
histogram(sample1, nint = 6)
```

**Note to teacher:** Show students how to save an RScript. Inform students that they can take notes on their RScript by including a hashtag (also known as a pound sign or #) at the beginning of the note. Data scientists refer to these types of notes as “code comments” or simply “comments”. See image below.



The Script will be stored in the files tab. To run each function individually, place your cursor on the line and hit the Run button. To run multiple lines of code at once, highlight them and hit Run.

20. Allow the students to access their school computers now to start creating their own simulations in an RScript using die B. Students can pair up, if needed. Have them begin by asking RStudio to roll the die 100 times. They should note their output from both the `tally()` function and the histogram. They can then compare the results to those from step 14 above. Are they similar? Can they determine which die is unfair yet? *Answers will vary by class. The results will be similar, but not exact. With the sample sizes of each simulation being fairly small, we cannot see a clear difference between the two dice yet.*
21. Let the students explore by changing the number of times RStudio rolls the dice. Remind them that the goal is to determine which of the two dice is biased. The sample sizes need to be very large in order for them to see a clear difference between the 2 histograms. The pattern becomes more visible when `times = 2000`.

**Note:** The maximum value for `times` within the `roll_die()` command is 2000. Simulations can be combined using the concatenate function `c()`. For example, suppose `s1` represents 2000 rolls of die A and `s2` is a second sample of 2000 rolls for die A. To combine these two samples the following can be used `more_rolls <- c(s1, s2)`.



22. When students have had enough time to make a decision regarding which dice is unfair and how, engage the class in a discussion to verify that everyone agrees. *Answer: Die B is unfair; Die A is fair. Die B favors the number 3.*
23. Then, steer the conversation towards why the sample size affected the results. *Answer: The sample size needed to be large because the difference between the probabilities of the die rolls was very small. In order to detect small differences, we must have larger sample sizes.*

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students will consider a four-sided die and imagine rolling it 20 times. They should sketch a histogram of (a) the ideal, expected outcome, (b) an outcome that they think is “realistic,” and (c) an outcome they might see if the die were biased to produce more 4’s.

## Lesson 10: Marbles, Marbles...

### Objective:

Students will understand that random events vary, so that the percentage predicted by a probability isn't exact, but varies. Students practice converting percentages to proportions.

### Materials:

1. For each student team: 50 marbles – 20 of one color, and 30 of another color

**Note:** Marbles can be substituted for other materials, such as small blocks, as long as they are the same size.

### Vocabulary:

proportion, percentage, event, with replacement, without replacement

**Essential Concepts:** There are two ways of sampling data that model real-life sampling situations: with and without replacement. Larger samples tend to be closer to the “true” probability.

### Lesson:

1. Remind students that, during the previous two lessons, they learned how to estimate probabilities for a population with the help of simulations to create sample data. Both lessons had nice, prepackaged functions already available in RStudio, which made the simulations fairly quick and easy – in Lesson 8, the **rflip()** function was used to simulate flipping a coin; and in Lesson 9, the **roll\_die()** function was used to simulate rolling one of two dice.
2. But what if we don't have a nice function to perform a simulation for us? Can we create our own method? Yes! We will actually learn to create a simulation from scratch during Lab 2C.
3. Ask students:
  - a. If you have a bag of 50 marbles, where 20 of them are blue and 30 of them are red, what is the probability of drawing a red marble? *Answer: 30/50 or 60%.*
4. Select a student to answer the question. Ask the class if they agree or disagree. If they agree, ask them to raise their hand. If there are students who disagree, lead a class discussion until a consensus is reached.
5. Ask students to share their strategies on how to convert the **proportion** into a **percentage**. As strategies are being shared, students should take notes in their DS journals. Review how to turn fractions into percentages, if necessary.
6. Ask students:
  - a. What if we actually drew out one marble, recorded its color, then replaced it back in the bag, and did this 10 times? *Answers will vary by class.*
  - b. Would the percentage of red marbles in this sample necessarily be exactly the same as the probability? Identify that each time a marble is drawn, we are creating an **event**. *Answers will vary by class.*
7. Distribute the bags of marbles to each team. Ask each team to:
  - a. Shake the bag of marbles.
  - b. Draw one marble out of the bag.
  - c. Record the marble's color in their DS journal.
  - d. Replace the marble back into the bag. Inform them that this is called sampling **with replacement**. Ask them to consider what “with replacement” means and discuss with the class. *Answer: “With replacement” means that after you select a marble from the bag, you have to put it back into the bag (replace it) before you select another marble.*

They should draw 10 marbles from the bag and record the observed colors.

8. Do a *Whip Around* to find out how many times each team drew a red marble out of their 10 draws. Have them calculate the corresponding sample proportions. For example, if one team drew 7 red marbles out of their 10 draws, their sample proportion is  $7/10 = 0.70$  (which is the sample as a sample percentage of 70%).

9. Ask students why the proportions are perhaps different from each other and from the “true” probability of drawing a red marble?



10. Have the student teams continue drawing marbles, one at a time and with replacement, until they have 50 events recorded. Discuss the following questions:

- How many times did they draw a red marble out of these 50? *Answers will vary by class.*
- What’s the corresponding sample proportion? Is it closer to the true probability than when you only drew 10 marbles? *Answers will vary by class. But they should notice that as the sample size increases, the sample proportion gets closer to the true population proportion.*



11. Engage students in a discussion about how the sample size affects the sample proportion. They should see that as they draw more marbles, their sample probability gets closer and closer to the true probability. If we were to continue drawing marbles forever, in the long run, our sample proportion should equal our true probability.

12. Have students consider what it might mean to sample **without replacement**. How would they do that with their bag of marbles? *Answer: “Without replacement” means that after you select a marble from the bag, you never put it back into the bag (do not replace it). Instead, you simply select another marble from the bag immediately. Students should recognize that, by not replacing the marble each time, the probabilities will change. This means each draw from the marble bag is NOT independent from another draw because removing one marble impacts the next event.*



13. *Exit Slip*: Based on this lesson, ask students to describe a sample, an event, and replacement.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Next Day

### Lab 2C: Which song plays next?

Complete Lab 2C prior to Lesson 11.

## Lab 2C: Which song plays next?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### A new direction

- For the past two labs, we've looked at ways that we can summarize data with numbers.
  - Specifically, you learned how to describe the *center*, *shape* and *spread* of variables in our data.
- In this lab, we're going to *estimate the probability* that a rap song will be chosen from a playlist with both rap and rock songs, if the choice is made at random.
  - The playlist we'll work with has 100 songs: 39 are rap and 61 are rock.

### Estimate what ... ?

- To *estimate the probability*, we're going to imagine that we select a song at random, write down its genre (*rock* or *rap*), put the song back in the playlist, and repeat 499 more times for a total of 500 times.
- The statistical investigative question we want to address is: *On average, what proportion of our selections will be rap?*
- **(1) Why do we put a song back each time we make a selection?**
- **(2) What would happen in our little experiment if we did not do this?**

### Calculating probabilities

- Remember that a *probability* is the long-run proportion of time an event occurs.
  - Many probabilities can be answered exactly with just a little math.
  - The probability we draw a single rap song from our playlist of 39 rap and 61 rock songs is **39/100, 0.39** or **39%**.
- Probabilities can also be answered exactly if we were willing to randomly select a song from the playlist, write down its *genre*, place the song back in the list, and repeatedly do this *forever*.
  - Literally, *forever*...
  - But we don't have that much time. So we're only going to do it 500 times which will give us an *estimate of the probability*.

### Estimating probabilities

- You might ask, *Why are we estimating the probability if we know the answer is 39%?*
  - Sometimes, probabilities are too hard to calculate with simple division as we did above. In which case, we can often program a computer to run an experiment to estimate the probability.
  - We refer to these programs as *simulations*.
- The techniques you learn in this lab could be applied to very simple probability calculations or very hard and complex calculations.
  - In both cases, your *estimated* probability would be very close to the *actual* probability.

### Getting ready

- Simulations are meant to mimic what happens in real-life using randomness and computers.
  - Before we can start simulating picking songs from a playlist, we need to simulate that playlist in **R**.

- Simulate our 39 rap songs using the repeat `rep()` function and assign the vector name `rap`.  

```
rap <- rep("rap", times = 39)
```
- Look in the *Environment* pane for the vector containing your rap songs.
- (3) Write and run a similar line of code to simulate the rock songs in our playlist of 100.

### Put the songs in the playlist

- Now that we've got some different songs, we need to combine them together.
  - To do this, we can use the combine function `c()` in R.
- (4) Fill in the blanks to combine your different songs:  

```
songs <- ____(rap, _____)
```
- And with that, our playlist of songs should be ready to go.
  - Type songs into the console and hit enter to see your individual songs.

### Pick a song, any song

- Data scientists call the act of choosing things randomly from a set *sampling*.
  - We can randomly choose a song from our playlist by using:  

```
sample(songs, size = 1, replace = TRUE)
```
- (5) Run this code 10 times and compute the proportion of "rap" songs you drew from the 10.
  - Vocabulary Check: A **proportion** is a fraction of the whole.
    - For example, if 2 rap songs were drawn from the 10, the *proportion* would be 2/10
    - It is more common to express a *proportion* as a decimal, in this case, 0.20
    - It is even more common to express a *proportion* as a percentage, 20%
- (6) Once everyone in your class has computed their *proportions*, calculate the *range of proportions* (the largest *proportion* minus the smallest *proportion*) for your class and write it down.

### Now do() it some more

- Instead of running the same line of code multiple times ourselves we can use R to `do()` multiple repetitions for us.
  - (7) Fill in the blanks below to do the `sample` code from the previous slide 50 times:  

```
do(____) * sample(____, ____ = ____, ____ = ____)
```
- Recall that we need to store our results to be able to perform analysis.
- (8) Write and run code *assigning* the 50 selected songs the name `draws` and then *View* your file.
- (9) What is the variable name?
  - R defaulted to naming the variable based on the function used. You may use the data cleaning skills you learned in Lab 1F to `rename` the variable if you wish.
- (10) Fill in the blank below to *tally* how often each genre was selected:  

```
tally(~____, data = draws)
```
- (11) Compute the proportion of "rap" songs for your 50 draws and find out if the *range* for your class's proportions is bigger or smaller than when we drew 10 songs.

## Proportions vs. Probability

- To review, so far in this lab we've:
  - Simulated a “playlist” of songs.
  - Repeatedly simulated drawing a song from the playlist, noting its genre and placing it back in the playlist.
  - Computed the proportion of the draws that were "rap".
- These proportions are all *estimates* of the theoretical probability of choosing a rap song from a playlist.
  - As we increase the number of draws, the *range* of proportions should shrink.

*When using simulations to estimate probabilities, using a large number of repeats is better because the estimates have less variability and so we can be confident we're closer to the actual value.*

## Non-random Randomness

- We've seen that random simulations can produce many different outcomes.
  - Some estimated probabilities in your class were smaller/larger relative to others.
- There are instances where you might like the same random events to occur for everyone.
  - We can do this by using `set.seed()`.
- For example, the output of this code will always be the same:

```
set.seed(123)
sample(songs, size = 1, replace = TRUE)
[1] "rap"
```

## Playing with seeds

- With a partner, choose a number to include in `set.seed` then redo the simulation of 50 songs.
  - Both partners should run `set.seed(____)` just before simulating the 50 draws.
  - The blank in `set.seed(____)` should be the same number for both partners.
  - (12) What value of `set.seed` did you and your partner use and what was the proportion of "rap" songs you obtained?
- Redo the 50 simulations one last time but have each partner choose a different number for `set.seed(____)`.
  - (13) Are the proportions still the same? If so, can you find two different values for `set.seed` that give different answers?

## On your own

- Suppose there are 1,200 students at your school. 400 of them went to the movies last Friday, 600 went to the park and the rest read at home.

*If we select a student at random, what is the probability that this student is one of those who went to the movies last Friday?*

- (14) Write and run code estimating the probability that a randomly chosen student went to the movies using 500 simulations.
  - (15) Write down both the code and the output that estimates the probability that a randomly chosen student went to the movies using 500 simulations. You might find it helpful to write your answer in an R Script (File -> New File -> R Script).
  - Include `set.seed(123)` in your code before you do 500 repeated samples.

## Lesson 11: This AND/OR That

### Objective:

Students will understand how AND/OR probabilities are defined and will be able to use frequency tables to compute these probabilities.

### Materials:

1. *Compound Probabilities* handout (LMR\_U2\_L11)
2. Blue sticky notes
3. Gold sticky notes
4. Four signs on the board labeled: *Pickles*, *Mayonnaise*, *Both*, *None* (in that exact order, and equally spaced across the length of the board)

### Vocabulary:

compound probabilities

**Essential Concepts:** What does “A or B” mean versus “A and B” mean? These are compound events and two-way tables can be used to calculate probabilities for them.

### Lesson:

1. Remind the students that they have been learning about estimating probabilities of single events based on sample proportions. Inform them that, today, they will learn how to calculate proportions when multiple events happen.
2. Review the basic idea of computing probabilities; in other words, the number of outcomes we are interested in divided by the total number of outcomes possible.



3. Pose the questions below to the class.

**Note:** They do not need to come up with specific answers; this is a time for them to make suggestions.

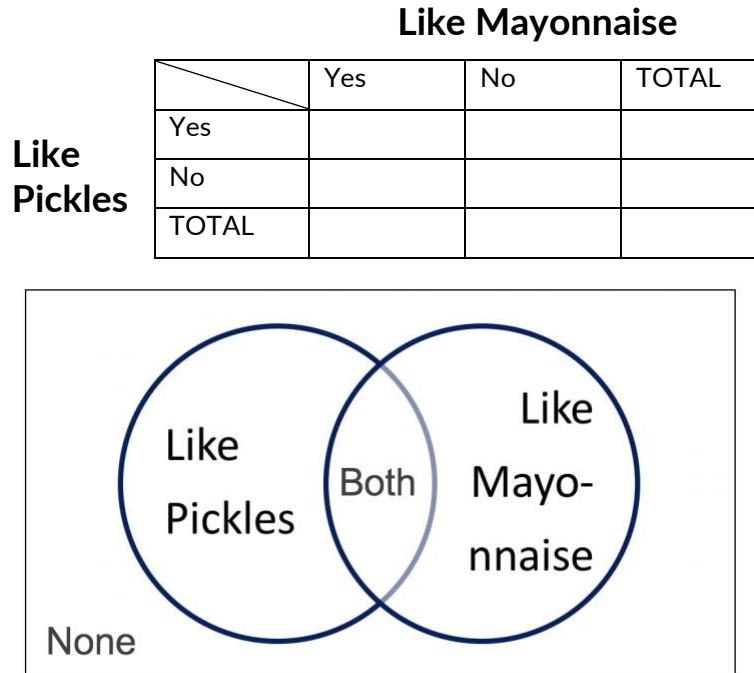
- a. How would we compute the probability of two outcomes occurring at the same time? For example, what is the probability that a randomly chosen student likes both pickles AND mayonnaise? *Answers will vary.*
  - b. How would we compute the probability of either of two outcomes occurring? For example, what is the probability that a randomly chosen student likes either pickles OR mayonnaise? *Answers will vary.*
4. For both questions, steer the students towards using the definition from step 2 above. That is, we want the students to realize that they can count the number of people that qualify for the given circumstance and divide by the total number of people to calculate a probability.
  5. In order to define AND/OR probabilities, students will participate in an activity where they are grouped by their food preferences.
  6. Divide the board into 4 groups and write the words “Pickles”, “Both”, “Mayonnaise”, and “None”, in that order, from left to right.
  7. Ask for 10 volunteers to stand by the word that represents their preferences. That is, if they only like pickles, they should stand by the word “Pickles”. If they like both pickles and mayonnaise, they should stand by the word “Both”.

**Note:** If all 4 groups do not have at least one student in them, select a few more students to stand at the board.

8. Ask the remaining students (those still seated) to count the total number of people standing by the board and have a student volunteer share the answer with the class. *Answers will vary by class.*

9. Next, create a 2-way frequency table like the one below to organize the values of student preferences as follows:
- Counts for students who like both go in the Yes/Like Mayonnaise and Yes/Like Pickles box.
  - Counts for students who like none go in the No/Like Mayonnaise and No/Like Pickles box.
  - Counts for students who like only mayonnaise go in the Yes/Like Mayonnaise and No/Like Pickles box.
  - Counts for students who like only pickles go in the No/Like Mayonnaise and Yes/Like Pickles box.

**Note:** A Venn diagram like the one below may be used as well, depending on student understanding and at teacher discretion.



10. Next, ask the students sitting down the following questions:

- a. How many students like both pickles AND mayonnaise? *Answers will vary by class.*
- b. What is the probability that a randomly selected student at the board likes both pickles AND mayonnaise? *Answers will vary by class. The probability should be calculated by dividing the number of people who are standing under "Both" (number given in step 9(a) above) by the number of students at the board (number given in step 8 above).*

$$\frac{\text{\# students under "Both"}}{\text{\# students standing at the board}}$$

11. Now, ask a student from the audience:

- a. How many students like pickles? *Answers will vary by class.*
- Note:** Avoid phrasing the question with "Students that like ONLY pickles." Students need to see that students who like "Both" items also belong to the groups liking the individual items. If students mistakenly report the number of students who like ONLY pickles, ask the people at the board to raise their hands if they like pickles and then ask the mistaken student to recount.
- b. What is the probability that a randomly selected student at the board likes pickles? *Answers will vary by class. The probability should be calculated by dividing the number of people who are standing under "Pickles" and "Both" by the total number of students at the board.*

$$\frac{(\text{\# students under "Pickles"}) + (\text{\# students under "Both"})}{\text{\# students standing at the board}}$$



12. Finally, ask one more student from the audience:

- a. How many students like pickles OR mayonnaise? *Answers will vary by class.*

**Note:** Avoid phrasing the question with “Students that like ONLY pickles OR ONLY mayonnaise.”

If students mistakenly report the number of students who like ONLY pickles plus the students who like ONLY mayonnaise, ask the people at the board to raise their hands if they like either pickles or mayonnaise (all students at the board should raise their hand except for the students who like “None”) and then ask the mistaken student to recount.

- b. What is the probability that a randomly selected student at the board likes pickles OR mayonnaise? *Answers will vary by class. The probability should be calculated by dividing the number of people who are standing under “Pickles”, “Mayonnaise”, and “Both” by the total number of students at the board.*

$$\frac{(\# \text{ students under "Pickles"}) + (\# \text{ students under "Mayonnaise"}) + (\# \text{ students under "Both"})}{\# \text{ students standing at the board}}$$

13. Informs students that AND/OR probabilities are called **compound probabilities**. In teams, have students record their own definitions of AND/OR probabilities based on the activity they just completed. *Answer: A compound probability is the probability of some combination of events occurring.*

14. Distribute the *Compound Probabilities* handout (LMR\_U2\_L11).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Compound Probabilities**

**Instructions:**  
As a class, fill out the following 2-way frequency table about ice cream preferences. Then answer the questions below.

|                    |       | Preferred Ice Cream Flavor |           |            | TOTAL |
|--------------------|-------|----------------------------|-----------|------------|-------|
|                    |       | Vanilla                    | Chocolate | Rocky Road |       |
| Sports Involvement | Yes   |                            |           |            |       |
|                    | No    |                            |           |            |       |
|                    | TOTAL |                            |           |            |       |

1. Show your work for each part of this question. What is the probability of randomly selecting:

- A student involved in sports?
- A student who is not involved in sports and prefers rocky road?
- A student who is involved in sports and likes vanilla, or a student who is not involved in sports and prefers chocolate?

LMR\_U2\_L11

15. Pass out a blue sticky note to each student who plays a sport and a gold sticky note to each student who does not play a sport.

16. Draw the table from the worksheet on the board (make it large and legible).

17. Have each student who plays a sport hold up their sticky note. Count them and record the number of students who play a sport in the appropriate row of the TOTAL column in the table.

18. Have each student who does not play a sport hold up their sticky note. Count them and record the number of students who do not play a sport in the appropriate row of the TOTAL column in the table.

19. Ask each student which of the following ice cream flavors they most prefer (each student must choose exactly one option): Vanilla, Chocolate or Rocky Road.

- Have the students write their ice cream preference on their sticky note.
- Fill out the remainder of the table by asking each group of students, those who play a sport and those who do not, to hold up their preference.
- Make sure the totals for preferred ice cream flavors and sports involvement add up to the same number.

20. Instruct the students to work in pairs to answer the questions on the *Compound Probabilities* handout (LMR\_U2\_L11).

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

**Homework & Next Day**

If not completed in class, students should finish the *Compound Probabilities* handout (LMR\_U2\_L11).

**Lab 2D: Queue it up!**

Complete Lab 2D prior to the Practicum.

## Lab 2D: Queue it up!

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Where we left off

- In the last lab, we looked at how we can use computer simulations to compute estimates of simple probabilities.
  - Like the probability of drawing a song genre from a playlist.
- We also saw that performing *more* simulations:
  - Took *longer* to finish.
  - Had estimates that *varied less*.
- In this lab, we'll extend our simulation methods to cover situations that are more complex.
  - We'll learn how to estimate their probabilities.
  - We also look at the role of sampling *with* or *without replacement*.

### Back to songs

- **(1) Write and run code simulating a playlist of songs containing 30 "rap" songs, 23 "country" songs and 47 "rock" songs.**
  - **Assign the combined playlist the name `songs`.**
- **(2) Write and run code simulating choosing a single song 50 times. Then use your simulated draws to estimate the probability of choosing a rap song.**
  - The "true" (theoretical) probability of choosing a *rap* song in this case is `0.30`.
  - **(3) Write a sentence comparing your estimated probability to the "true" probability.**

### With or Without?

- So far, you've selected songs *with replacement*.
  - We called it that, because each time you made a selection, you started with the same playlist. That is, you chose a song, wrote down its data, and then placed it back on the list.
- It's also possible to select *without replacement* by setting the `replace` option in the `sample` function to `FALSE`.
- **(4) Write and run code taking a sample of size 100 from our playlist of songs *without replacement*. Assign this sample the name `without`.**
  - **(5) Run `tally(without)` and describe the output.**
- **(6) Does something similar happen if you sample *with replacement*?**
  - Notice that the tilde `~` was not needed with the `tally` function. This is because `without` was not a variable within a data frame but rather a vector which acts like a lone variable.
- **(7) What happens if `size = 101` and `replace = FALSE`?**

## Sample with? Or without?

- Imagine the following two scenarios.
  1. You have a coin with two sides: *Heads* and *Tails*. You're not sure if the coin is fair and so you want to estimate the probability of getting a *Head*.
  2. A child reaches into a candy jar with 10 *strawberry*, 50 *chocolate* and 25 *watermelon* candies. The child is able to grab three candies with their hand and you're interested in the probability that all three candies will be chocolate.
- (8) Which of these scenarios would you sample *with replacement* and which would you sample *without replacement*? Why?
- (9) Write down the line of code you would run to *sample* from the candy jar. Assume the simulated jar is named *candies*.

## Simulations at work

- In reality, songs from a playlist are chosen without replacement.
  - This way, you won't hear the same song several times in a row.
- Let's write a more realistic simulation and estimate the probability that if we select two songs at random, without replacement, that both are rap songs.
  - (10) Write and run code using the *do* function to perform 10 simulated samples of size 2, without replacement and assign the simulations the name *draws* and then View your file. Use *set.seed(1)*.
  - (11) What are the variable names? What happened in the first simulation? Did any of your 10 simulations contain two rap songs?

## Simulations and probability

- To estimate the probability from our simulations, we need to find the proportion of times that the event we're interested in occurs in the simulations.
- In other words, we need to count the number of times the desired events occurred, divided by the number of attempts we made (the number of simulations).
- The next slides will show you two ways to do this.

## Counting similar outcomes

- One way we can estimate the probability of drawing two songs of the *same* genre is to use the following trick to count the number of *rap* songs in each of the 10 simulations:

```
mutate(draws, nrap = rowSums(draws == "rap"))
```

- (12) Let's break down the code above by running each part of the code one piece at a time. As you run each line of code below describe the output.

```
draws == "rap"
rowSums(draws == "rap")
mutate(draws, nrap = rowSums(draws == "rap"))
```

- Remember to assign a name to your mutated dataset.

## Counting other outcomes

- Another method we can use to estimate the probability of complex events is to use the following 2-step procedure:
  1. Subset or filter the rows of the simulations that match our desired outcomes.
  2. Count the number of rows in the subset and divide by the number of simulations.
- The result that you obtain is an estimate of the probability that a specific combination of events occurred.
- We'll see an example of this method on the next slide.

### Step 1: Creating a subset

- **(13) Fill in the blanks below to:**
  1. Create a subset of our simulations when both draws were "rap" songs.
  2. Count the number of rows in this subset.
  3. And divide by the total number of repeated simulations.

```
draws_sub <- filter(draws, ___ == "rap", ___ == "rap")
nrow(____) / ____
```

## Estimating probabilities

- **(14) Write and run code performing 500 simulations of sampling 2 songs from a playlist of 30 "rap", 23 "country" and 47 "rock" songs. You might consider running `set.seed` so that your results can be reproduced:**
- **(15) Calculate and write down the estimated probabilities for the following situations:**
  1. You draw two "rap" songs.
  2. You draw a "rap" song in the first draw and a "country" song in the 2nd.
- **(16) Create a histogram that displays the number of times a "rap" song occurred in each simulation. How often were zero rap songs drawn? A single rap song? Two rap songs?**

## On Your Own

- Using what you've learned in the previous two labs, answer the following question by performing two computer simulations with 500 repetitions a piece:  
**(17) If we draw 5 songs from a playlist of 30 rap, 23 country and 47 rock songs, how does the estimated probability of all 5 songs being rap songs change if we draw the songs with or without replacement?**
- For each simulation:
  - Create a histogram for the number of rap songs that occurred for each of the 500 repetitions.
- **(18) Describe how the distribution of the number of rap songs changes depending on if we use replacement or not.**

## **Practicum: Win Win Win**

### **Objective:**

Students will create and combine simulations to assess probabilities.

### **Materials:**

1. *Win Win Win Practicum* (LMR\_U2\_Practicum\_Win\_Win\_Win)

### **Practicum Win Win Win**

The California lottery has a game called the *Daily 3*.

- It consists of 3 numbers between 0 - 9 that are drawn daily.
- The numbers are drawn *with replacement*.
- Winners are usually awarded a couple hundred dollars.
- To win the maximum amount of money, players must correctly choose the numbers that are drawn, in order.

Based on what you learned in *Lab 2C* and *Lab 2D* (*Which song plays next?* and *Queue it up!*) and using the rules of the *Daily 3*, you need to:

1. Write down the code to correctly simulate the *Daily 3* once.
2. Use your code to simulate the *Daily 3* 500 times.
3. Compute the estimated probability of getting the first 2 numbers of the *Daily 3* correct.
4. Should the estimated probability of correctly guessing the last 2 numbers of the *Daily 3* be less than, the same as, or more than guessing the first 2 numbers? Why?
5. In teams of 4:
  - a. Each team member chooses 3 numbers for the *Daily 3*.
  - b. Each team member simulates the *Daily 3* game 500 times.
  - c. Within your group, combine the team simulation estimates to estimate the probability of winning the *Daily 3*.
6. Write and submit a one-page report. Your report should include the code.

# What Are the Chances That You Are Stressing or Chilling?

Instructional Days: 8

## Enduring Understandings

Permutations of data provide a model that shows us how the world behaves if chance is the only reason for differences between groups or for associations between variables. If our actual observation is a rare permutation, this suggests that chance is not a good explanation for the difference or association. On the other hand, if the actual observation is a common permutation, this suggests that chance may be a valid explanation. Differences between permuted data and actual data suggest that the chance model can be rejected and there is a dependent relationship between two variables.

## Engagement

Students will read the Huffington Post article titled *Don't Take My Stress Away* to set the stage for the Stress/Chill Campaign. High school students who expected, and wanted, to feel stressed out by school wrote this article. The article is found at:

[http://www.huffingtonpostdont/jack-cahn/dont-take-my-stress-away\\_b\\_2090203.html](http://www.huffingtonpostdont/jack-cahn/dont-take-my-stress-away_b_2090203.html)

## Learning Objectives

### Statistical/Mathematical:

S-IC 2: Decide if a specified model is consistent with results from a given data-generating process, e.g., using simulation.

S-IC 6: Evaluate reports based on data.

### Data Science:

Understand that a chance model serves as an indicator of whether or not associations in the actual data are due to chance (understand why a plot might appear to have a trend but may actually be the result of randomness). Understand that simulations provide a way of comparing expected chance outcomes to real outcomes in order to determine if a model and actual data appear consistent. Learn about merging datasets by understanding the structure of both datasets and the logic of the way they will be combined.

### Applied Computational Thinking using RStudio:

- Permutations of data, determining if actual data is similar to permuted data.
- Merge multiple datasets together based on a common variable.
- Create permutations using a merged dataset.

### Real-World Connections:

In media, citizens read about results and scientific studies in which treatments are applied. In real life, one can ask the question: Does this happen by chance? Understanding chance helps us interpret media reports of scientific and medical findings.

## Language Objectives

1. Students will engage in partner and whole group discussions to express their understanding of chance occurrences and their probabilities.
2. Students will use complex sentences to write informative short reports that use data science concepts and skills.
3. Students will write informative reports that use graphical representations and numerical summaries.

## Data File or Data Collection Method

### Data Files:

1. Students' *Personality Color* survey: data(colors)
2. Students' *Stress/Chill* campaign data
3. Titanic: data(titanic)
4. Horror Movie: data(slasher)

### Data Collection Method:

**Stress/Chill Participatory Sensing Campaign:** Students will monitor how they feel at different times of the day – whether they are “stressing” or “chilling”. Along with how they feel, they will make observations regarding other factors, such as being alone or with others, what they are doing at that moment, and why they are doing that activity.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 12: Don't Take My Stress Away!

### Objective:

Students will read the Huffington Post article titled *Don't Take My Stress Away* to spark their interest about how they spend their time and will continue to read reports critically to look for claims that may or may not be based on data.

### Materials:

1. Article: *Huffington Post's Don't Take My Stress Away* found at: [http://www.huffingtondont.com/jack-cahn/dont-take-my-stress-away\\_b\\_2090203.html](http://www.huffingtondont.com/jack-cahn/dont-take-my-stress-away_b_2090203.html)
2. Data collection devices
3. *Stress/Chill Campaign* guidelines (LMR\_U2\_Campaign\_StressChill)

**Essential Concepts:** Generating statistical questions is the first step in a Participatory Sensing campaign. Research and observations help create applicable campaign questions.

### Lesson:

1. Become familiar with the *Stress/Chill Campaign Guidelines* (shown at the end of this lesson), particularly the questions, to help guide students during the campaign (see Campaign Guidelines in Teacher Resources).



2. Ask students the following questions and conduct a brief share out of their responses.
  - a. Do you know anyone who seems to be always stressed or anyone who seems to be always chilled? *Answers will vary by class.*
  - b. What are some observations you have made that make that person extremely stressed or chilled? *Answers will vary by class.*



3. Inform students that they will be learning about some high school students who view stress as a part of life in the Huffington Post article titled *Don't Take My Stress Away*.

4. Provide students the link to the article and allow time for them to read it: [http://www.huffindontpost.com/jack-cahn/dont-take-my-stress-away\\_b\\_2090203.html](http://www.huffindontpost.com/jack-cahn/dont-take-my-stress-away_b_2090203.html)

5. As they read the article, students should note whether they agree or disagree with the authors and should write down their comments and/or reactions to the article in their DS journals.



6. Ask student pairs to share if they agree or disagree with the authors of the article and why. Conduct a *Share Out* of student responses.
7. Inform students that for this unit, we will be investigating how stressed or chilled they are at certain times of the day.
8. Students will collect data using the *Stress/Chill* Participatory Sensing campaign. They will add the *Stress/Chill* campaign to their list of available campaigns either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org>.
9. Ask students to complete their first survey.
10. After students have completed their first survey, use a random number generator to generate two random times a day for the next 6 days (RStudio example given below). It is recommended that you create 6 sets of random numbers so that students are polled at different times each day. Example for RStudio (assuming students are awake between the hours of 7:00am and 11:00pm):

```
> sample(7:23, size = 2, replace = FALSE)
```

**Note:** If a time falls within the school day, it is up to the discretion of the teacher to use this time or not.

11. Based on the times generated, ask students to set reminders for the next 6 days. Students without a mobile device may set reminders using a method available to them.

12. Focus students' attention on the *Stress/Chill* survey questions – you may display the questions on the Campaign Guidelines document (LMR\_U2\_Campaign\_StressChill). Ask students to generate three statistical questions that could be answered using the *Stress/Chill* data.
13. Then, ask them to write down in their DS journals some predictions about what they think they will see after they collect some data.

#### **Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### **Homework**

For the next 6 days, students will collect data for the *Stress/Chill* campaign either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org>.

## Campaign Guidelines – Stress/Chill

### 1. The Issue:

People report being more and more stressed every day. This trend is extending beyond adults, it is also reported by children and teenagers. The amount of work for which people are responsible has been increasing. To understand what makes us feel stressed, some important questions to ask are:

- 1) What factors affect my stress/chill level?
- 2) Do different personality types have different things that make them happy/sad?
- 3) Do you like to be alone or with people?
- 4) Is your stress/chill level a function of the environment in which you are in?

### 2. Objectives:

Upon completing this campaign, students will have compared groups and gained an understanding of variability within and between groups. They will have learned how to conduct and use permutations to model variability, perform informal inference, and how to do simulations to make predictions.

### 3. Survey Questions:

Use a random number generator to generate two random times a day for the next 6 days, including a weekend if possible. It is recommended that you create 6 sets of random numbers so that students are polled at different times each day. If a time falls within the school day, it is up to the discretion of the teacher to use this time or not.

| Prompt                                                                                                       | Variable      | Data Type   |
|--------------------------------------------------------------------------------------------------------------|---------------|-------------|
| Where are you?<br>-school<br>-work<br>-home<br>-public place<br>-others' house<br>-commuting                 | where         | categorical |
| Why are you here (in one word)?                                                                              | why           | text        |
| How many people are you with (not counting yourself)?                                                        | howmanypeople | integer     |
| Who are you with?<br>-alone<br>-friends and/or family<br>-teacher and classmates<br>-strangers<br>-coworkers | who           | categorical |
| How stressed are you feeling right now (3 is very stressed, 0 is not stressed at all)?                       | stress        | integer     |
| AUTOMATIC                                                                                                    | location      | lat, long   |
| AUTOMATIC                                                                                                    | time          | time        |
| AUTOMATIC                                                                                                    | date          | date        |

**When?** Surveys are taken two to three times per day at pre-determined randomly selected times.

**How Long?** About two weeks. Ideally, two of these days include a weekend.

#### **4. Motivation:**

Students must understand that they need to keep collecting data. Use the Plot App to look at the data after the first day and have a discussion.

Ask: Why were most people stressed? Guide students along the way.

Ask students to predict the following: What is your stress/chill level in the evening versus morning? Does it change every day? How about during the weekend? What is the difference between groups?

Data collection: After the first day, use the Campaign Monitoring tool to see who has collected the most data.

#### **5. Technical Analysis:**

Students will use RStudio.

#### **6. Guiding Questions:**

- 1) Have students generate predictions and check up on their predictions.
- 2) What's the typical stress/chill level of the class across the campaign?
- 3) What's my typical stress/chill level and how does it compare to the whole class?
- 4) Do the stress/chill levels vary by weekday or weekend or the type of people you are with?
- 5) Under which conditions is my stress/chill level affected?
- 6) Encourage students to generate their own questions.

#### **7. Report:**

Students will complete the Stress/Chill Practicum. They will analyze their stress/chill data using data analysis skills and RStudio skills learned in the unit.

### Lesson 13: The Horror Movie Shuffle

#### Objective:

Students will understand that, just by chance, we will see differences between two groups. They will understand that these differences are usually small. Specifically, they will learn that we can determine if outcomes are due to chance for categorical variables by calculating differences in the proportions between two groups.

#### Materials:

1. 3" x 5" cards (1 per student)

#### Vocabulary:

chance, simulations, randomness, shuffle

**Essential Concepts:** We can “shuffle” data based on categorical variables. The statistic we use is the difference in proportions. The distribution we form by shuffling represents what happens if chance were the only factor at play. If the actual observed difference in proportions is near the center of this shuffling distribution, then we could conclude that chance is a good explanation for the difference. But if it is extreme (in the tails or off the charts), then we should conclude that chance is NOT to blame. Sometimes, the apparent differences between groups is caused by chance.

#### Lesson:

1. **Data Collection Monitoring:** Display the IDS Campaign Monitoring Tool, found at <https://portal.thinkdataed.org>. Click on **Campaign Monitor** and sign in.
2. Inform students that you will be monitoring their data collection. Ask:
  - a. Who has collected the most data so far? *See User List and sort by Total.*
  - b. How many active users are there? How many inactive users are there? *Click on pie chart.*
  - c. How many responses were submitted yesterday and today? *See Total Responses.*
  - d. How many responses have been shared? How many remain private? *Click on pie chart.*
  - e. Using TPS, ask students to think about what they can do to increase their data collection.
3. Conduct a discussion about the data that has been collected.
4. Have students recall what they have learned about **chance** (see Lesson 8). *Synonyms: possibility, prospect, expectation, unintentional, unplanned. The actual definition of chance is “a possibility of something happening.”*
5. Explain that we can use **simulations** to show that sometimes, when we think two groups are different, the difference is really just because of chance, or **randomness**, and does not mean anything.
6. Remind students that a simulation is a model for creating random outcomes. Randomness means that something just happens without a specific order.
7. In pairs, ask students to name situations where two groups could be compared, and then have the students record these situations in their DS journals. Some examples include:
  - *Men earn more money than women for some work.*
  - *Basketball players are faster runners than baseball players.*
  - *Los Angeles students are smarter than \_\_\_\_\_.*
  - *UCLA football players are better athletes than USC football players.*
  - *You and a friend flipped a coin 10 times, and you got more “heads”.*



8. Then, ask students to write next to each situation whether they think the differences are either real or due to chance. Sometimes differences between two groups are real, but sometimes they might just be due to chance. They will be learning ways to tell the difference. *Answers will vary but the purpose here is to have a discussion about student perceptions.*



9. Explain to the class that we are interested in finding out who will survive by the end of a horror movie. Ask the students:

- a. Do you think men and women have an equal likelihood of surviving by the end of a horror movie? *Answers will vary by class.*

10. Have a few students share out their opinions along with their reasoning.

11. Inform the students that they will be pretending to be actors from horror movies during today's lesson.

12. Explain that data from horror movies (sometimes called slasher films) were collected of 485 characters from 50 films. For each character, 2 variables were recorded: Gender and Survival. The values for Gender were "Female" and "Male". The values for Survival were "Dies" and "Survives".

**Note:** The source for this data uses gender in a binary manner ("Female" and "Male"), presumably because the actors are playing characters in each movie.

|          | Gender |      |
|----------|--------|------|
| Survival | Female | Male |
| Dies     | 172    | 228  |
| Survives | 50     | 35   |
| Total    | 222    | 263  |

Notice that there were more male characters than female characters and that most characters in slasher films do not survive.

13. From this data, the proportion of survivors was calculated for each gender. In other words, for all female characters, the number of female survivors was divided by the total number of females. Similarly, for all male characters, the number of male survivors was divided by the total number of males.

$$\frac{\#(\text{"Female"} \ \& \ \text{"Survives"})}{\#(\text{"Female"})} \quad \text{or} \quad \frac{\#(\text{"Male"} \ \& \ \text{"Survives"})}{\#(\text{"Male"})}$$



14. The percent of females who survived by the end of a horror movie was about **23%**, and the percent of males who survived by the end of a horror movie was about **13%**. Ask the students:

- a. Is this what you expected (refer back to the discussion from step 9 above)? *Answers will vary by class. If students thought males would survive more often, then these results would be unexpected. If students thought females would survive more often, then these results would be exactly what they expected. If students thought there was an equal likelihood of survival, these results would also be surprising.*
- b. What is the difference in the proportions of survival rates between genders? What does this mean in the context of surviving a horror movie? *Answer: The difference is 23% - 13% = 10%. This means that 10% more women characters survived than men.*
- c. Is this difference "big" or "small"? How can they define what is "big" and what is "small"? *Answers will vary by class. Upon first glance, it may seem like 10% is a big difference, but we do not know for sure.*

15. Explain that they will participate in an activity to determine if the 10% difference seen in the actual dataset is big or small. This will help them determine if there really is a difference in survival rates for males versus females, or if the 10% difference was just due to chance.

16. Split the class into two groups, 46% of them on one side of the room and the other 54% on the other side of the room. Tell the smaller group they have been assigned to play female characters in the horror movie (regardless of gender) and tell the larger group that they have been assigned to play male characters in the horror movie (regardless of their gender). Once those groups have been created, have the class calculate the number of students in each group that would have survived a horror film using the actual proportions given in step 14 above.

For example: For a class of 30 students:

- 46% of 30 ( $0.46 \times 30$ )  $\approx$  14 students representing female characters.
- Of those 14 female characters, 23%, or 3 ( $0.23 \times 14 \approx 3$ ), are survivors.
- The remaining 16 students ( $30 - 14 = 16$ ) represent male characters.
- Of those 16 male characters, 13%, or 2 ( $0.13 \times 16 \approx 2$ ), are survivors.

17. Each group should then decide which students will be survivors. Using the 3" x 5" cards, students should write either "dies" or "survives" on their card.

For example (continued from above):

*Three of the females are survivors; so 3 female characters from the group should write "survives" on their cards. The rest of the group should write "dies" on their cards.*

*Two of the male characters are survivors; so 2 males from the group should write "survives" on their cards. The rest of the group should write "dies" on their card.*



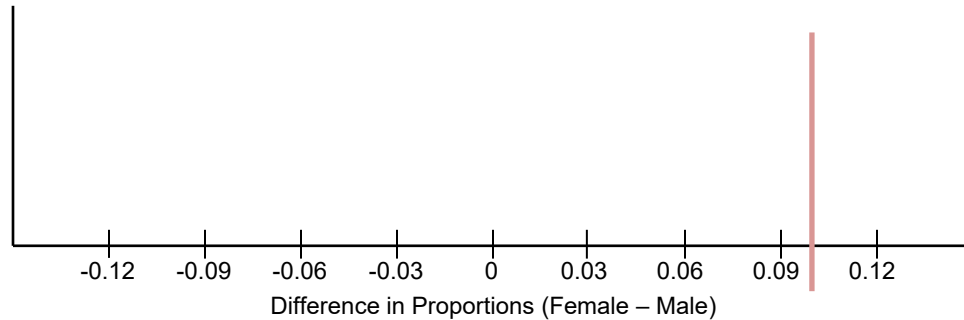
18. Explain to students that IF there really is no difference between genders in horror films, then the characters who survived would only have done so by chance. In other words, males and females would have an equal likelihood of surviving. Have students discuss the following questions:
- a. How many total people in our class are survivors? What is the total proportion of survivors? *Answers will vary by class. Using the example above, there would be a total of 5 survivors from the class of 30 students. The proportion of survivors would be  $5/30 \approx 0.17 = 17\%$ .*
  - b. How many of the survivors would we expect to be male? How many would we expect to be female? *Answers will vary by class. Using the example above, we would expect to see 17% of males and 17% of females survive since that was the overall proportion of survivors. So, we would expect  $0.17 \times 16 \approx 3$  male survivors, and  $0.17 \times 14 \approx 2$  female survivors.*
19. Collect all of the 3" x 5" cards from the students and explain that you are going to **shuffle** the cards and redistribute them so that their genders have no influence on whether or not they survive the horror movie.
20. Visibly shuffle the survives/dies cards to create a random shuffle. Once the cards have been well-shuffled, pass them back out to the students face down. After all the cards are given out, each group should identify the number of people that are survivors and calculate the corresponding proportion of the survivors.
21. On the board, create a table to display the proportions of survivors for each gender, and include a column for the difference (female survivors – male survivors). Fill in the table with the values the students found in step 20 above.

**Note:** The first row has been filled in with the example data from above BEFORE the shuffles have taken place. Exact numbers were not used so that the proportions would match the actual horror movie dataset.

| # of Female Survivors | # of Male Survivors | Proportion of Female Survivors | Proportion of Male Survivors | Difference in Proportions (Female – Male) |
|-----------------------|---------------------|--------------------------------|------------------------------|-------------------------------------------|
| 3.22                  | 2.08                | $3.22/14 = 0.23$               | $2.08/16 = 0.13$             | $0.23 - 0.13 = 0.10$                      |
|                       |                     |                                |                              |                                           |

22. Note that values in the "Difference in Proportions" column can be positive or negative because sometimes more women will survive, and other times more men will survive.

23. Draw a dotplot on the board labeled “Difference in Proportions.” Include a vertical line at 0.10 (10%) to represent the actual difference in gender survival rates in real horror movies (see example below).



24. Using the information from steps 20 and 21 above, place a dot at the corresponding value for the shuffled data’s difference in proportions. Ask the students:
- How does this difference compare to the actual dataset’s difference of 10%? *Answers will vary by class. Most likely, the difference in proportions will be much smaller than 10%. In fact, the difference in proportions will be centered around 0.*
25. Repeat steps 19-24 from above a few more times (depending on how much class time you have available).
26. Ask the students to record their responses to the following questions:
- What was the biggest difference we saw from our shuffles? What was the smallest? *Answers will vary by class.*
  - What do you think this dotplot would look like if we shuffled our survival cards 1000 times? *Answer: The dotplot would look roughly symmetric and centered around 0, meaning that if there were no relationship between a character’s gender and whether or not they survive, the difference in proportions would typically be 0.*
27. Have a discussion about how the actual difference in gender survival (10%) is rarely seen when we assign “survives” or “dies” just by chance (aka when shuffling). What does this mean in terms of who will die in actual horror movies? *Answer: Since we never (or rarely) saw a 10% difference in the proportions of female survivors versus male survivors, it seems that horror movies actually favor female survivors.*
28. Ask the students:
- If you were going to be cast in a horror movie, would you want to be a male character or a female character? *Answer: You would want to be a female character because they are more likely to survive by the end of the film.*
29. Inform the students that they will learn how to shuffle in RStudio in order to determine if an event is real or simply due to chance.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework & Next Day

For the next 5 days, students will collect data for the *Stress/Chill* campaign either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org>

### Lab 2E: The Horror Movie Shuffle

Complete Lab 2E prior to Lesson 14.

## Lab 2E: The Horror Movie Shuffle

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Playing with permutations

- There is a common belief that women in *slasher* films are more likely to survive than men.
- This lab will focus on the statistical investigative question: *Are women in slasher films more likely to survive until the end of the film than men?*
- To answer this question, we'll learn how to use permuted data to gauge how likely an event occurs by chance.
- **(1) To begin, write and run code using the `data` function to load the `slasher` data file.**
  - The data contains information about 485 characters from a random sample of 50 *slasher* horror films.

### Initial thoughts...

- To familiarize yourself with the data, answer the following:
  - **(2) How many variables and observations are contained in the data and what are the possible values of the variables?**
  - **(3) Which gender had more survivors? Write down a few sentences as to how you came to your conclusion. Be sure to look at both the *counts* and *percentages* of survivors in each group before deciding.**
  - **(4) Calculate the difference between the percentage of females who survived and the percentage of males who survived. Is the difference large enough to conclude that women tend to survive more often than men?**

### Tally whoa ... !

- Something you might have noticed is that these two lines of code aren't equivalent:

```
tally(~gender | survival, data = slasher, margin = TRUE)
tally(~survival | gender, data = slasher, margin = TRUE)
```
- The first line of code takes the group of *survivors* and tells us how many of them were `Male` or `Female`.
- The other takes the group of *females / males* and tells us how many of them `Dies` or `Survives`.
- **(5) The last question on the previous slide can be answered using the line of code below. Why?**
  - Pro-tip: Include the option `format = "percent"` to obtain a two-way table with percentages.

```
tally(~survival | gender, format = "percent", data = slasher, margin = TRUE)
gender
survival Female Male
Dies 77.47748 86.69202
Survives 22.52252 13.30798
Total 100.00000 100.00000
```

## Examining differences

- When we're comparing the difference between two quantities, such as survival rates of slasher films, it can be difficult to decide how *different* two values need to be before we can conclude that the difference didn't just happen by chance.
  - To help us decide when a difference is not due to chance, we'll use repeated random shuffling.
- By using repeated random shuffling, we'll estimate how often our *actual* difference occurs by *chance*.

## Do the shuffle!

- When we shuffle data, we use our original dataset as a starting point.
  - (6) Run the following and write down the resulting table on a piece of paper.

```
tally(~survival | gender, data = slasher)
```

- (7) Now run the following to randomly reassign each `survival` status to each observation. Compare the resulting table to the one you wrote down.

```
tally(~shuffle(survival) | gender, data = slasher)
```

## Let's compare ...

- (8) How many people, in total, survived the slasher film before shuffling? How many people survived after shuffling?
- (9) How has shuffling our data changed the percentage of women who survived compared to men who survived?
  - (10) Is the difference in percentages from your shuffled data larger or smaller than the difference from the original data? Interpret what this means.
- (11) Explain why shuffling our data one time is not enough to decide if the difference seen in our *actual* data occurs by chance or not.

## Detecting differences

- To help us decide if the difference in percentages in our *actual* data occurs by chance or not, we can use the `do()` function to shuffle our data many times and see how often our *actual* difference occurred by chance.
- Run the following lines of code:

```
set.seed(7)
shuffled_outcomes <- do(10) * tally(~shuffle(survival) | gender,
 format = "percent",
 data = slasher)
View(shuffled_outcomes)
```

- (12) In how many simulations did a higher percentage of males survive than females?
- (13) What is the largest difference in percentages of survival between males and females?
- (14) What patterns are emerging from these simulations?
- Ten simulations is not enough. Use the code above and perform 500 shuffles. Assign your 500 shuffles the name `shuffled_outcomes`. Use `set.seed(1)`.

## Now what?

- The next step to find out how often our *actual* difference occurs by chance is to compare it to the differences in our shuffled data.
- To compute the differences for each shuffle we can use the `mutate` function.
  - (15) Fill in the blanks to add a new column that contains the difference between `Survives.Female` and `Survives.Male` to our `shuffled_outcomes` data.

```
shuffled_outcomes <- mutate(shuffled_outcomes, diff = ____ - ____)
```

## Time to decide

- (16) Write and run code creating a histogram of the differences in our `shuffled_outcomes` data. Based on your plot, answer the following:
  - (17) What was the typical difference in percentages between men and women survivors?
- Include a vertical line in your histogram of the actual difference by running the code below:

```
add_vline(vline = 22.52252 - 13.30798)
```
- (18) Does the actual difference occur very often by chance alone?
- (19) Does `gender` play a role in whether or not a character will survive in a slasher film? Explain your reasoning.
- (20) If you wanted to survive in a slasher film, would you want to play a female character or a male character?

## Summary

- By shuffling the `survival` label, we made the proportion of males and females who survived the slasher film random.
  - The males and females survived by chance alone.
- If surviving the film occurred purely by chance, most of the time the difference in survival proportions would be close to zero.
  - Notice how most values in the histogram occur close to zero.
- When we look to see how often our actual difference occurs in our shuffled data, if the actual difference doesn't occur very often then perhaps there is something more going on than just chance alone ...

## On your own

- Carry out another 500 simulations but this time shuffle the `gender` variable instead of the `survival` variable.
  - Include the code `set.seed(1)` before your 500 simulations to make your answer reproducible.
- (21) Does shuffling the `gender` variable instead of the `survival` variable change your answer to the question?
- (22) Does `survival` play a role in a character's `gender`? Why or why not?

## Lesson 14: The Titanic Shuffle

### Objective:

Students will continue to understand that, just by chance, we will see differences between two groups. They will understand that these differences are usually small.

### Materials:

1. *LMR\_Titanic\_Strips*

**Advance preparation required:** Cut strips based on number of students (see step 8 in lesson)

2. Poster paper

**Note:** It may be helpful to have the x-axis scale pre-drawn on the posters (-£30 to £30)

3. Markers

**Essential Concepts:** We can also “shuffle” data based on numerical variables. The statistic we use is the difference in medians. The distribution we form by this form of shuffling still represents what happens if chance were the only factor at play. When differences are small, we suspect that they might be due to chance. When differences are big, we suspect they might be ‘real’.

### Lesson:

- 
1. Remind students that they previously learned how to determine if a difference is due to chance by shuffling based on categorical variables (gender and survival).
  2. Display the dotplot created during Lesson 13 of the difference in proportions between female and male survivors of horror movies. Remind the students that, “by chance,” the differences were typically zero. Most of the time, they were pretty small. Sometimes they were bigger, but that was rare and this tells us that if we see “small” differences, we might think they are due to chance. But if we see “big” differences, they are not.
  3. Lead a short discussion about what students think small and big differences mean. Make sure they answer in units (which are percentage points for the horror movie data). So, for example, a “big” difference might be 5 percentage points (but don’t let them just say “5”).
  4. Inform students that, during today’s lesson, they will learn how to determine if there is a difference between groups when a numerical variable is involved.
  5. In particular, they will assume the roles of passengers in the *Titanic* for today’s lesson. In case some students may not know about the *Titanic*, ask a volunteer to share what he/she knows.
  6. Explain that, at its time, the *Titanic* was the largest cruise ship ever built and was declared to be unsinkable. However, on its first voyage, it sank and was one of the worst maritime disasters in history. About 40% of passengers survived; however, your chances of survival depended very much on your age, sex, and wealth.
  7. Inform the students that we are going to look at whether the amount of money a passenger paid for his/her cabin (the fare price) had anything to do with whether or not he/she survived.
  8. Each student will need a strip from the *LMR\_Titanic\_Strips* file—see below for instructions.

#### **Advance preparation required:**

The *Titanic Strips LMR* contains data from 40 actual passengers on the *titanic*. Each strip represents the data from one passenger: the left-hand side shows the fare paid and right-hand side contains the survival information of that passenger after the collision. Cut the LMR into strips such that the fare price is attached to the survivor status for each of the 40 observations.

| Titanic Strips |          |
|----------------|----------|
| £7.75          | Survivor |
| £26.00         | Survivor |
| £56.93         | Survivor |
| £7.75          | Survivor |
| £80.00         | Survivor |
| £26.55         | Survivor |

LMR\_Titanic\_Strips

9. 40 strips were created for large classes. If your class has less than 40 students, assign the students to two groups such that roughly 40% of them are in the survivor group ( $15/40 = 37.5\% \approx 40\%$ ), and the rest are in the victim group. If your class is small (smaller than 10), then put the students in two equal sized groups. The split does not have to be exactly 40%, as the actual value from our data is 38.8%.
10. Inform the smaller group that they are the survivors and distribute a survivor strip with its corresponding fare to each student. Set aside any leftover strips. Tell them that the price on the strip represents the amount of money paid for their ticket to board the *Titanic*. Notify them that £20 in 1912 is worth about £2952 today.
11. Divulge to the larger group that they, unfortunately, are the victims and distribute a victim strip with its corresponding fare to each student. Set aside any leftover strips.
12. Ask each group to create a boxplot of their fare prices on a poster. Lead a quick discussion comparing the two boxplots visually. Then, ask each group to calculate the median fare for their group.
13. As a class, find the difference between the median fares for the two groups.

median of "Survivor" fares – median of "Victim" fares

For example:

*If all 15 survivor cards and all 25 victim cards are used, the difference in medians would be:*

$$£26.00 - £13.00 = £13.00$$

14. Explain that one of the controversies of the *Titanic* disaster was that some people felt that the rich people were given better access to the lifeboats than were the poor, so rich people were more likely to survive. Note that the data represented on the fare cards are only a subset of the actual *Titanic* data, which had over 800 passengers. However, the data were randomly selected from the real data and are considered representative of the 800 passengers.



15. In pairs, ask students to discuss the following:
  - a. Based on the data from our dotplots, do you think rich people were more likely to survive? In other words, did passengers who paid more for their tickets have a better chance of survival? *Answer: Yes, there is evidence that rich passengers survived more often than poorer passengers. The median difference between the fare prices of the survivors and the victims is 13.00 (see step 13 above). Most survivors had higher fare prices than the victims, so the distribution of survivor fares is shifted to the right and is more right-skewed.*
16. Share out a couple of responses with the whole class.

17. Have students tear their strip such that they separate the fare from the outcome (survivor or victim). Collect only the outcomes and randomly shuffle them. Students will keep their fare. Distribute the shuffled outcome strips face down to the students. Once everyone has a new outcome strip, ask students to turn their outcome strip over and re-group based on their new survival status.



18. Ask the students:

- Why do we shuffle the survivor/victim strips and not the fare strips? *Answer: We want to know if the price someone paid for his/her ticket affects whether or not he/she survived. So, when we shuffle, we assume that fare price has nothing to do with survival, so the prices should be irrelevant.*
- What do you think the median fare difference of our shuffled groups will be? *Answer: The median fare difference of the shuffled groups should be close to 0, meaning that there should be NO difference in fare price for the survivors and the victims. Everyone would have the same chances of surviving, regardless of their ticket price.*



19. Have each group calculate the median fare price for their new groups. Then, ask:

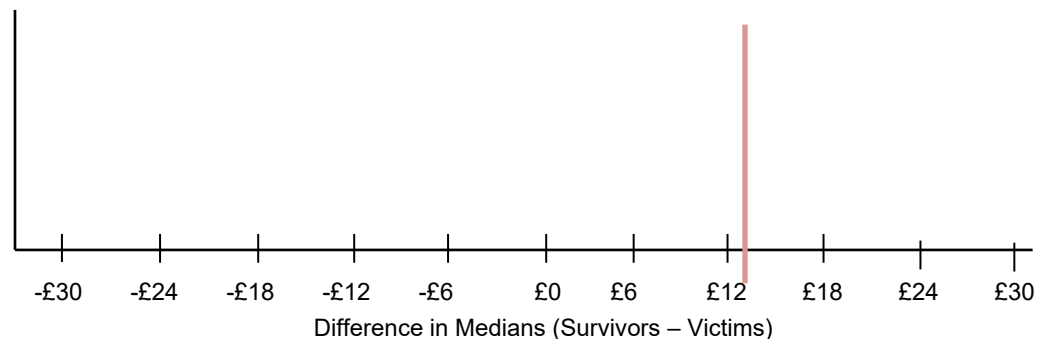
- Do you think this difference, of \_\_\_ dollars, is real or due to chance? *Answers will vary by class. Since the data were shuffled, any difference should be due to chance.*

20. On the board, create a table to display the median fare prices for each group, and include a column for the difference (median “Survivor” fare – median “Victim” fare). Fill in the table with the values the students found in step 13 above.

**Note:** The first row has been filled in with the example data from above BEFORE the shuffles have taken place.

| Median Fare Price of Survivors | Median Fare Price of Victims | Difference in Medians (Survivors – Victims) |
|--------------------------------|------------------------------|---------------------------------------------|
| £26.00                         | £13.00                       | £26.00 - £13.00 = £13.00                    |
|                                |                              |                                             |

21. Note that values in the “Difference in Medians” column can be positive or negative because sometimes the survivors will pay more for their tickets, and other times the victims will pay more for their tickets.
22. Draw a dotplot on the board labeled “Difference in Medians.” Include a vertical line at £13.00 (or whatever value was calculated in step 13 above by the class) to represent the actual difference in the median fare prices between the survivors and the victims (see example below).



23. Using the information from steps 19 and 20 above, place a dot at the corresponding value for the shuffled data’s difference in medians. Ask the students:

- How does this difference compare to the actual difference of £13.00 (from step 13 above)? *Answers will vary by class. Most likely, the difference in medians will be much smaller than £13.00. In fact, the difference in medians will be centered around 0.*

24. Remind students that small differences might be due to chance and big differences typically mean that there is a “real” difference between groups. In this case, a big difference might mean that the rich passengers were more likely to survive. And a small difference might mean that survival was just a matter of plain luck.



25. Repeat steps 17-23 from above a few more times (depending on how much class time you have available).

26. In pairs, ask students to discuss whether they think the real difference in median fare prices they calculated in step 13 above (£13.00 if all cards were used) is small or large. *Answers will vary by class. Guide students to look at the MAD value of the distribution of differences in median fares.*

27. Explain that one way that we can decide what is “large” or “small” is by creating cut-off values that we think are too far away from the center of the distributions of differences. In general, we can assign a rule that states that any difference in median fare prices that is greater than 2 MAD values above or below the median is considered unusual. This means that any value in the outer edges of the plot would indicate that a passenger’s ticket price impacted his/her chances of survival.

28. Inform students that they will use RStudio to shuffle the actual *Titanic* data of all 800 passengers during the next class and can decide if the difference in survival rates of rich passengers and poor passengers was real, or just due to chance.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework & Next Day

For the next 3 days, students will collect data for the *Stress/Chill* campaign either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org>

### **Lab 2F: The Titanic Shuffle**

Complete Lab 2F prior to Lesson 15.

## Lab 2F: The Titanic Shuffle

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Previously ...

- In the previous lab, we learned that by using a **do**-loop and the **shuffle** function, we could simulate randomly shuffling our data many times.
  - This helps us determine how likely it is that a difference between groups is due to chance.
- For this lab, we will extend these ideas to *numerical* variables by using random shuffling and numerical summaries.
- The question we will investigate in this lab is:

*Is there any evidence to suggest that those who survived paid a higher fare than those who died?*

### The Titanic

- The Titanic was a ship that sank en route to the U.S.A. from England after hitting an Iceberg in 1912.
  - At the time, it was claimed that the Titanic was *unsinkable* ... it wasn't ... because it did.
- **(1) Write and run code using the **data** function to load the **titanic** passenger and survival data.**
- **(2) Write and run code creating a boxplot of the **fares** paid by passengers and facet the plot based on whether the passenger survived or not.**
  - **(3) Based on the plot, do you believe that passengers who paid a higher fare on the Titanic were more likely to survive? Explain why and describe how certain you are of being correct.**

### The search begins!

- **(4) Write and run code to visualize the distribution of fares paid.**
- **(5) Which numerical summary might be preferred to describe the *typical* value?**
- **(6) What was the *typical* fare paid by survivors? Non-survivors? How much more did the typical survivor pay?**

### Do the shuffle!

- **(7) Write and run code using the **do** and the **shuffle** functions to **shuffle** the passenger's survival status 500 times.**
  - Use the previous lab if you need some help on how to do this.
  - For each shuffle, compute each group's **median** fare paid.
  - Assign your shuffled data the name **shuffled\_survival**.
- **(8) After shuffling your data, write and run code using the **mutate** function to create a variable called **diff** which is the **median** fare of survivors minus the **median** fare of non-survivors.**
  - Assign your mutated data the name **shuffled\_survival** again.

### Put your simulations to use

- **(9) Using your shuffled data, answer the research question we posed at the beginning of the lab.**  
*Is there any evidence to suggest that those who survived paid a higher fare than those who died?*

- (10) Write up your answer as a statistical analysis. Create a plot and explain how the plot supports your conclusion. Be sure to also explain why shuffling your data is important.

### Comparing Mean Fares

- What about if instead of calculating the median fare price for each group after a shuffle, we calculated the mean fare price and took the difference (mean\_survivor – mean\_victim)?
- (11) If we did this 500 times, what do you predict the distribution of differences will look like?
- (12) Write and run code using the `do` and the `shuffle` functions to shuffle the passenger survival status 500 times.
  - For each shuffle, compute each group's mean fare paid.
  - After shuffling your data, use the `mutate` function to create a variable called `diff` which is the mean fare of survivors minus the mean fare of non-survivors.
- (13) What does the shuffled data reveal? Does the answer to the research question below change when using the mean fares instead of the median fares?

*Is there evidence to suggest that those who survived paid a higher fare than those who died?*

## Lesson 15: Tangible Data Merging

### Objective:

Students will learn how to merge two datasets and ask statistical questions about the merged data.

### Materials:

1. *Tangible Data Merging* file (LMR\_U2\_L15)  
**Advance preparation required:** Needs to be cut into strips (see step 5 of lesson)
2. Copy paper in two colors  
**Advance preparation required:** To distinguish between Dataset 1 & 2 (see step 5 of lesson)

### Vocabulary:

merge

**Essential Concepts:** We can enhance the context of a statistical problem by merging related datasets together. To merge data, each dataset must have a “unique identifier” that tells us how to match up the lines of the data.

### Lesson:

1. Inform students that they are going to examine the research question “Does the personality color test really work?” To answer this, we’re going to examine whether the different color groups actually differ on particular beliefs or attitudes, or if these differences might just be due to chance. In particular, we are going to use the *Stress/Chill* data to see if there is evidence that the “colors” actually differ.
2. Show students the variables in each of these datasets. Give students time to brainstorm statistical questions of interest with their teams and record their questions in their DS journals. Encourage them to think of two- and three-variable questions.
3. Conduct a share out of some of the questions students came up with. *Examples include: (1) Do people whose predominant color is Gold tend to stress more than people whose predominant color is Blue? (2) Is there a difference between the sorts of things that stress out the different personality colors?*
4. In order to answer the above questions, we will need to merge our class’s 2 datasets together (*Personality Color* and *Stress/Chill*). In order to do this, we will be practicing how to merge datasets today.
5. Print out the material from the *Tangible Data Merging* file (LMR\_U2\_L15). Use a different color of paper for each of the two datasets. For example, Dataset 1 could be on plain white paper and Dataset 2 could be on blue paper. Cut the paper by creating horizontal strips of each observation of data. For example, from the first page of Dataset 1, you would create 12 different strips of paper, one for each observation.

**Note:** The screenshot below shows the first 5 observations but there are 12 on the first page.

Tangible Data Merging

Instructions for the teacher:

- a. Use a different color of paper for each of the two datasets. For example, Dataset 1 (pg. 1 & 2) could be on plain white paper, and Dataset 2 (pg. 3 & 4) could be on blue paper.
- b. Cut the paper by creating horizontal strips of each observation of data. For example, for page 1, you would create 12 different strips of paper, one for each observation.

Dataset 1

| Birth Month | Zip Code | Age | ID Number | Favorite Movie         |
|-------------|----------|-----|-----------|------------------------|
| January     | 90064    | 21  | 1742      | The Notebook           |
| Birth Month | Zip Code | Age | ID Number | Favorite Movie         |
| August      | 90291    | 26  | 1936      | Dumb & Dumber          |
| Birth Month | Zip Code | Age | ID Number | Favorite Movie         |
| August      | 90004    | 26  | 1324      | Moulin Rouge           |
| Birth Month | Zip Code | Age | ID Number | Favorite Movie         |
| February    | 90024    | 19  | 1937      | Hunger Games           |
| Birth Month | Zip Code | Age | ID Number | Favorite Movie         |
| November    | 90291    | 30  | 1434      | Breakfast at Tiffany's |

LMR\_U2\_L15

6. Hand each student in the class a strip of paper. Ask them to try to find someone with the other dataset (i.e., a person with a different colored strip of paper) that they can “match up,” or **merge**, with.
7. For example, a student with the first row of data listed below from Dataset 1 might want to match up with the second row of data listed below from Dataset 2 because a person who is 21 has probably graduated high school.

| Birth Month | Zip Code | Age | ID Number | Favorite Movie |
|-------------|----------|-----|-----------|----------------|
| January     | 90064    | 21  | 1742      | The Notebook   |

| Zip Code | ID Number | Birth Month | Siblings | Education   |
|----------|-----------|-------------|----------|-------------|
| 91331    | 1352      | August      | 2        | High School |

8. However, they should notice that they cannot just make guesses about a person’s characteristics in order to match up the data. They should realize that only 3 of the variables are the same in both datasets: *Birth Month*, *Zip Code*, and *ID Number*.
  - a. Since multiple people have the same *Birth Month*, discuss why this may not be the best variable to merge with. *Answer: Multiple people are born in January, so we would have no way of differentiating between those people.*
  - b. The same is true for the *Zip Codes* variable. Although there are less repeats with *Zip Codes*, we still see some overlap between observations.
  - c. So, the only *UNIQUE* identifier in both datasets is *ID Number*. So the students should end up in pairs at the end of the exercise – a student from Dataset 1 is matched with the student from Dataset 2 that has the same *ID Number*.
9. Have the students write about the experience of tangible data merging in their DS journals and ask:
  - a. Why is it important to have at least one unique identifier for both datasets? *Answer: It is the only way to know which information belongs to which person. We want to make sure we do not match up observations (in this case, people) incorrectly because that will compromise any analysis we do later.*
10. Inform students that they will learn to merge datasets using RStudio during the next lab.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework & Next Day

Students will collect data for one more day for the *Stress/Chill* campaign either through the IDS ThinkData Ed App or via web browser at <https://portal.thinkdataed.org/>

### Lab 2G: Getting It Together

Complete Lab 2G prior to the Practicum.

## Lab 2G: Getting It Together

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Putting data together

- In the labs so far, we've only ever looked at individual data files.
  - But often times, we gain additional insights by including additional information from a separate dataset.
- In this lab, we will learn how to merge information from our *personality color* data with our *stress/chill* data.
- **Export, upload, import your Personality Color dataset and name it colors.**
- **Then, export, upload, import your Stress/Chill dataset and name it stress.**

### Looking at Stress/Chill

- We would like to analyze the research question:  
*How do people's personality colors and/or sports participation affect their stress levels?*
- We already have data about *personality color* and a separate dataset about *stress*.
  - What we don't have is a single dataset with information from both ... yet.
- We'll start then by strategizing how to merge our data together.

### Deciding how to merge

- Before we merge data, we need to decide *how* we plan to merge it:
- We can *stack* our datasets, that is, take one dataset's rows and add them to the bottom of the other dataset.
- We can also *join* our datasets horizontally. This is where we take one dataset's columns and add them to the end of the other dataset's columns based on matching an *ID* variable.
  - The *ID* variable will have entries that we use to *match* observations in both datasets.
- **(1) To answer the research question of interest, would it make more sense to stack or join our colors and stress data?**

### Finding variables in common:

- Look at the **names** of the variables in each dataset.
  - To merge different datasets together, we need to find variables they have in common.
- **(2) Which variables do the datasets have in common?**
- **(3) Which variable would make sense to merge the datasets together with? Why not the others?**

### Caution required

- Whether *stacking* or *joining*, we need to be careful when we merge data:
- When *stacking* data, we need to be absolutely certain that the variables we're stacking represent the exact same measurements.
  - We wouldn't want to stack **height** in meters and **height** in inches, for instance (without converting one to the other).
- When *joining* data, we need to make sure that the *id* variable in our primary dataset matches to *one and only one* observation in the joining data.
  - Otherwise, **R** won't know which observation to match to.

## Getting ready

- Our goal is to add the variables from the `colors` data onto the `stress` data.
- Start by ensuring that every user `.id` in the `colors` data is unique.
  - If there's a duplicate, have your teacher remove the duplicate from your class's Web Response Manager and then re-export, upload, import your `colors` data.
- (4) After we add the data from `colors` to `stress`, how many rows should our merged data have? Write this number down.

## Putting them together

- We can use the `merge` function to join our datasets together using the variables that appear in both sets.
- (5) Fill in the blanks below to join the information from the `colors` data onto the `stress` data.

```
merge(____, ____, by = "____")
```
- Assign this merged dataset the name `stress_colors`.
  - Make sure your data has the same number of observations that you wrote down on the previous slide.

## Saving your data:

- View your merged data and make sure nothing appears to be blatantly wrong with it.
- (6) Why didn't we stack the rows of data instead?
- (7) What happens if you swap the order of the datasets in the `merge` function?
- Use the code below to save our `stress_colors` data for later use.

```
save(stress_colors, file = "stress_colors.Rda")
```
- Be sure to look in the *Files* tab to make sure your data was saved.

## Moving on

- In the next practicum, we'll begin analyzing our merged data. In the meantime:
- (8) Write and run code making a few plots using variables from the `stress` data and *facet* or *group* the plots based on variables from the `colors` data.
- (9) Write down the most interesting discovery you made by just exploring your data. Write out how you found your discovery and interpret what it means for the people in your class.
- (10) With our `colors` data, we could answer questions about the *typical* color scores in your class. Why can we no longer answer this question in our `stress_colors` data?

## **Practicum: What Stresses Us?**

### **Objective:**

Students will use RStudio to make graphical representations or numerical summaries of their *Stress/Chill* and *Personality Color* data to answer research questions.

### **Materials:**

1. *What Stresses Us? Practicum* (LMR\_U2\_Practicum\_What\_Stresses\_Us)

### **Practicum**

#### **What Stresses Us?**

We made a dataset that combined our *Stress/Chill* data with our *Personality Color* data. You will use this data to answer the following research questions:

- Do color personalities really predict a person's personality?
- Do people with different personality colors tend to have different stress levels?

Based on the merged data, you need to:

1. Write a one-page report to address these research questions. Use the Data Cycle. Your analysis should include both numerical methods (means, medians, etc.) and graphical methods (plots). The research questions are fairly broad, and you should first think of simpler statistical questions you could ask that would address these research questions.
2. In your report, be sure to:
  - a. Provide the plot(s) and numerical summary (or summaries).
  - b. Describe what the plot shows.
  - c. Explain why you chose to make that particular plot.
  - d. Explain how the plot and numerical summary answers your statistical question.
3. Present your report to another member of the class who is not in your team.
  - a. Make sure to include any relevant plots or numerical summaries that you use to make *any* conclusion in your analysis.

**Note:** You may use the scoring guide in Unit 1 to give you an idea of how to score the Practicum.

# What's Normal?

Instructional Days: 5

## Enduring Understandings

Students learn that the Normal curve can be used as a model that describes many real phenomena. Drawing plots of the Normal curve over histograms helps data scientists determine if the distribution represented by the histogram is close to Normal. The Normal curve suggests that one is more likely to obtain values that are close to typical (average), which are found in the center of the curve, and less likely to obtain values that are extreme and farther away from typical.

## Engagement

Students will learn about the Normal curve by watching the first 35 seconds the New York Times Video “Bunnies, Dragons, and the Normal World” found at:

<http://www.nytimes.com/video/science/10000002452709/bunnies-dragons-and-the-normal-world.html>.

## Learning Objectives

### Statistical/Mathematical:

S-ID 4: Use the mean and standard deviation of a data set to fit it to a normal distribution and to estimate population percentages. Understand that there are data sets for which such a procedure is not appropriate. Use calculators and RStudio to estimate areas under the normal curve.

S-IC 6: Evaluate reports based on data.

### Data Science:

Learn to eyeball Normal distributions and overlay a Normal curve on a histogram; learn to simulate draws from a Normal distribution, and the impact of sample size; learn that estimating probabilities with a model leads to stable estimates; and estimate probabilities by finding the area under the Normal curve using RStudio.

### Applied Computational Thinking Using RStudio:

- Use software to find the area under a Normal curve
- Use software to compare sample distributions (with histograms, for example) with the Normal distribution and make a decision as to whether the distribution appears Normally distributed.
- Draw random samples from a Normal distribution using software.

### Real-World Connections:

The Normal curve is used to make inferences about a population. The model makes it possible to estimate the probability of occurrence of any value of a Normally distributed variable. For example, heights are Normally distributed. Using a Normal curve, we can find the probability of that a person would be a height of 6' 2".

## Language Objectives

1. Students will use mathematical vocabulary to explain Normal distributions and their attributes.
2. Students will compare and contrast standard deviation with mean absolute deviation (MAD).
3. Students will write informative short reports that use data science concepts and skills.

## Data File or Data Collection Method

### Data Files:

1. CDC: `data(cdc)`
2. Titanic: `data(titanic)`

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 16: What is Normal?

### Objective:

Students will learn what a Normal distribution is and learn how to identify a Normal distribution.

### Materials:

1. Video: *New York Times*' "Bunnies, Dragons, and the Normal World" found at: <http://www.nytimes.com/video/science/100000002452709/bunnies-dragons-and-the-normal-world.html>  
**Note:** Show only the first 41 seconds of the video.
2. Graphics from the *Normal Plots* file (LMR\_U2\_L16)
3. Projector to display plots
4. 3" x 5" cards (1 per student)

### Vocabulary:

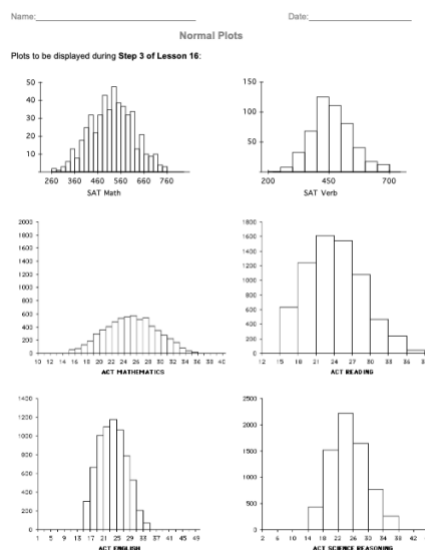
bell-shaped, normal curve, normal distribution

**Essential Concepts:** The Normal curve, also called the Gaussian distribution and the "bell curve," is a model that describes many real-life distributions and is usually called the Normal Model.

### Lesson:

1. Remind students that in Unit 1, Lesson 11 (*What Shape Are You In?*), they sorted histograms into groups based on their shapes.
2. The *Normal Plots* file (LMR\_U2\_L16) contains some of the unimodal **bell-shaped** distributions from the original handout of that lesson (*Sorting Histograms* handout (LMR\_U1\_L11)).

**Note:** You do not need the original handout from Unit 1 – all relevant plots have been compiled in the *Normal Plots* file (LMR\_U2\_L16) for accessibility. Six plots are included: SAT Math, SAT Verb, ACT Mathematics, ACT Reading, ACT English, and ACT Science Reasoning.



LMR\_U2\_L16



3. Display the group of bell-shaped distributions from page 1 of the *Normal Plots* file (LMR\_U2\_L16) to the class and ask the students:
  - a. What characteristic does this particular group share? **Answer:** *All of these plots are unimodal (one mode/peak) and symmetric.*
  - b. Inform students that these types of distributions are often referred to as bell-shaped. Why might this term be used? **Answer:** *The histograms look very similar to the shape of a bell.*

4. To show the similarities between the shape of a bell and the shape of these distributions, a clip-art image (shown here) has been included in the *Normal Plots* file (LMR\_U2\_L16) on page 2.



5. Explain that this shape occurs often in real-life. It occurs so often that it's been given its own name: the **normal curve**, or **normal distribution**. Can the students think of distributions where they have seen Normal curves in previous labs? *A possible answer might be the height and weight histograms for the cdc data.*

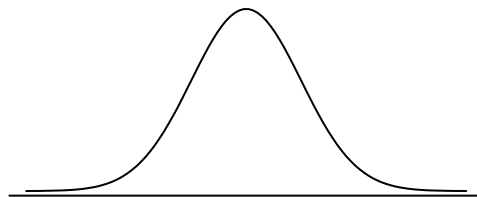


6. To give some more background on the normal distribution, play the New York Times video titled "Bunnies, Dragons, and the Normal World" found at:  
<http://www.nytimes.com/video/science/100000002452709/bunnies-dragons-and-the-normal-world.html>

**Note:** Show only the first 41 seconds of the video.

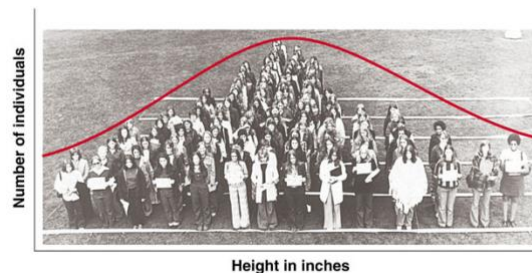
7. Discuss that the normal curve has a very precise mathematical definition, which is pretty complex. But the result is a curve that looks like the one in the "Bunnies" video. In general, the curve looks like the plot shown below.

**Note:** You can either draw the diagram below on the board or display it via a projector – the image can be found on page 2 of the *Normal Plots* file (LMR\_U2\_L16).



8. Explain that normal distributions are good for describing some populations of people. For example, people's heights are often considered to be normally distributed. Display the famous Frank Anscombe photograph (shown below) via a projector. The graphic can be found on page 2 of the *Normal Plots* file (LMR\_U2\_L16). Inform the students that this photo was taken of a group of randomly selected college women who stood in height order.

Tobin/Dusheck, Asking About Life, 2/e  
Figure 16.6



Copyright © 2001 by Harcourt, Inc. All rights reserved.



9. Next, lead a discussion about why the normal curve is a good fit to the histogram in the above picture. Notice that more people are near the center of the distribution, and fewer are in the outer edges, or tails. Engage the class in a conversation using the following probing questions:

- a. Notice that there is a peak in the center of the distribution. What height do you think is at the center? *Answer: The average height of American women is approximately 5'5" (5 feet, 5 inches) tall. We might therefore expect the average height of this group to be close to 5'5" as well.*
- b. Why are more people in the center, and less people in the edges, or tails, of the distribution? *Answer: The center represents the mean height. Most women will fall somewhere close to the mean and may be a few inches shorter or taller than it. However, less people are likely to be MUCH shorter or MUCH taller than the mean. For example, we would not expect to see many women who are 4'10" tall nor would we expect to see many women who are 6'0" tall.*



10. Explain that the normal curve is a good description of a distribution when it makes sense that there is a single 'typical' value with random deviations above and below that value. Ask students:

- a. Why does this make sense with heights but not with incomes? *Answer: With heights, we expect most people to be near the average, with some deviations above and below the mean (people who are taller or shorter than the mean height). However, with incomes, the distribution will have more deviations that are above the typical value, since there is no upper bound for a maximum income (ex. Bill Gates, Warren Buffett, etc.).*
- b. Are there more real-life examples, other than height, that students think might follow a normal distribution? *Answers will vary by class. Some examples include: (1) scores on standardized math and reading tests (like the SAT and ACT), (2) IQ scores, and (3) body temperatures.*
- c. Does it matter that the curve drawn on the photograph does not match exactly to the women's heights? *Answer: No. We often refer to the curve as the "normal model" because the curve is just a model of the true population distribution. So, even though the red curve is not exactly the same as the women's heights, it is a close enough approximation of the shape of their heights.*

**Note to teacher:** The main role the normal distribution has historically played has been in modeling errors (most measurements will be close to the actual value while larger errors occur less often) and sample means.

11. Inform the students that, during the next few lessons, they will be learning more about the normal distribution. In particular, they will learn about a new measure of spread used to describe a normal distribution, how to calculate probabilities from this distribution, and how to randomly sample from this distribution.



12. **Cheat Card:** Distribute an index card to students and ask them to create a cheat card that will help them remember information about the normal curve.

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework

Students will complete their cheat cards if they were not able to finish in class.

## Lesson 17: A Normal Measure of Spread

### Objective:

Students will learn that standard deviation is another way to measure variability.

### Materials:

1. *How Far Apart?* handout (LMR\_U2\_L4) – completed during Lesson 4
2. *How Far Apart? (with standard deviation – SD)* handout (LMR\_U2\_L17)
3. Projector to display visuals using RStudio
4. RScript with the functions in this lesson

### Vocabulary:

standard deviation (SD)

**Essential Concepts:** The standard deviation is another measure of spread. This is commonly used by statisticians because of its role in common models and distributions, such as the Normal Model.

### Lesson:

1. In their DS journals, ask students to create a two-column table and label the left-column as *Measures of Center (Central Tendency)* and the right column as *Measures of Spread (Dispersion)*.
2. In pairs, ask students to recall methods they have learned so far for measuring center and measuring spread in distributions.  
Measures of Center: *Answer: mean (average or typical value), median*  
Measures of Spread: *Answer: mean absolute deviation (MAD), interquartile range (IQR)*
3. Share out a pair's explanation and ask the rest of the pairs to agree or disagree. If there is disagreement, hold a class discussion until the lists are correct.
4. Point out that a measure of center or a measure of spread depicts one value for a distribution. Ask student pairs to discuss the following question:
  - a. What does the value of each measure tell us about the data in the distribution? *Possible answer: A measure of center tells us the value that is typical, or in the center. A measure of spread tells us how variable, or how spread apart, the data are.*
5. Next, ask students to add the term **standard deviation (SD)** to their *Measures of Spread* column.
6. Inform students that the standard deviation of a distribution is another way to measure spread, or variability. The standard deviation is similar to the mean absolute deviation (MAD).
7. Ask students to recall the formula for calculating the MAD:

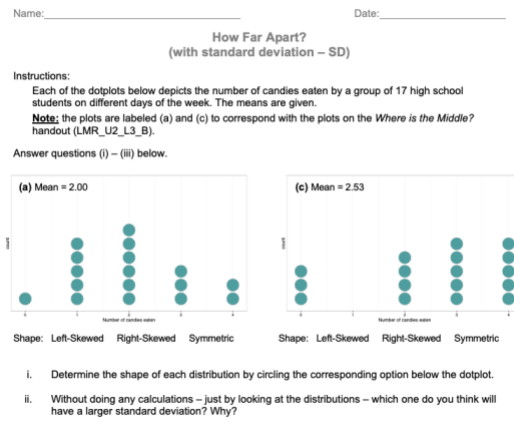
$$MAD = \frac{\sum_{i=1}^n |x_i - \bar{x}|}{n}$$

8. While the MAD measures the absolute distance of each data point from the mean, the standard deviation squares the distances of each data point from the mean. Both methods result in positive measurements because distance is always positive.
9. Ask students to recall that they calculated MAD values in the *How Far Apart?* handout (LMR\_U2\_L4) during Lesson 4 of this unit.
10. Show and discuss the formula for calculating the standard deviation of a dataset:

$$s = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n}}$$

**Note to teacher:** There are different formulas for the standard deviation. We are presenting the simpler one, which divides by  $n$ . In AP Statistics (or college introductory statistics), students will learn that if they are using a sample of data to estimate the standard deviation for the population, then dividing by  $n - 1$  is a better estimator than dividing by  $n$ . But this technically requires a lot of scaffolding and leads to little understanding, and so we will stick with the simpler version. (In some books, this is called the “population value of the standard deviation” and the  $n - 1$  version is called the “sample estimate of the standard deviation.”)

11. Guide the class to complete the *How Far Apart? (with standard deviation -- SD)* handout (LMR\_U2\_L17) to calculate standard deviations of the dotplots using the formula listed above.



LMR\_U2\_L17

*Answers: Plot (a) – SD = 1.0847 candies; Plot (c) – SD = 1.3770 candies*

12. As a whole group, ask students to compare and contrast the standard deviations with the MAD values for the two plots in the handout.

*Similarities between SD and MAD:*

- *Measure the same idea: variability, or spread*
- *Are based on looking at the “deviations” from the mean: the difference between an observation and the mean*
- *Uses the “typical” deviation*

*Differences between SD and MAD:*

- *The MAD uses the absolute value and finds the average of the absolute deviations from the mean*
- *The SD uses the square of each deviation from the mean, and finds the average of the squares*
- *The SD takes the square root of the average of the squares*

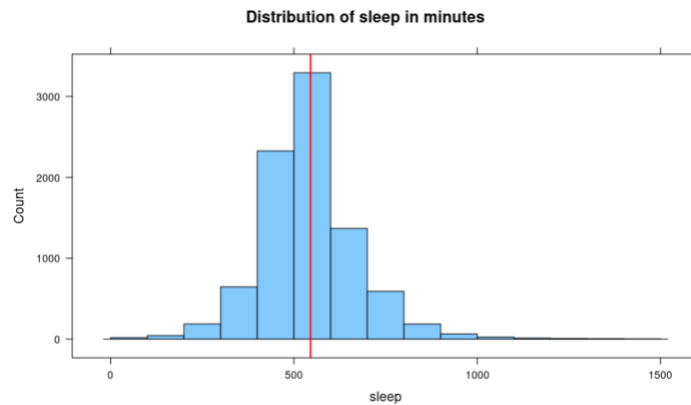
13. Ask students why they think the SD takes the square root of the average of the squares. *Possible response: Taking the square root of the average of the squares returns the measurements to their original units instead of square units.*

14. To reinforce students’ conceptual understanding of standard deviation, student teams will estimate the standard deviation for a few numerical distributions and explain the reasoning for their estimate. Load and view the **atus** data, then run the following functions one by one:

```
> histogram(~sleep, data = atus, breaks = seq(0, 1500, by=100),
 main = "Distribution of sleep in minutes")

> sleep_mean <- mean(~sleep, data = atus)

> add_vline(vline = sleep_mean)
```



15. Zoom in on the visual and give student teams a few minutes to discuss what they estimate the standard deviation to be. Have the reporter from each team report their estimate using the following sentence frame:

“The time spent sleeping (in minutes) typically varies from the mean by \_\_\_\_\_ minutes.”

16. Reveal the actual standard deviation by running the function:

```
> sd(~sleep, data = atus)
```

17. Choose a reporter from a student team that had a good approximation to explain their reasoning.

18. Repeat this process with a few more numerical variables. Functions are provided below.

Household Activities

```
> histogram(~hh_activities, data = atus, nint = 13)
> hh_activities_mean <- mean(~hh_activities, data = atus)
> add_line(vline = hh_activities_mean)
```

“The time spent on household activities typically vary from the mean by \_\_\_\_\_ minutes.”

```
> sd(~hh_activities, data = atus)
```

Socializing

```
> histogram(~socialize, data = atus, breaks = seq(0, 2000, by = 100))
> social_mean <- mean(~socialize, data = atus)
> add_line(vline = social_mean)
```

“The time spent socializing (in minutes) typically varies from the mean by \_\_\_\_\_ minutes.”

```
> sd(~socialize, data = atus)
```

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 18: What's Your Z-Score?

### Objective:

Students will understand that a z-score can be used to measure how far away – or how many standard deviations – an observation is away from the mean. Typically, z-scores will range between -3 and +3. For simulations involving shuffling, if we compute a z-score that lies far away from the mean, then we might conclude that the outcome was not due to chance. If we see a z-score that lies close to the mean, then we might conclude it was by chance.

### Materials:

1. Projector to display RStudio function
2. RScript with all of the functions in this lesson
3. A ruler with centimeter marks on it

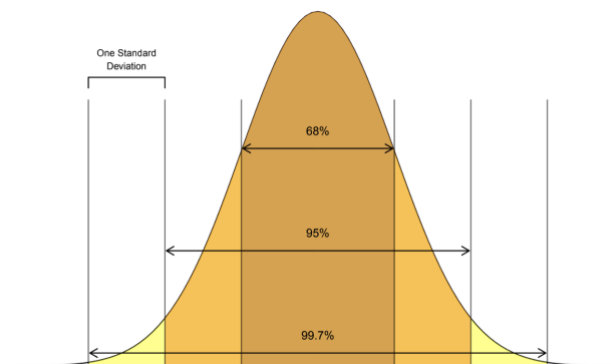
### Vocabulary:

z-score, standardized score, Empirical Rule

**Essential Concepts:** Z-scores offer us a way to measure how extreme a value is, regardless of the units of measurement. Typically, z-scores will range between -3 and +3, so values that are at or more extreme than -3 and +3 standard deviations are considered extremely rare.

### Lesson:

1. Ask students to recall what they remember about normal distributions.  
*Answer: Normal distributions are unimodal and symmetric and are often referred to as bell-shaped. Some real-life examples of variables that produce normal distributions are people's heights, scores on standardized tests, and body temperatures.*
2. Display the following statement to students: "All normal distributions are bell-shaped, but not all bell-shaped distributions are normal." Then inform students that normal distributions have special properties.
3. Display the image below and introduce the **Empirical Rule**, which states:
  - Approximately 68% of the observations in a normal distribution fall within one standard deviation of the mean
  - Approximately 95% of the observations in a normal distribution fall within 2 standard deviations of the mean
  - Approximately 99.7% of the observations in a normal distribution fall within 3 standard deviations of the mean



4. Open RStudio and project for students to see. Load the babies dataset, named **Gestation**, by following these steps:
  - Enter `data(Gestation)` in the Console
  - You should see **Gestation** located in your Environment
  - Enter `View(Gestation)` in the Console

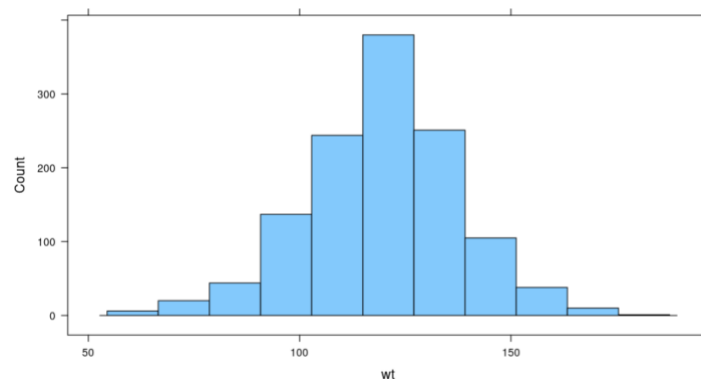
Display the help documentation by typing **?Gestation**. Ask student teams to predict which of the variables in the “Gestation” dataset they think might be normally distributed. Choose a couple of teams to share out.

Description of numerical variables:

- wt – birth weight (in ounces)
- gestation – length of pregnancy (in days)
- parity – 0 if baby was first born, 1-13 otherwise
- age – mother’s age (in years)
- ht – mother’s height (to the last completed inch)
- wt.1 – mother’s weight (in pounds)
- dage – father’s age (in years)
- dht – father’s height (to the last completed inch)
- dwl – father’s weight (in pounds)

5. Create histograms using the variables shared by student teams. There are a few variables that look normally distributed, such as birth mother’s heights. We will investigate the babies’ birth weights.

**histogram(~wt, data = Gestation)**



6. Ask students:

- a. Does the distribution of baby birth weights look approximately normal? Explain. *Answer: The distribution of baby birth weights is unimodal, roughly symmetric, and somewhat bell-shaped, so it might be approximately normal.*



- b. What do you approximate the mean weight of the distribution to be? How about the standard deviation? *Answers will vary. Use this as a check for understanding of standard deviation as well as estimating the mean using the balancing point concept. See next step for calculating the actual mean weight and standard deviation.*

7. Use RStudio to calculate the actual mean and standard deviation.

**mean\_wt <- mean(~wt, data = Gestation)**

**sd\_wt <- sd(~wt, data = Gestation)**

|         |                  |
|---------|------------------|
| mean_wt | 119.576860841424 |
| sd_wt   | 18.2364518669726 |



8. Have students draw a number line with seven equally spaced intervals and label it “Baby birth weight in ounces.” Make sure students leave about 5 centimeters of space above the number line to draw a normal curve. Have students label the middle tick mark with the mean baby weight (round to the nearest tenth of an ounce = 119.6 ounces). Then ask students:

- a. What weight is one standard deviation above the mean? *Answer: A baby whose weight is 137.8 ounces is one standard deviation above the mean baby weight.*

- b. What weight is one standard deviation below the mean? *Answer: A baby whose weight is 101.4 ounces is one standard deviation below the mean baby weight.*

Have students label their number line with these values.



9. Have students continue filling their number line with the corresponding weights that are two and three standard deviations from the mean. *Answer: A baby who weighs 156 ounces is two standard deviations above the mean weight, and a baby who weighs 174.2 ounces is three standard deviations above the mean weight. A baby who weighs 83.2 ounces is two standard deviations below the mean weight, and a baby who weighs 65 ounces is three standard deviations below the mean weight.*

10. Ask students: If the distribution of baby weights is approximately normal, what percentage of babies weigh between 101.4 and 137.8 ounces? *Answer: If the distribution of baby weights is approximately normal, about 68% of babies should be between 101.4 ounces and 137.8 ounces.*

11. Use RStudio to confirm if indeed the distribution of baby weights is approximately normal.

```
> one_sd_wt <- filter(Gestation, wt > 101.4, wt < 137.8)
```

*Answer: In this sample of 1236 observations, there are 861 babies whose weights are one standard deviation from the mean, so  $861/1236 = 0.697$ . This means that around 69.7% of the weights of babies in this sample fall within one standard deviation from the mean baby weight. This is close to 68%, so it seems that the distribution of baby weights is approximately normally distributed.*

**Note:** If you continue this process for this sample, you will find that the distribution of baby weights is normally distributed as defined by the Empirical Rule. In this sample,  $1171/1236 = 94.7\%$  of the baby weights fall within two standard deviations of the mean, and  $1229/1236 = 99.4\%$  of the baby weights fall within three standard deviations of the mean.

12. Now that it has been verified that a normal distribution is an appropriate model for this distribution, have students draw a normal curve above the number line. Suggested method to obtain a decent normal curve:

- Step 1: Draw a dot 4 centimeters above the mean height.
- Step 2: Draw dots 2.4 cm above the heights that are 1 standard deviation from the mean.
- Step 3: Draw dots 0.36 cm above the heights that are 2 standard deviations from the mean.
- Step 4: Draw dots right above the number line for the heights that are 3 standard deviations from the mean.
- Step 5: Connect the dots with a smooth curve.



13. Tell students that we are using this normal curve as a model to represent the distribution of all baby weights. This will allow us to make comparisons, draw conclusions, and make predictions about baby weights. Let's see:

- a. What percentage of babies weigh less than 119.6 ounces? Explain. *Answer: About 50% of babies weigh less than 119.6 ounces. Since normal distributions are symmetric, the mean and the median are about the same. Since the median divides a distribution into equal halves, in this case so does the mean.*
- b. What percentage of babies weigh between 119.6 and 137.8 ounces? *Answer: About 34% of babies weigh between 119.6 and 137.8 ounces. According to the Empirical rule, 68% of the observations fall within one standard deviation of the mean, and since normal distributions are symmetric, the area under the curve from the mean to one standard deviation is half of 68% or 34%.*
- c. What percentage of babies weigh more than 137.8 ounces? *Answer: About 16% of babies weigh more than 137.8 ounces. From part a and b above, we know that  $50\% + 34\% = 84\%$  of babies weigh less than 137.8 ounces, so  $100\% - 84\% = 16\%$  weigh more than 137.8 ounces.*

14. Explain that statisticians use something called a **z-score** to compare values. A z-score tells us how many standard deviations away from the mean an observation is. Another name for z-score is a **standardized score**.

15. Introduce the formula for calculating a z-score and discuss what each symbol means.

$$z = \frac{x - \bar{x}}{s}$$

16. Explain that z-scores answer the question: “How typical is x?” If x is the same as the typical value (the mean), then  $z = 0$ . If x is one standard deviation away from the mean, then  $z = -1$  or  $+1$ . Remind students from the normal curve that as you move farther from the center (from the mean), there are fewer observations. Therefore, a large z-score is considered an unusual value.



17. Ask: Should we be concerned if a baby weighs 100 ounces? Let's find out.

18. Have students calculate the z-score for a baby that weighs 100 ounces:

$$z = (100 - 119.6) / 18.2 = -1.08$$

Ask the class:

- What does a negative z-score mean? *Answer: A negative z-score means the x value is below the mean. This means that the weight is below average.*
  - What does a positive z-score mean? *Answer: A positive z-score means the x value is above the mean. This means that the weight is above average.*
  - What is the most negative z-score you think we will find? What is the most positive z-score? *Answer: Typically, values in a normal distribution rarely fall outside two or three standard deviations from the mean. So, if our data is purely by chance, we probably won't see any values that are less than -3 or greater than +3.*
19. Ask students: “Where does a baby that weighs 100 ounces fall within the distribution of baby weights?” Have students find 100 ounces on the x-axis of the normal curve and draw a vertical line from the x-axis until it intersects the normal curve. Have them shade the area under the curve to the left of the vertical line.
20. Tell students that the shaded area represents a percentile in the distribution. A percentile is the exact value in which the desired proportion of observations lie below the specific value in a distribution. Use RStudio to calculate the percentile.

$$\text{pnorm}(100, \text{mean} = 119.6, \text{sd} = 18.2) = 0.140$$

21. Doctors report percentiles to describe a child's development compared to other children their age. For a baby that weighs 100 ounces, a doctor would report the following: “The baby is at the 14<sup>th</sup> percentile in weight.” This means that the baby weighs more than 14% of all babies. This is between one and two standard deviations from the mean so there is no cause for concern with this baby's weight.

**Note:** A z-score can also be used to calculate a percentile, but since a z-score is a standardized score, the mean of the distribution would be zero and the standard deviation would be one.

$$\text{pnorm}(-1.08, \text{mean} = 0, \text{sd} = 1) = 0.140$$

22. Inform the class that they will be using RStudio in the next labs to practice using normal models.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Next 2 Days

## Lab 2H: Eyeballing Normal

## Lab 2I: R's Normal Distribution Alphabet

Complete Labs 2H and 2I prior to the End of Unit Project.

## Lab 2H: Eyeballing Normal

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### What's normal?

- The **normal distribution** is a curve we often see in real data.
  - We see it in people's blood pressures and in measurement errors.
- When data appears to be *normally distributed*, we can use the *normal model* to:
  - Simulate *normally distributed* data.
  - Easily compute probabilities.
- In this lab, we'll look at some previous datasets to see if we can find data that are roughly normally distributed.

### The normal distribution

- The normal distribution is *symmetric about the mean*:
  - The **mean** is found in the very center of the distribution.
  - And the curve looks the same to the left of the mean as it does on the right.
- **Use the following to draw a normal distribution:**

```
plotDist('norm', mean = 0, sd = 1)
```

### The mean and sd of it

- To draw a normal curve, we need to know exactly 2 things:
  - The **mean** and **sd**.
- The **sd**, or **standard deviation**, is a measure of spread that's similar to the **MAD**.
- **(1) Which part of the normal curve changes when the value of the mean changes?**
- **(2) Which part of the normal curve changes when the value of the sd changes?**
- *Hint:* Try changing the **mean** and **sd** values in the **plotDist** function.

### Finding normal distributions

- **Load the cdc data and answer the following:**
- **(3) Think about the height and weight variables. Based on what you know about these variables, which of the variables do you think have distributions that will look like the normal distribution?**
  - **(4) Make histograms of these variables. Which ones look like the normal distribution?**
  - *Hint:* To help answer this question, try including the option **fit = "normal"** in the histogram function. You might also try faceting by **sex**.

### Using normal models

- Data scientists like using normal models because it often resembles real data.
  - *But not EVERYTHING is normally distributed.*
- As a data scientist in training, you must decide when a normal model seems appropriate.
  - No model is ever perfect 100% of the time.
  - If you choose a model, you should be able to justify why you chose it.

### On your own

- For each of the following, determine which, if any, appear to be normally distributed. Explain your reasoning:
  - (5) The difference in **percentages** between male and female survivors in a slasher film for 500 random shuffles.
  - (6) The difference in **median** fares between survivors and non-survivors on the Titanic for 500 random shuffles.
  - (7) The difference in **mean** fares between survivors and non-survivors on the Titanic for 500 random shuffles.
- Hint: Refer to Lab 2E and 2F.

## Lab 2I: R's Normal Distribution Alphabet

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Where we're headed

- In the last lab, you were able to overlay a normal curve on histograms of data to help you decide if the data's distribution is close to a normal distribution.
  - We also saw that calculating the **mean** of random shuffles also produces differences that are normally distributed.
- In this lab, we'll learn how to use some other **R** functions to:
  - Simulate random draws from a normal distribution.
  - Calculate probabilities with normal distributions.

### Get set up

- **(1) Start by loading the **titanic** data. Write and run code calculating the **mean age** of people in the data but **shuffle their survival status 500 times**.**
  - **Assign this data the name **shfls**.**
- **(2) After creating **shfls**, write and run code using **mutate** to add a new variable to the dataset.**
  - This new variable should have the name **diff** and should be the **mean age** of those who survived minus those who died.
- **(3) Finally, write and run code calculating the **mean** and **sd** of the **diff** variable.**
  - **Assign these values the name **diff\_mean** and **diff\_sd**.**

### Is it normal?

- Before we proceed, we need to verify that our **diff** variable looks approximately normally distributed.
  - **(4) Is the distribution close to normal? Explain how you determined this. Describe the center and spread of the distribution.**
  - **(5) Compute and write down the mean difference in the age of the *actual* survivors and the *actual* non-survivors.**

### Using the normal model

- Since the distribution of our **diff** variable appears normally distributed, we can use a normal model to estimate the probability of seeing differences that are more extreme than our actual data.
- **(6) Draw a sketch of a normal curve. Label the mean age difference, based on your shuffles, and the actual age difference of survivors minus non-survivors from the actual data. Then shade in the area, under the normal curve, that is *smaller* than the actual difference.**
- **(7) Fill in the blanks to calculate the probability of an even smaller difference occurring than our actual difference using a normal model.**

```
pnorm(____, mean = diff_mean, sd = ____)
```

### Extreme probabilities

- The probability you calculated in the previous slide is an estimate for how often we expect to see a difference smaller than the actual one we observed, by chance alone.

- (8) If you wanted to instead calculate the probability that the difference would be larger than the one observed, we could run (fill in the blanks):

```
1 - pnorm(____, mean = diff_mean, sd = ____)
```

### Simulating normal draws

- We can simulate random draws from a normal distribution with the `rnorm` function.
- (9) Fill in the blanks in the following two lines of code to simulate 100 heights of randomly chosen men. Assume the mean height is 67 inches and the standard deviation is 3 inches.

```
draws <- rnorm(____, mean = ____, sd = ____)
```

- (10) Fill in the blank below to plot your simulated heights with a histogram.

```
histogram(draws, fit = ____)
```

### P's and Q's

- We've seen that we can use `pnorm` to calculate *probabilities* based on a specified *quantity*.
  - Hence, why we call it “P” norm.
- Now we'll see how to do the opposite. That is, calculate the *quantity* for a specific *probability*.
  - Hence, why we'll call this a “Q” norm.
- (11) How tall can a man be and still be in the shortest 25% of heights if the mean height is 67 inches with a standard deviation of 3 inches?

```
qnorm(____, mean = ____, sd = ____)
```

### On your own

Conduct one of the statistical investigations below:

- Using the `titanic` data:
  - (12) Were women on the Titanic typically younger than men? Write and run code using a histogram, 500 random shuffles and a normal model to answer the question.
- Using the `cdc` data:
  - (13) Using 500 random shuffles and a normal model, how much taller would the typical male have to be than the typical female in order for the difference to be in the upper 1% by chance alone?
  - (14) How can we use this value to justify the claim that the average Male in our data is taller than the average Female?

## **End of Unit 2 Project: Comparing Groups Using Our Own Data**

### **Objective:**

Students will apply their learning of the first and second units of the curriculum by completing an end of unit project.

### **Materials:**

1. *End of Unit 2 Project: Comparing Groups Using Our Own Data* (LMR\_U2\_End\_of\_Unit\_Project)

## **End of Unit 2 Project Comparing Groups Using Our Own Data**

Available datasets:

1. Food Habits
2. Time Use
3. Stress/Chill
4. Personality Color

Your mission is to ask and answer a statistical question using at least one dataset above.

1. Your question must include a comparison of two distinct groups.
2. Your analysis should address whether any observed differences are real or could be simply due to chance.
3. You should use at least two of the following methods to answer your question with appropriate explanations:
  - ☐ Merge data
  - ☐ Create simulations
  - ☐ Calculate probabilities based on simulations
  - ☐ Use a Normal model
  - ☐ Shuffle/permute data

You will have 5 days to complete this project with your assigned partner. You need to:

- ☐ Prepare a presentation (both partners need to participate) that includes:
  - o A 4-slide, 5-minute presentation.
  - o An explanation of why you think your statistical question is interesting.
  - o An interpretation of supporting plots and summaries that answer your question.
  - o A reasoning of whether you think the outcome might be due to chance.
- ☐ Submit a 2-4 page typed, double-spaced summary of your analysis.



### **Project Assignment Sequence:**

- ☐ Day 1: Decide on a statistical question with assigned partner.
  - ☐ Get approval from teacher.
- ☐ Day 2: Working day for analysis – create plots and numerical summaries.
- ☐ Day 3: Working day for analysis – create presentation (4 slides maximum).
- ☐ Day 4: Presentations.
- ☐ Day 5: Presentations.



# Introduction to Data Science

## Unit 3

# Introduction to Data Science

## Daily Overview: Unit 3

| Theme                                         | Day   | Lessons and Labs                                               | Campaign            | Topics                                                      | Page |
|-----------------------------------------------|-------|----------------------------------------------------------------|---------------------|-------------------------------------------------------------|------|
| Testing, Testing...<br>1, 2, 3...<br>(7 days) | 1     | Lesson 1: Anecdotes vs. Data                                   |                     | Reading articles critically, data                           | 217  |
|                                               | 2     | Lesson 2: What is an Experiment?                               |                     | Experiments, causation                                      | 220  |
|                                               | 3     | Lesson 3: Let's Try an Experiment!                             |                     | Random assignments, confounding factors                     | 223  |
|                                               | 4     | Lesson 4: Predictions, Predictions                             |                     | Visualizations, predictions                                 | 225  |
|                                               | 5     | Lesson 5: Time Perception Experiment                           |                     | Elements of an experiment                                   | 227  |
|                                               | 6     | Lab 3A: The Results Are In!                                    |                     | Analyzing experiment data                                   | 229  |
|                                               | 7     | Practicum: TB or Not TB?                                       | Time Perception     | Simulation using experiment data                            | 230  |
| Would You Look at That?<br>(4 days)           | 8     | Lesson 6: Observational Studies                                |                     | Observational study                                         | 234  |
|                                               | 9     | Lesson 7: Observational Studies vs. Experiments                |                     | Observational study, experiment                             | 236  |
|                                               | 10    | Lesson 8: Monsters that Hide in Observational Studies          |                     | Observational studies, confounding factors                  | 238  |
|                                               | 11    | Lab 3B: Confound It All!                                       |                     | Confounding factors                                         | 242  |
| Are You Asking Me?<br>(9 days)                | 12    | Lesson 9: Survey Says...                                       |                     | Survey                                                      | 245  |
|                                               | 13    | Lesson 10: We're So Random                                     |                     | Data collection, random samples                             | 248  |
|                                               | 14    | Lesson 11: The Gettysburg Address                              |                     | Sampling bias                                               | 252  |
|                                               | 15    | Lab 3C: Random Sampling                                        |                     | Random sampling                                             | 257  |
|                                               | 16    | Lesson 12: Bias in Survey Sampling                             |                     | Bias in survey sampling                                     | 259  |
|                                               | 17    | Lesson 13: The Confidence Game                                 |                     | Confidence intervals                                        | 262  |
|                                               | 18    | Lesson 14: How Confident Are You?                              |                     | Confidence intervals, margin of error                       | 265  |
|                                               | 19    | Lab 3D: Are You Sure about That?                               |                     | Bootstrapping                                               | 267  |
|                                               | 20    | Practicum: Let's Build a Survey!                               |                     | Survey design with non-leading questions                    | 270  |
| What's the Trigger?<br>(5 days)               | 21    | Lesson 15 Ready, Sense, Go!                                    |                     | Sensors, data collection                                    | 272  |
|                                               | 22    | Lesson 16: Does it have a Trigger?                             |                     | Survey questions, sensor questions                          | 275  |
|                                               | 23    | Lesson 17: Creating Our Own Participatory Sensing Campaign     |                     | Participatory Sensing campaign creation                     | 277  |
|                                               | 24    | Lesson 18: Evaluating Our Own Participatory Sensing Campaign   |                     | Statistical questions, evaluate campaign                    | 280  |
|                                               | 25^   | Lesson 19: Implementing Our Own Participatory Sensing Campaign | Class Campaign—data | Mock-implement campaign, campaign creation, data collection | 282  |
| Webpages<br>(5 days)                          | 26    | Lesson 20: Online Data-ing                                     | Class Campaign—data | Data on the internet                                        | 285  |
|                                               | 27    | Lab 3E: Scraping Web Data                                      | Class Campaign—data | Scraping data from the internet                             | 288  |
|                                               | 28    | Lab 3F: Maps                                                   | Class Campaign—data | Making maps with data from the internet                     | 290  |
|                                               | 29    | Lesson 21: Learning to Love XML                                | Class Campaign—data | Data storage, XML                                           | 292  |
|                                               | 30+   | Lesson 22: Changing Format                                     | Class Campaign—data | Converting XML files                                        | 297  |
| End of Unit Project<br>(5 days)               | 31-36 | End of Unit 3 Project: Analyzing Our Own Campaign Data         | Class Campaign      | Statistical questions, analyzing our data                   | 300  |

^=Data collection window begins.

+=Data collection window ends.

## IDS Unit 3: Essential Concepts

### **Lesson 1: Anecdotes vs. Data**

Data beat anecdotes. In science, we need to closely examine the quality of evidence in order to make sound conclusions. Anecdotes can contain personal bias, might be carefully selected to represent a particular point of view, and, in general, may be completely different from the general trend.

### **Lesson 2: What is an Experiment?**

Science is often concerned with the question “What causes things to happen?” To answer this, controlled experiments are required. Controlled experiments have several key features: (1) there is a treatment variable and a response variable, and we wish to see if the treatment causes a change that we can measure with the response variable; (2) There is a comparison/control group; (3) Subjects are assigned randomly to treatment or control (randomized assignment); (4) Subjects are not aware of which group they are in (a ‘blind’). This may require the use of a placebo for those in the control group; and (5) those who measure the response variable do not know which group the subjects were in (if both 4 and 5 are satisfied, this is a ‘double blind’ experiment).

### **Lesson 3: Let’s Try an Experiment!**

Randomized assignment is required to determine cause-and-effect.

### **Lesson 4: Predictions, Predictions**

Designing an experiment requires making many decisions, including what to measure and how to measure it.

### **Lesson 5: Time Perception Experiment**

Designing and carrying out an experiment helps us answer specific statistical questions of interest.

### **Lesson 6: Observational Studies**

Observational studies are those for which there is no intervention applied by researchers.

### **Lesson 7: Observational Studies vs. Experiments**

Experiments are not always possible because of various factors such as ethics, cost limitations, and feasibility.

### **Lesson 8: Monsters that Hide in Observational Studies**

Confounding factors/variables make it difficult to determine a cause-and-effect relationship between two variables.

### **Lesson 9: Survey Says...**

Surveys ask simple, straightforward questions in order to collect data that can be used to answer statistical questions. Writing such questions can be hard (but fun)!

### **Lesson 10: We’re So Random**

Another popular data collection method involves collecting data from a random sample of people or objects. Percentages based on random samples tend to “center” on the population parameter value.

### **Lesson 11: The Gettysburg Address**

Statistics vary from sample to sample. If the typical value across many samples is equal to the population parameter, the statistic is “unbiased.” Bias means that we tend to “miss the mark.” If we don’t do random sampling, we can get biased estimates.

### **Lesson 12: Bias in Survey Sampling**

Bias concerning survey sampling includes identifying sampling methods that may lead to biased samples, recognizing potential over- or under-representation in samples, and acquiring skills to choose more reliable sampling techniques.

### **Lesson 13: The Confidence Game**

There is uncertainty when we estimate population parameters. Because of this, it is better to give a range of plausible values, rather than a single value.

### **Lesson 14: How Confident Are You?**

The margin of error expresses our uncertainty in an estimate. The estimate, plus or minus the margin of error, gives us an interval in which we are very confident the true value lies.

### **Lesson 15 Ready, Sense, Go!**

Sensors are another data collection method. Unlike what we have seen so far, sensors do not involve humans (much). They collect data according to an algorithm.

### **Lesson 16: Does it have a Trigger?**

A key feature that distinguishes the way sensors collect data from more traditional approaches is that sensors collect data when a 'trigger' event occurs. In Participatory Sensing, this event is something we humans agree upon beforehand. Every time that trigger happens, we collect data.

### **Lesson 17: Creating Our Own Participatory Sensing Campaign**

Creating a Participatory Sensing Campaign requires that survey questions must be completed whenever they are "triggered." Research questions provide an overall direction in a Participatory Sensing Campaign.

### **Lesson 18: Evaluating Our Own Participatory Sensing Campaign**

Statistical questions guide a Participatory Sensing Campaign so that we can learn about a community or ourselves. These campaigns should be evaluated before implementing to make sure they are reasonable and ethically sound.

### **Lesson 19: Implementing Our Own Participatory Sensing Campaign**

Practicing data collection prior to implementation allows optimization of a Participatory Sensing Campaign.

### **Lesson 20: Online Data-ing**

We stretch students' conception of data, to help them see that many web pages present information that can be turned into data.

### **Lesson 21: Learning to Love XML**

XML is a programming language that we use with our campaigns. We create basic XML "tags" in the code, which help us store data in a format we understand.

### **Lesson 22: Changing Format**

Converting XML to spreadsheet format helps us better understand and view our data.

# Testing, Testing...1, 2, 3...

Instructional Days: 7

## Enduring Understandings

An experiment is a data collection method in which the effects of different treatments on an outcome of interest are measured. In an experiment, a treatment is applied to subjects and then observations about the effect of the treatment are made. To isolate the effects from unexplained variation, randomization (or chance) assignment to treatments is applied.

## Engagement

Students will view Hans Rosling's video *How Not to Be Ignorant About the World* and will participate in his interactive quiz in order to learn how anecdotes and personal experience can influence what we know and, alternatively, how data provides basis for evidence. The video can be found at:

[https://www.ted.com/talks/hans\\_and\\_ola\\_rosling\\_how\\_not\\_to\\_be\\_ignorant\\_about\\_the\\_world](https://www.ted.com/talks/hans_and_ola_rosling_how_not_to_be_ignorant_about_the_world)

## Learning Objectives

### Statistical/Mathematical:

- S-IC B-5: Use data from a randomized experiment to compare two treatments; use simulations to decide if differences between parameters are significant.
- S-IC 1: Understand statistics as a process for making inferences about population parameters based on a random sample from that population.
- S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.
- S-IC 6: Evaluate reports based on data.

### Focus Standards for Mathematical Practice for All of Unit 3:

- SMP-1: Make sense of problems and persevere in solving them.
- SMP-4: Model with mathematics.
- SMP-8: Look for and express regularity in repeated reasoning.

### Data Science:

Understand that differences between the measured outcomes of the treatment and control groups in an experiment can be tested. Understand the roles of randomization and of random sampling in statistical inference.

### Applied Computational Thinking:

- Test for differences between experimental groups.
- Create graphical representations to compare data between experimental groups.
- Write code to randomly assign subjects to treatment groups.

### Real-World Connections:

Experiments are used to ensure safety and efficacy of medicines, reliability of electronics and structural materials and find patterns in human behavior.

## Language Objectives

1. Students will compare and contrast anecdotes with data.
2. Students will construct summary statements about their understanding of data, how it is collected, how it is used, and how to work with it when creating and conducting experiments.
3. Students will read informative texts to evaluate claims based on data and anticipate visualizations of that data.
4. Students will describe causal relationships.

## Data File or Data Collection Method

### Data File:

1. Students' *Time Perception* experiment data.

### Data Collection Method:

Students will gather data generated through a simple experiment.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 1: Anecdotes vs. Data

### Objective:

Students will learn the difference between anecdotes and data. They will begin to read articles critically to discern whether the evidence presented is based on anecdotes or data.

### Materials:

1. Hans Rosling's video *How Not to Be Ignorant About the World* found at [https://www.ted.com/talks/hans\\_and\\_ola\\_rosling\\_how\\_not\\_to\\_be\\_ignorant\\_about\\_the\\_world](https://www.ted.com/talks/hans_and_ola_rosling_how_not_to_be_ignorant_about_the_world)
2. *Miracle at the KK Café* Article (LMR\_Miracle\_Cafe)
3. *Can Trophy Hunting Actually Help Conservation?* Article (LMR\_Trophy\_Hunting\_Conversation)

### Vocabulary:

anecdote, data

**Essential Concepts:** Data beat anecdotes. In science, we need to closely examine the quality of evidence in order to make sound conclusions. Anecdotes can contain personal bias, might be carefully selected to represent a particular point of view, and, in general, may be completely different from the general trend.

### Lesson:



1. Prepare your video player to show the first 5 minutes and 23 seconds of Hans Rosling's video *How Not to Be Ignorant About the World* found at: [https://www.ted.com/talks/hans\\_and\\_ola\\_rosling\\_how\\_not\\_to\\_be\\_ignorant\\_about\\_the\\_world](https://www.ted.com/talks/hans_and_ola_rosling_how_not_to_be_ignorant_about_the_world)

2. Ask students to play along as they watch the video. Each time Hans Rosling asks the audience to choose an answer to each of the three questions, pause the video for about 5 seconds and ask students to write down what they think is the answer to each question.



3. After viewing the video, engage students in T-I-P-S (see strategies) with the questions below:
  - a. Why did the chimps at the zoo score better than the people? *Answer: Anyone can select the correct answer just by chance.*
  - b. On the second question, Hans Rosling says that "everyone is aware that there are countries and there are areas where girls have great difficulties and they are stopped when they go to school." How could this information influence the answer choice? *Answer: Personal knowledge and experiences can influence what we think we know.*
  - c. Why do you think only a few people know the correct answer to these three questions? *Answer: People do not know enough about the data that can help them answer these questions.*
4. Display the following statements to students:

- "My skin glows more...I feel pretty confident." Melissa for Proactiv®
- "Within four months, I'd lost a grand total of 63 pounds\* and was down to my goal weight." Marianne G. for Nutrisystem®
- "The customer service is obnoxious. The employees are patronizing, smug, and intractable." Seymour773 for Bank of America®



5. Discuss each statement with students by asking the following questions:
  - a. Is \_\_\_\_\_ a good product? *For example, is Proactiv® a good skin product?, is Nutrisystem® a good diet program?, is Bank of America® a good bank?*
  - b. Do you think this person's experience is "typical"? Why? *Answer: Maybe it is typical but maybe not. Their own experience might be very different.*
  - c. Do you think the company chose this person? How do you know? *Answer: Each company may have chosen the first 2 statements because they were a success. In the case of the Seymour773, a competing company may have chosen his experience to make them appear better.*

- d. What about all the other people? How many were successes, how many failures? *Answer: We don't know for sure.*
  - e. How could we answer such questions? *Answer: Collect data!*
6. Inform students that the statements are called testimonials and they are examples of **anecdotes**. Anecdotes are stories that someone tells about his/her own experience or the experience of someone he/she knows. Anecdotes are good for some things like witness statements in a police report but are not useful for reaching conclusions about groups of people because the assumptions they are based on are not always true. Their claims are easily debunked. Many anecdotes do not equal data.

**Note to teacher about witness statements:** A lot of evidence suggests that witness testimony needs to be examined very closely. "As perhaps the single most effective method of proving the elements of a crime, eyewitness testimony has been vital to the trial process for centuries. However, the reliability of eyewitness testimony has recently come into question with the work of organizations such as The Innocence Project, which works to exonerate the wrongfully convicted. This thesis examines previous experiments concerning eyewitness testimony as well as court cases in which eyewitnesses provided vital evidence in order to determine the reliability of eyewitness testimony as well as to determine mitigating or exacerbating factors contributing to a lack of reliability." Information gathered from [digitalcommons.liberty.edu](http://digitalcommons.liberty.edu)

- 7. On the other hand, **data** are a series of observations, measurements, or facts. Data are information and tell a story.
- 8. Quickly survey students about whether the video they watched at the beginning of class is based on anecdotes or data. Then inform students that they will analyze two articles to find out if their claims are based on anecdotes or data.
- 9. Students will read one of two articles, *Miracle at the KK Cafe* (LMR\_Miracle\_Cafe) or *Can Trophy Hunting Actually Help Conservation?* (LMR\_Trophy\_Hunting\_Conversation) to analyze whether the claims each makes are based on anecdotes or data.



## SF WEEKLY

### News

NEWS • BAY VIEW

Like 19 Tweet Save

### Miracle at the KK Cafe

Peanut milk can strengthen AIDS patients, heal gums, soothe chronic skin diseases, and prevent baldness, or so it is said. And it's found only one place on Earth.

By **Matt Smith**  
Wednesday, May 8 2002

Comments 1

Its creators call it peanut milk. It tastes, unfortunately, exactly like its name. It is made from a secret recipe of just four ingredients: peanuts, rice, sugar syrup, and water. It contains 12 grams of fat per 8-ounce serving. It costs \$4.50 a



LMR\_Miracle\_Cafe

## The Last Word on Wildlife

21st Century Government Affairs Solutions for the Wildlife Sector

### CAN TROPHY HUNTING ACTUALLY HELP CONSERVATION?

January 21, 2014 by Andrew Wyatt

(<http://conservationmagazine.org/2014/01/can-trophy-hunting-reconciled-conservation/>)

CAN TROPHY HUNTING ACTUALLY HELP CONSERVATION? January 15, 2014  
(<http://conservationmagazine.org/conservation-science-news/>)

Re-blogged from Conservation Magazine: <http://conservationmagazine.org/2014/01/can-trophy-hunting-reconciled-conservation/comment-page-1/#comment-22621>  
(<http://conservationmagazine.org/2014/01/can-trophy-hunting-reconciled-conservation/comment-page-1/#comment-22621>)

Can trophy hunting ever be a useful tool in the conservationist's toolbox? On the surface, the answer would appear obvious. It seems as if the killing of an animal – especially an endangered one – for sport is directly contradictory to the goal of ensuring the survival of a species. The question has been

LMR\_Trophy\_Hunting\_Conversation

**Note:** These articles, although published in 2002 and 2014, effectively demonstrate the difference in using anecdotes versus data to support claims. Please don't hesitate to utilize any additional articles that are more up-to-date or pertain to subjects that are particularly aligned with your student interests.

10. Ask students to number themselves off as 1 or 2. Students whose number is 1 will read *Miracle at KK Café* and those whose number is 2 will read *Can Trophy Hunting Actually Help Conservation?*
11. Ask students to find a partner with the same number.
12. Before reading, ask students what they think their article will be about. Have students pair-share their thoughts.
13. During reading, students will take turns reading each paragraph out loud to each other. The student listening will verbally summarize what his/her partner just read.
14. After reading, student pairs will answer the following questions in their DS Journals:
  - a. What was the article about?
  - b. What claim(s) was/were this article making? Cite examples from the article.
  - c. Was this article based on anecdotes or data? Cite examples from the article.
  - d. How believable are the claims?
15. After answering the questions, students will find a partner with a different number. Each student will report to their new partner the following information about the article he/she read:
  - a. The name and publisher of the article.
  - b. His/her response to the four questions in their DS journal.
16. Quickly survey students about which article was based more on anecdotes and which one was based more on data. Ask a couple of students to explain their choices and give examples.

*Miracle at KK Café makes claims that are anecdotal. Students may cite a customer's claim as an example of an anecdote. Can Trophy Hunting Actually Help Conservation? uses data to make their claims. Students may refer to a statistic used in the article as an example.*
17. Class discussion: **Data Beat Anecdotes!** Ask students to come up with reasons why this statement is true. Have them come up with situations where you **have** to have an anecdote. For example, if asked what it's like to walk on the moon, only a few people would be able to tell us.



#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students will do a *Last Word Review* for the words DATA and ANECDOTE.

*Last Word Review:* Write the word vertically. Students come up with a word or phrase for each letter of the word. Each letter of the word should summarize something about what the students learned about the topic.

## Lesson 2: What is an Experiment?

### Objective:

Students will learn about the elements of an experiment and the meaning of “causation”. Students will learn to distinguish claims of causation from claims of association.

### Materials:

1. Video: MythBusters' Are Hands-Free Calls Really Safer While Driving?  
<https://www.youtube.com/watch?v=C-RAKWdKDEk>

**Note:** If video is not found using link, please use a search engine (e.g., Google Video) and type “MythBusters Are Hands-Free Calls Really Safer While Driving” to find it. The clip is a little over 9 minutes in length.

### Vocabulary:

experiment, research question, subjects, treatment, treatment group, control group, random assignment, outcome, statistic

**Essential Concepts:** Science is often concerned with the question “What causes things to happen?” To answer this, controlled experiments are required. Controlled experiments have several key features: (1) there is a treatment variable and a response variable, and we wish to see if the treatment causes a change that we can measure with the response variable; (2) There is a comparison/control group; (3) Subjects are assigned randomly to treatment or control (randomized assignment); (4) Subjects are not aware of which group they are in (a ‘blind’). This may require the use of a placebo for those in the control group; and (5) those who measure the response variable do not know which group the subjects were in (if both 4 and 5 are satisfied, this is a ‘double blind’ experiment).

### Lesson:

1. Display the following headlines to students:
  - a. Stop Global Warming: Become a Pirate
  - b. Lack of sleep may shrink your brain
  - c. Early language skills reduce preschool tantrums
  - d. Dogs walked by men are more aggressive
2. Discuss each headline by asking the following questions:



- a. What is the headline implying with its wording? *Answer: 1a is implying that you can stop global warming by becoming a pirate, 1b is implying that it's possible to shrink your brain if you aren't getting enough sleep, 1c is implying that having early language skills will decrease preschool tantrums, 1d is implying that dogs are more aggressive when they've been walked by men.*
- b. Is it implying causation or association? *Answer: Discuss definitions of causation and association. Causation means there is a cause-and-effect relationship between variables. For example, heat causes water to boil; whereas association or correlation means that high values of one variable tend to be associated with high values of the other (or high values tend to be with low values). However, this is not necessarily cause-and-effect at play. For example, blanket sales in Canada are associated with brush fires in Australia - not because Canadian blankets cause the fires, but because Canadian winters cause blanket sales, and Canadian winters are Australian summers, which cause fires. 1a, 1c and 1d are implying causation and 1b is implying association.*
- c. How can you tell the difference between causation and correlation? What words stand out in these headlines? *Answers will vary but some terms for causation include: cause, increase/decrease, benefits, impacts, effect/ affect, etc.; and for correlation include: get, have, linked, more/ less, tied, connected, etc. In 1a, “become” stands out; in 1b, “may” stands out; in 1c, “reduce” stands out; in 1d, “are” stands out.*

**Note:** Answers may not be clear-cut in all situations. The purpose here is to get students thinking about “association” and “cause and effect” mean in the real-world and understand that association does not imply cause and effect.

- d. Change each causal version of a headline into a non-causal version and vice versa. *Answers will vary but an example for 1a is to instead say Global Warming linked to increase of pirates.*
3. Introduce the MythBusters video clip by answering the following questions, in teams, for their headline “Are Hands-Free Calls Really Safer While Driving?”
  - a. What is the headline implying with its wording? *Answer: That hands-free calls might not be safer than holding a phone while driving.*
  - b. Is it implying causation or correlation? How do you know? *Answer: Causation because “really safer” implies that driving while on a hands-free call causes you to not drive safely.*
  - c. How can we determine if this is true? *Answer: Split the class into groups and have each team come up with a way to determine if this is true. Each group should assume that they get to examine 30 people.*
-  4. Show the MythBusters video clip called Are Hands-Free Calls Really Safer While Driving? The clip can be found at:  
<https://www.youtube.com/watch?v=C-RAKWdKDEk>
5. Focus students on the following guiding questions and ask them to take notes as they watch the video clip:
  - a. How did the MythBusters design the investigation? *Answer: They conducted a driving test. In this case, one group is driving while on a hands-on call (because in the eyes of the law, this is no different than driving without being on a call). The treatment group is driving while on a hands-free call. They determined that “safe” meant passing the driving course by following the turn-by-turn navigation correctly and having no accidents. Passing the driving course was considered a “success.”*
  - b. What steps did they take? *Answer: They split the 30 subjects into two groups, 15 in the one group and 15 in the treatment group, and had them drive a car in a simulator.*
  - c. How is this different than your team’s headline responses in step 3 from above? Did the Mythbusters investigation match what you thought it would be about? *Answers will vary depending on students’ assumptions about the investigation prior to watching the video.*
6. After viewing the clip, inform students that the MythBusters have just conducted an **experiment**, which is one method of data collection.
7. We begin with a brief introduction into “what is an experiment” but the definition will be developed over the next several lessons.
8. Guide students to identify the elements of an experiment by referring back to the video clip:
  - a. **Research Question**—the question to be answered by the experiment  
*Are Hands-Free Calls Really Safer While Driving?*
  - b. **Subjects** – people or objects that are participating in the experiment  
*The 30 adults.*
  - c. **Treatment** – the variable that is deliberately manipulated to investigate its influence on the outcome; this is sometimes known as the explanatory, or independent, variable  
*Having drivers on a call without a cellphone in their hand. Using the research question as our guide, the question implies that we are experimenting with the hands-free calling, therefore that is the treatment as we are comparing it to hands-on calling.*
  - d. **Treatment group** – the group of subjects that receive the treatment  
*The 15 people who used the hands-free phone.*
  - e. **Control group** – the group that does not receive a treatment  
*The remaining 15 people who used the hands-on phone.*

- f. **Random assignment** – subjects are randomly assigned to either the treatment or control group  
*It's unclear if random assignment was used in this video (if it was, we were not told so), but we're going to assume that people were randomly assigned to each group.*
- g. **Outcome** – the variable that the treatment is meant to influence; this is sometimes known as the response, or dependent, variable  
*Success – whether or not a person was a safe driver by passing the test.*
- h. **Statistic** – a number such as a mean or proportion used to summarize our findings for the control and treatment groups  
*In this case, the MythBusters gave raw data. A statistic would have been that 6.7% (1 out of 15) of the control group passed the test. Similarly, 6.7% of the treatment group passed the test.*

**Note:** In this experiment, and in those found in the IDS curriculum, we use a treatment and a control group. However, a control group is not a *necessary* element of an experiment. Sometimes it is more appropriate to have two treatment groups with no control group (e.g., medical professionals testing different doses of drugs). The effect that is being studied will dictate whether to feature a control group or not.



9. Display the following questions on the board or projector. Using T-I-P-S, ask students to discuss them.
  - a. Why did the MythBusters follow all of these steps to design their experiment? *Answer: In order to determine if driving while on a hands-free call was really safer than driving while holding a cellphone on a call.*
  - b. We don't know how MythBusters chose who would be in the treatment group and who would be in the other group. Suppose that less-experienced drivers were assigned to one group and the more-experienced drivers ended up in the treatment group. Would you believe in the conclusions? *Answer: No, because less-experienced drivers would probably be more prone to mistakes (because they're inexperienced) than more-experienced drivers. Explain that this – another explanation for the cause-and-effect – is caused by a confounding variable.*
  - c. Explain that in order to make the two groups as similar as possible, experimenters usually assign subjects randomly. How might we randomly assign about half of the subjects to the treatment and half to the control? *Answer: We might flip a coin, and those who get Heads go to Treatment.*
  - d. Why would random assignment improve the MythBusters study? *Answer: Because then the two groups would be more similar. So we wouldn't have a confounding variable to worry about.*
10. **Emphasize that without random assignment, we cannot determine causation because we are not comparing two similar groups.**

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Lesson 3: Let's Try an Experiment!

#### Objective:

Students will explore the importance of randomized assignment in experiments. They will understand that without random assignment, there might be confounding variables and will be able to suggest possible confounding variables.

#### Materials:

1. Measuring Tape

#### Vocabulary:

confounding factors

**Essential Concepts:** Randomized assignment is required to determine cause-and-effect.

#### Lesson:

1. Inform students that they will be exploring the question “Why do we need randomized assignment?” by conducting an experiment. Tell students that you have a treatment that can make people taller. Explain that the class will be divided into two groups, one group will get the treatment, and one group will not. The group that does not receive the treatment will be the control group. After the treatment, they will measure the groups to see which is taller. Now divide the class into two groups by placing the taller students in the treatment group and the shorter students in the control group.

Remember that in an experiment we typically have a treatment group and a control group. Some experiments do not require a before and after as long as the two groups being compared are similar (many medical studies use this method for experiments). In this case, we will run the experiment and *then* compare average height of the treatment group to the control group.

2. Tell them that after the treatment group takes the treatment, your statistic to compare groups will be to measure the heights. If the treatment group is taller, then the treatment must have worked. There are two possible outcomes to dividing the class this way:
  - a. The students will protest (as they should) and you can start a discussion as to why this is not a good way to divide the class.
  - b. OR the students don't protest and you continue with the experiment. The treatment should be something silly, like waving a ruler in front of the person's face or by asking them to chant “grow, grow, grow!” three times. After treatment, measure the heights of each group and ask them if they think this is good evidence (**do not say “proves”**) that the treatment is effective.
3. Regardless of the outcome, students should recognize that by putting the taller students in one group, the outcome was pre-determined, since they were taller to begin with due to genetics or other factors. This is an example of a **confounding factor**. Confounding factors are variables that provide an alternative explanation of the effect of the treatment on the outcome variable.
4. Ask students: “How should students be put into groups?”
5. Discuss various other methods of grouping students. *Students should be able to recognize that you shouldn't use any characteristics to decide the groups.*
6. Continue discussion of other ways to decide the groups. Use the following questions as a guide:
  - a. What about flipping a coin?
  - b. What will the height balance look like? *Answer: Each group should have about the same balance as the class, though not exactly.*
  - c. Why is it important that the groups be similar? *Answer: Because otherwise, something else might be the cause of the response changing.*



7. Inform students that today the class will begin to design their own experiment using what they have learned over the last few lessons. The question they will investigate is:

**How does our perception of time change when exposed to a stimulus?**

8. They will be trying to determine the length of one minute without the use of time-aids. In their experiment, they will subject some students to a stimulus and others to no stimulus. They will then analyze the data to determine if subjecting students to a stimulus affects the perception of how long a minute of time lasts.
9. In their DS journals, ask students to answer the following questions about the elements of their experiment:
  - a. What is the research question we're interested in addressing?
  - b. Who are the subjects that will be participating in the experiment?
  - c. How should we randomly assign the subjects into treatment and control groups? (See step 12 below for an RStudio method that the teacher can use)
  - d. What is the outcome variable that we will be measuring? What unit of measurement should we use?

**Note:** Students will decide on a treatment to apply to each group on the following day.



10. As a class, discuss the responses to the questions above (step #9, a-d) and come to a consensus for each question's answer.
11. Inform the class that they will be using the answers they have agreed upon as the final design of the class's experiment.
12. At the end of the class, the students should be assigned to the treatment or control groups using the randomization method they chose as a class in step #9c from above.

**Note:** One method to determine group assignment would be to use the class roster and the **sample()** function in RStudio. The students have a number that corresponds to their placement on the roster (i.e. student 1's last name most likely starts with an A, and then we move alphabetically through the roster). You can then use RStudio to randomly select which half of the numbers/students will be assigned to the treatment group.

```
> sample(1:30, size = 15, replace = FALSE)
```

13. Students will conduct the experiment in the next lesson.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 4: Predictions, Predictions

### Objective:

Students will continue to read articles critically. They will anticipate visualizations about the data that will be collected from the class experiment and make predictions about the outcome.

### Materials:

1. Article: PsyBlog's *Time Perception: How Your Brain Experiences Time* found at: <https://www.spring.org.uk/2025/01/time-perception-brain.php>
2. *Experiment Predictions* handout (LMR\_ U3\_L4)

### Vocabulary:

theory

**Essential Concepts:** Designing an experiment requires many decisions, including what to measure and how to measure it.

### Lesson:



1. Students will read the article *Time Perception: How Your Brain Experiences Time* found at: <https://www.spring.org.uk/2025/01/time-perception-brain.php>.
2. They will read the article critically to answer the following questions (displayed or written on the board):
  - a. Who was observed and what were the variables measured? *Answer: People and their perceptions of time.*
  - b. What statistical questions were the researchers trying to answer? *Answer: How is time perception affected by different stimuli or factors?*
  - c. Who collected the data? *Answer: Researchers collected the data.*
  - d. How were the data collected? *Answer: Case studies, behavioural experiments, and advanced imaging technologies are mentioned as methods of data collection.*
  - e. What claim(s) did the article make? *Answer: The article claimed that there are multiple factors involved in time perception such as psychological factors, physiological factors, and environmental factors.*
  - f. What statistics could be found by the ideas presented in this article? *Answers may vary. A possible answer might be to ask for the statistic of people who report a change in time perception as they get older.*



3. In their teams, ask students to share their responses from reading the *Time Perception: How Your Brain Experiences Time* article and agree on the responses as a team.
4. Do a quick *Whip Around* of the responses (see step #2 from above for responses).
5. Remind students that they designed a class experiment during the previous lesson but did not select an actual treatment. As a class, decide on a treatment to use for the experiment. Students can use the methods found in the article for inspiration or come up with something novel on their own.

**Note:** Stimuli examples include music (genres determined by the class), lights off, physical activity (e.g., holding arms out), relaxation/meditation techniques, heads down, eyes closed, etc. Ensure that the experiment can be completed in one 50–60-minute class period. Treatments requiring excessive preparation time (e.g., running a mile) are less than ideal.

6. Before they conduct the experiment, students will test their theories by making predictions about the data and the outcomes. A **theory** is an idea used to explain a situation.
7. Display the class experiment's research question:

**How does our perception of time change when exposed to a stimulus?**

8. Take a poll of the students who believe that there will be differences in the estimate of the length of a minute between the treatment and control groups. The remaining students, then, do not believe that there will be differences.
9. Then, ask those students who believe there are differences, how small or large they think the difference will be.



10. Distribute the *Experiment Predictions* handout (LMR\_U3\_L4) and, in pairs, have students discuss and complete the answers for the handout.

**Note:** What will the distribution of time perceptions look like? *Answer: The distributions will likely have more points that are closer to 60 seconds but will also have values that are shorter and longer than 60 seconds. Appropriate plots to use will include histograms, dotplots or boxplots.*

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Experiment Predictions**

**Instructions:**  
Answer the following before conducting your time perception experiment. Remember, the variable that we're measuring is the number of seconds that actually elapse until each person believes one minute has passed.

1. In the boxes below, draw a plot of what you predict the distribution of each group's data will look like. Be sure to add numbers and labels.

| Treatment |
|-----------|
|           |
| Control   |
|           |

LMR\_U3\_L4

11. Using *Anonymous Author*, select student work to share with the whole class.
12. Give student teams time (about 2 minutes) to discuss each product that is shared/presented.
13. Teams will offer their thoughts using a modified *Two Cents* strategy where, instead of two cents, each team will receive one cent (or a token) and, in order to turn it in, the team will have to make comments or ask questions about the student work that is being shared. Call on teams until you have collected every cent. This ensures that all teams contribute to the discussion.
14. Inform students that they will conduct the experiment in which they will estimate the length of time of one minute during the next lesson.



#### **Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## **Lesson 5: Time Perception Experiment**

### **Objective:**

Students will engage in a collectively designed experiment.

### **Materials:**

1. RStudio's **stopwatch()** function
2. IDS ThinkData Ed App or Browser-Based survey-taking tool

**Essential Concepts:** Designing and carrying out an experiment helps us answer specific statistical questions of interest.

### **Lesson:**

1. Begin the lesson by eliciting the elements of an experiment from students (they may refer back to their DS journals for their responses from Lesson 2).
2. Inform students that they will be using RStudio to get a precise measurement of their estimate. Ask for a student volunteer.
3. Demonstrate the stopwatch function using RStudio by typing in the following code:  

```
> stopwatch()
```
4. Then, ask the student volunteer to stand in front of your computer and get ready to estimate the length of time of one minute without looking at a clock. Once he/she thinks a minute has passed, ask him/her to press the enter/return key on the keyboard to see the result of the estimate.
5. Inform students that you have just demonstrated how they will measure their one-minute estimates.
6. Begin conducting the experiment by reviewing the research question:  
**How does our perception of time change when exposed to a stimulus?**
7. Refer back to the experiment design. Review the specific treatment that the subjects in the treatment group will receive. If necessary, demonstrate to the treatment group how to do the experiment. For example, if standing with open arms is the stimulus, the estimate begins when the student starts the **stopwatch()** function and engages in the stimulus, and ends when the subject presses enter/return in RStudio to stop the timer.
8. For the control group, the students can simply sit at their desks with their eyes closed. Each student will run the **stopwatch()** function and stop the timer when they believe a minute has elapsed.
9. Conduct the experiment in its entirety. Use team roles to ensure the experiment is done correctly.
10. Have each student use a computer and the **stopwatch()** function to record her/his estimate of one minute. Ensure each student records her/his estimate in the DS journal.
11. When the experiment is completed, have students enter their data in the *Time Perception* survey found in the Survey Taking Tool at <https://portal.thinkdataed.org> or by using the IDS ThinkData Ed App in their iOS or Android devices.
12. Inform students that they will be analyzing their experiment results in *Lab 3A: The results are in!*



**Note to Teacher:** As a check for understanding, engage the class in a discussion about experiments. Use the following questions as a guide to assess student understanding.

- 1) When is random assignment used? Why is it important? *Answer: Random assignment is used when you wish to determine whether a treatment causes changes in an outcome variable. It's important because it creates a "balance" of the groups so that the only way the groups differ, on average, is that one gets the treatment and one does not. Thus, if there is a change in the outcome variable, only the treatment could have caused it.*

- 2) Below are some headlines. Determine which ones are causal. If not causal, re-write so that it is. If causal, state why it's causal.
- Straight A's in high school may mean better health later in life. *Answer: Not causal, re-writing answers will vary.*
  - Murder rates affect IQ test scores: Study. *Answer: Causal, explanations will vary.*
  - Microbe linked to Alzheimer's Disease. *Answer: Not causal, re-writing answers will vary.*
  - Luckiest people "born in summer." *Answer: Causal, explanations will vary.*
- 3) Why is a control group important? *Answer: The control group is important because it allows us to measure the effects of the treatment group with an untreated comparable group. Without the control group, we don't know what would have happened if we had done nothing. For example, think of a new vaccine for the flu. If there is no control group, and we see the treatment group improving, we will never know if they would have improved anyways, without the vaccine.*

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

**Next Day**

**Lab 3A: The results are in!**

Complete Lab 3A prior to Practicum.

### **Lab 3A: The results are in!**

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

#### **Conducting experiments**

- Previously in class, you conducted an experiment to gauge how a stimulus affected people's perception of time.
  - Some people were given a treatment, others were not.
- In this lab, we'll use the data cycle to analyze the *research question*:  
*Does the stimulus your class chose change people's perception of time?*

#### **Coming up with questions**

- **(1) Write down two statistical investigative questions that will help you answer the *research question* from the previous slide.**
- **Then, export, upload, import your experiment data into RStudio.**
  - If you're having trouble coming up with good statistical investigative questions, try loading the data and looking at the variables.
  - Ask yourself, *How would I use these variables to answer the research question?*

#### **Analyzing our data**

- **Create appropriate plots to answer your statistical investigative questions.**
- **Calculate appropriate numerical summaries to answer your statistical investigative questions.**
- **(2) Write down a few sentences interpreting your plots and summaries.**

#### **Wrapping it up**

- Is it possible your initial results occurred by chance alone?
  - **(3) Write and run code using repeated shuffling to determine how likely the typical difference between the two groups occurred by chance alone.**
  - **(4) Create a plot and use it to justify your answer.**
- What do you conclude about the *research question*?
  - **(5) Write a report using the plots and analysis you conducted to answer the *research question*.**
  - Be sure to describe how you conducted your experiment.

## **Practicum: TB or Not TB?**

### **Objective:**

Students will apply what they have learned about experiments and shuffling.

### **Materials:**

1. Computers
2. *Practicum: TB or Not TB* (LMR\_U3\_Practicum\_TB\_or\_Not\_TB)

### **Practicum TB or Not TB**

Experiments in the medical field that involve new treatments (new medications) are called clinical trials. You have received a dataset that shows the results from Sir Austin Bradford Hill's first randomized study in 1948 examining the effects of the antibiotic Streptomycin on 107 tuberculosis patients. You and a partner will use this dataset to find out if Streptomycin is an effective treatment for tuberculosis.



A short article about tuberculosis facts can be found at:

<https://www.cdc.gov/tb/communication-resources/tuberculosis-fact-sheet.html>

Since this is an experiment, answer the following questions below. You may need to research the answer to some of the questions.

- a. What is the research question?
- b. Who are the subjects that participated in the experiment?
- c. What is the treatment?
- d. Who is in the treatment group?
- e. Who is in the control group?
- f. How were the subjects assigned to each group?
- g. What population is this experiment representative of?
- h. What is the variable that we will be measuring?
- i. What is the outcome of this experiment?

To answer your research question, you and a partner will compare the outcome of the data with the outcomes given by a chance model (in which Streptomycin has no effect on TB).

1. First, scrape the data. Refer to the web scraping lab if you need to recall how to scrape data. To access Sir Hill's data, go to: <https://labs.thinkdataed.org/extras/webdata/tb.html>
2. Second, determine the percentages of subjects in the study that died and the percentages of the subjects that recovered for each group.
3. Third, assuming that the treatment had no effect, use the data to:
  - a. Calculate the percentage of people with tuberculosis we would expect to die.
  - b. Use the *expected* percentage for (a), above, to calculate the number of people we expect to die from the treatment group.
  - c. Compare the outcome from (b), the number of people we expected to die, to the number of people from the treatment group that *actually* died.
4. Then, if we assume that the outcome does not depend on the treatment, design and complete an appropriate simulation in RStudio using a chance model to replicate Sir Hill's study:
  - a. Shuffle the treatment and control labels 300 times; each time, calculate the percentage of treatment patients who "died". Plot the distribution of the 300 percentages. Refer to the simulation labs if you need to recall how to create a simulation.
  - b. Use the results from the chance model (shuffling) to determine whether (i.) or (ii.) below is the most reasonable explanation for the actual data in Sir Hill's study and state why:

- i. Streptomycin is a much better treatment for tuberculosis than bed rest. So, the outcome depends on the treatment.
- ii. The actual difference between treatments is due to chance; Streptomycin may not be effective on tuberculosis. So, it is possible that treatment and outcome are independent.

5. Can we say that Streptomycin **causes** the recovery of tuberculosis patients? Explain your answer.

Create a 4-5 slide, 5-minute presentation that shows your results. Be sure to include a detailed explanation of how you and your partner decided to conduct your simulation. Each person must participate in the presentation. In addition to the presentation, submit a 2-4-page, double-spaced summary of your analysis.

# Would You Look at That?

Instructional Days: 4

## Enduring Understandings

An observational study is a data collection method in which subjects are observed and outcomes are recorded. Unlike experiments, it may not be possible to assign subjects to treatment and control groups in observational studies, which impedes our ability to control for confounding factors. This means that researchers must rely on existing control and treatment groups to observe the outcomes. Observational studies can show associations in the data, but cause and effect relationships can only be concluded with experiments.

## Engagement

Students will participate in the *Observational Studies Activity* described in Lesson 5. They will record information that can be obtained through pictures. The data will then be analyzed to see if there are any variables related to the number of friends a person has on social media.

## Learning Objectives

### Statistical/Mathematical:

- S-IC 1: Understand statistics as a process for making inferences about population parameters based on a random sample from that population.
- S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.
- S-IC 6: Evaluate reports based on data.

### Data Science:

Understand that data from observational studies can help us find associations among variables. Explain why some variables that are not related in reality might look as though they are due to the presence of confounding factors.

### Applied Computational Thinking using RStudio:

- Download data from the Internet that was collected via an observational study.
- Clean dataset by adding variable names.
- Create scatterplots of two variables and determine possible relationships between them, as well as identify potential confounding variables.

### Real-World Connections:

Economists, psychologists, and biologists conduct observational studies to study human behavior. For example, observational studies are used in epidemiology to study outbreaks of illnesses and people's behavioral patterns.

## Language Objectives

1. Students will connect the data they gather through Participatory Sensing campaigns to observational studies.
2. Students will compare and contrast observational studies with experiments.
3. Students will describe the relationship between causation and association.

## Data File or Data Collection Method

### Data File:

1. *Lung Capacity of Children* dataset found at <https://raw.githubusercontent.com/thinkdataed/dataset/main/fev.csv>  
**Note:** The raw dataset can be found at: <https://jse.amstat.org/datasets/fev.dat.txt>  
**Note:** The documentation can be found at: <https://jse.amstat.org/v13n2/datasets.kahn.html>

Data Collection Method:

Students will record information about a set of high school students by observing characteristics given in a picture.

**Legend for Activity Icons**



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 6: Observational Studies

### Objective:

Students will learn that an observational study is a data collection method in which subjects are observed and outcomes are recorded. They will learn how to collect this type of data and make informal inferences about the results.

### Materials:

1. *Stick Figures Cutouts* (LMR\_U1\_L2) from Unit 1, Lesson 2  
**Advance preparation required:** Needs to be cut into cards (see step 1 in lesson)
2. *Turning Observations into Data* handout (LMR\_U3\_L6)

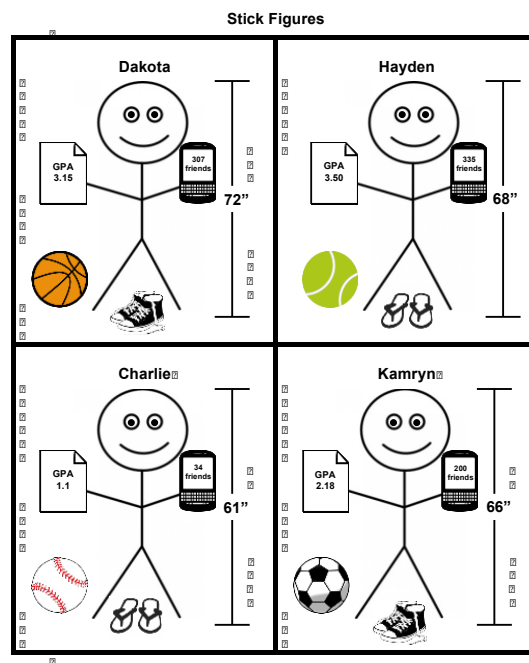
### Vocabulary:

observational study

**Essential Concepts:** Observational studies are those for which there is no intervention applied by researchers.

### Lesson:

1. From Unit 1, Lesson 2, redistribute one full set of 8 cards from the *Stick Figures* handout (LMR\_U1\_L2) to each student team.  
**Advance preparation required:** Print the *Stick Figures* handout (LMR\_U1\_L2). The handout can then be cut into the 8 cards. You will need enough sets of the cards for each student team to share a full set. For example, if there are 5 student teams in a class, then 5 copies of the file will need to be printed so that each team gets all 8 cards.
2. Have students recall that they used these cards in Unit 1, Lesson 2. When they used them in Lesson 2, the data was collected, recorded, and organized, but without particular structure to it.



LMR\_U1\_L2

- Then, distribute one copy per student of the *Turning Observations into Data* handout (LMR\_U3\_L6).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Turning Observations into Data**

**Part 1**

Name the 6 variables that can be recorded from the picture on your card. What variable names could you use to represent these? Record your answers in the correct column below.

| Numerical Variables | Categorical Variables |
|---------------------|-----------------------|
|                     |                       |
|                     |                       |
|                     |                       |

**Part 2**

Using the variable names you chose in Part 1, create a data table and use the first row to record the information about the person on your card.

Compare your card's values to your team members' and add their data to your table.  
If there are extra cards left over, record those observations as well.

|  |  |  |  |  |  |
|--|--|--|--|--|--|
|  |  |  |  |  |  |
|  |  |  |  |  |  |

LMR\_U3\_L6

- Every student from the team will then select one of the cards from the team's pile of 8 and should begin working through the *Turning Observations into Data* handout individually.
- As the students finish each part of the handout, they should compare their responses with their student teams.
- Go over the names of the variables in Part 1 by doing a quick Whip Around by teams. Then, select a couple of teams to share the information on the first row and one of the columns.
- Part 3 of the handout asks the students to consider the following research question:

**What determines the number of friends a person has on social media?**



- Once the students have completed the handout, discuss the variable that they thought was best associated with the number of friends on social media. *They should have seen that a person's GPA was related to the number of friends. More specifically, the higher a person's GPA, the more friends he/she had.*
- Ask a few students to share out their responses to the very last question: "Can you think of another variable (not necessarily given in the pictures) that might impact both the number of friends AND the variable you selected? Give an example and explain how it might impact each of the variables." *Answers will vary, but one example could be: a person's self-esteem level (if he/she is confident in school, his/her grades might be higher; higher confidence could also be a reason for a person having more friends).*



- Remind students that in the previous section, they learned about the elements of an experiment. In teams, ask students to discuss how collecting this data is similar or different from experiments. Then have a whole class discussion about this comparison, guiding students to realize that there were no assignments to groups and no treatment was applied. The subjects (i.e. the people displayed on the cards) were simply observed, and then information about them was recorded.
- Inform students that an **observational study** is a data collection method in which subjects are observed and outcomes are recorded. No treatment is applied to the subjects. Instead, researchers are simply watching something happen and have absolutely no control over it.
- In Lesson 7, students will learn more about the differences between experiments and observational studies and what conclusions they can make about each.

#### **Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 7: Observational Studies vs. Experiments

### Objective:

Students will learn how observational studies differ from experiments and will classify different research scenarios based on which method would be most appropriate. They will also learn about the roles of ethics, cost limitations, and feasibility when deciding between the two data collection methods.

### Materials:

1. *What Should We Do?* handout (LMR\_U3\_L7)

### Vocabulary:

ethics, cost limitations, feasibility

**Essential Concepts:** Experiments are not always possible because of various factors such as ethics, cost limitations, and feasibility.

### Lesson:

1. Remind students that in observational studies, we can never randomly assign subjects to treatment and control groups. Conversely, in experiments, we always need to have random assignment into these groups.
2. Pose the following question to students: “Why can’t we just always do experiments?” Have students discuss this question in their student teams and write down a few responses in their DS journals.
3. Inform students that a researcher wants to perform studies to answer the research questions below. In teams, have students come up with reasons for why an experiment would not be possible for each scenario.
  - a. Does smoking cause lung cancer? *Answer: Unethical. You cannot make people smoke cigarettes and then see if they have lung cancer later in life.*
  - b. Does drinking water from Mars keep you healthier than drinking water from Earth? *Answer: Cost. It would be incredibly expensive to design a space shuttle that can successfully transport people to Mars and have them live there for an extended period of time and most researchers would not have the funding to do this.*
  - c. Do people with higher IQ scores score better on the SAT than people with lower IQ scores? *Answer: Not feasible/not possible. You cannot randomly assign IQ scores to people because it is a measurement based on aptitude.*
4. Select three teams and assign a scenario above to each team. Ask each team to report out on their assigned scenario. As teams share, be sure to discuss the following issues regarding why we cannot always to experiments:
  - a. **Ethics:** Sometimes, experiments cannot be performed because it would be unethical to give certain treatments to subjects. For example, we could not inject an HIV infection into participants because the long-term effects might lead to death.
  - b. **Cost Limitations:** Sometimes, experiments would be very costly and much too expensive to perform. Some possibilities could be with technology.
  - c. **Feasibility,** impossible to randomize: In certain cases, you cannot perform an experiment because it is impossible to randomly assign people to particular groups. For example, you cannot assign a gender to a person.
5. Distribute *What Should We Do?* handout (LMR\_U3\_L7). In teams, students will identify whether the research question could best be answered via an experiment or an observational study.



- Once all student teams have completed the handout, assign one research question to each team to report out. As each response is shared, conduct a whole-class discussion to compare which data collection method was most appropriate for each research question. Ensure everyone understands the reasons each method was chosen before moving on to the next scenario.

**Note:** Page 2 of the handout is an answer key for teacher reference only.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**What Should We Do?**  
*Deciding between Experiments and Observational Studies*

For each research question posted in column one below, decide which type of data collection method would be best, or would be most appropriate. Circle the appropriate method in column two. In column three, explain why you chose this method. Be sure to include why the other method would not be appropriate.

| Research Question                                                                                        | Best Data Collection Method             | Why This Method? |
|----------------------------------------------------------------------------------------------------------|-----------------------------------------|------------------|
| Do people who live alongside freeways suffer from asthma more than people who do not live near freeways? | Observational Study<br>OR<br>Experiment |                  |
| How does being alone or with others affect a person's stress or chill ratings?                           | Observational Study<br>OR<br>Experiment |                  |
|                                                                                                          | Observational Study                     |                  |

LMR\_U3\_L7



- Next, student teams will generate three research questions on their own. They need to identify the best data collection method for answering their question and should provide an explanation. At least one of the three research questions should use an observational study for data collection.
- Using a share-out strategy, have the reporter of each team share one of their research questions with the rest of the class.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 8: Monsters that Hide in Observational Studies

### Objective:

Students will learn about confounding factors that may impact the results of an observational study, which is why causation can never be concluded with observational studies, only associations between variables.

### Materials:

1. Computers
2. *Spurious Correlations* website (<http://tylervigen.com/spurious-correlations>)

### Vocabulary:

cause, confounding factors, associated

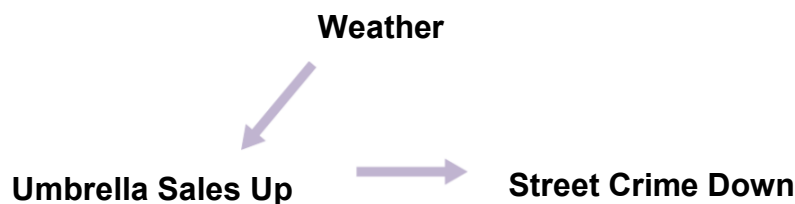
**Essential Concepts:** Confounding factors/variables make it difficult to determine a cause-and-effect relationship between two variables.

### Lesson:

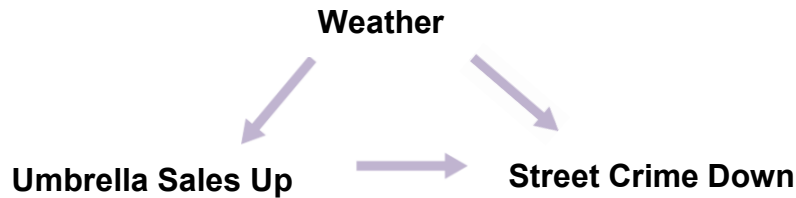
1. Ask students to recall that they looked at the relationship between a student's GPA and the number of friends that person has on social media during lesson 6. It seemed that students with higher GPAs had more friends than students with lower GPAs. But did this mean that the **cause** of a person's GPA is the number of friends they have? NO!
2. They also identified other variables that could have contributed to the relationship. These outside variables are called **confounding factors**. Remind students that we defined confounding factors in lesson 3 – confounding factors provide an alternative explanation of the effect of the treatment on the outcome variable. In this case, the variables are related to both the explanatory variable and the response variable in an observational study.
3. Propose the following statement to students: "Research suggests that a rise in umbrella sales leads to decreased street crime."
4. Allow the students to work in teams to think about possible confounding factors. They should choose a variable that is related to umbrella sales, and that might lead to decreased street crime. After they've come up with a few possibilities, use the following diagram progression to further explain the impact of confounding factors.
  - a. Step 1: Draw an arrow showing that "a rise in umbrella sales leads to decreased street crime" since that is what researchers have stated.

**Umbrella Sales Up**            **Street Crime Down**

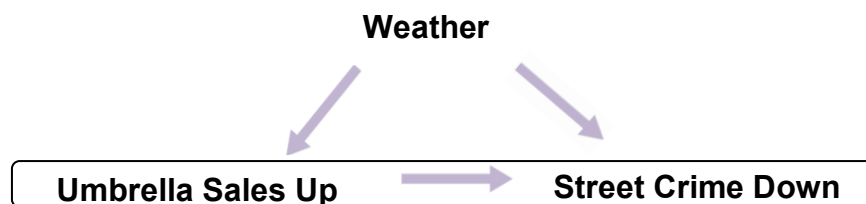
- b. Step 2: Include the variable that might be related to people buying more umbrellas (i.e., the confounding factor). For example, when the weather is rainy, people buy more umbrellas.



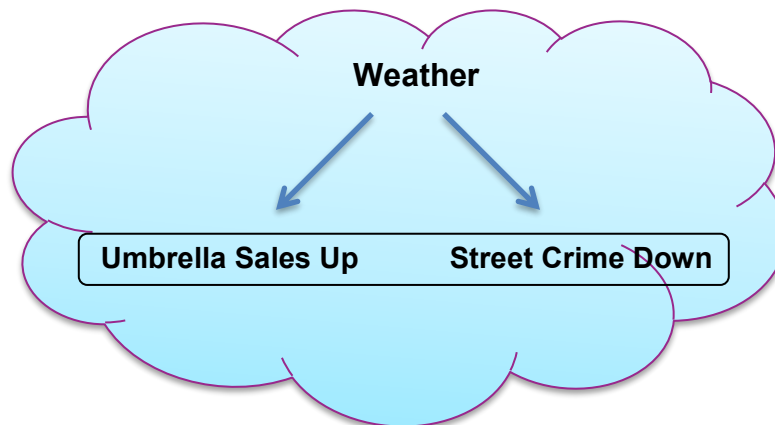
- c. Step 3: Draw an additional arrow from “Weather” to “Street Crime Down” because it is well known that when the weather is bad, people are less likely to be outside committing crimes.



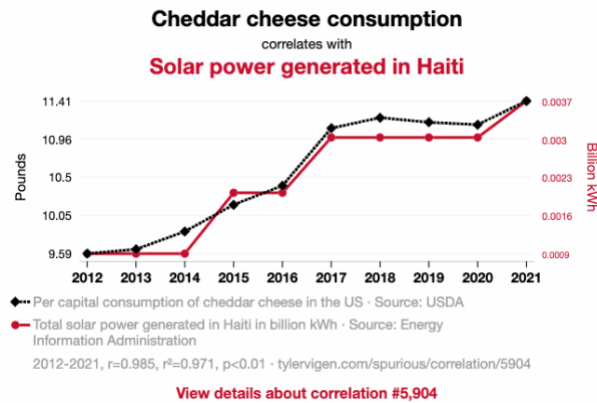
- d. Step 4: Remind students that the original claim was that “a rise in umbrella sales leads to decreased street crime.” However, we’ve now shown that maybe buying umbrellas is not the only thing that could be contributing to a decrease in street crime, which makes us question the link between the two variables.



- e. Step 5: Therefore, we have found a confounding factor with the variable “street crime”. This means we can erase the original “link” between a rise in umbrella sales and decreased street crime since there are outside variables interfering. We can’t say buying umbrellas *causes* decreased street crime, but we can say that a rise in umbrella sales is **associated** with decreased street crime. We can say that two variables are associated when the values of one variable relate to the values of the other in some way.



5. Once the students grasp what confounding factors are and how to identify them, introduce them to the website *Spurious Correlations* by Tyler Vigen. This site shows many explanatory and response variables that are randomly associated with each other. Spurious Correlations can be found at: <http://www.tylervigen.com/spurious-correlations>.



6. For the example given above, we see that as cheddar cheese consumption in the US increases, solar power generated in Haiti increases. Clearly, it does not make sense that if the US keeps consuming cheese, then Haiti will generate more solar power. It simply happened by chance (or a bizarre chain of confounding factors) that the two variables are related to each other.
7. Allow the students to explore the website on their own. They should choose a graph that interests them and answer the following questions in their DS journals:

**Note:** There are multiple pages of graphs, so they are not restricted to simply the homepage.

- a. What are the two variables shown in your graph?
- b. Is there a positive association or a negative association between the variables?
- c. Write an interpretation of this plot in the context of the data.
- d. Write the data points in a “spreadsheet format” in a form that RStudio could read. Each row should represent a point on the graph, and each column one of the two variables.

**Note:** You can get the details about the correlation by clicking on the “View details about correlation #...” link.

- e. By hand, make a scatterplot of the association. Describe whether the association seems strong or weak or moderate to you.
- f. Do you think that the explanatory variable *causes* the response variable? Explain.
- g. If you answered ‘no’ to f, then draw a diagram like in #4 with possible confounding factors.

**Note:** This can be difficult, depending on the graph chosen. Some factors to consider are weather, economy, fashion trends.

8. Example answers to step 7 from above are given below:



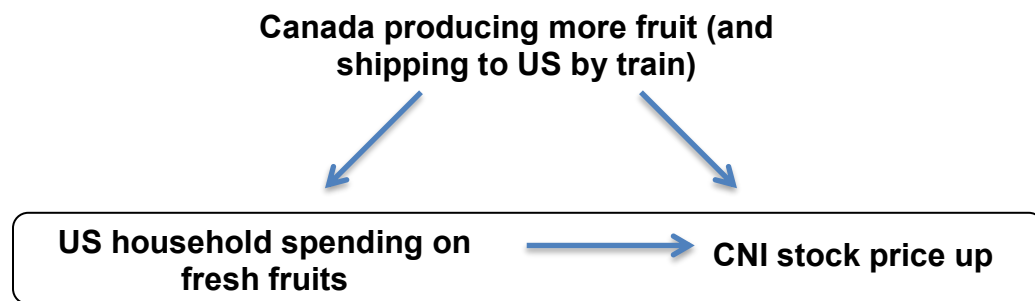
- a. What are the two variables shown in your graph? *Answer: Annual US household spending on fresh fruits and the Canadian National Railway Company's stock price (CNI).*

- b. Is there a positive association or a negative association between the variables? *Answer: There is a positive association because the lines have the same increasing shape.*
- c. Write an interpretation of this plot in the context of the data. *Answer: It seems that as more US households purchase fresh fruits, the Canadian National Railway Company's stock price (CNI) goes up.*
- d. Answers will vary. *For this example, clicking "View details about correlation #5,901" showed us the table of data used for the plot.*

| Annual US household spending on fresh fruits | CNI (stock price) |
|----------------------------------------------|-------------------|
| 178                                          | 8.03              |
| 171                                          | 6.9               |

etc.

- e. Can you conclude that the one variable *causes* the other? *Answer: No. Although the two variables are associated with one another, we do not have evidence to say that US households buying more fruit causes the CNI price to increase because the data do not come from a controlled experiment.*
- f. Draw a diagram like the one we did together earlier (in step 4 of lesson) with possible confounding factors. *Student diagrams should look similar to the one below:*



9. Once all students have selected a graph and answered the above questions, have them share their responses with a partner. They should explain why they thought their particular graph was interesting, how the two variables are related (directly or inversely), and whether or not they believe there is a causal link between the variables.
10. At the end of this lesson, students should understand that causation can only be concluded when an experiment is performed, but associations can be concluded from observational studies.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Next Day

### Lab 3B: Confound it all!

Complete Lab 3B prior to Lesson 9.

### **Lab 3B: Confound it all!**

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

#### **Finding data in new places**

- Since your first forays into doing data science, you've used data from two-sources:
  - Built-in datasets from RStudio.
  - Campaign data from the Campaign Manager.
- Data can be found in many other places though, especially online.
- In this lab, we'll read an *observational study* dataset from a website.
  - We'll use this data to then explore what factors are associated with a person's lung capacity.

#### **Importing our data**

- Rather than *export*-ing the data and then *upload*-ing and *importing*-ing it, we'll pull the data straight from the webpage into R.
- You can find the data online here:
  - (Right-click and select *Open in New Window*)  
<https://raw.githubusercontent.com/thinkdataed/dataset/main/fev.csv>
- **Click on the *Import Dataset* button under the *Environment* tab.**
  - Then click on the *From Text (readr)* option.
  - Type or copy/paste the URL into the box.
  - Click *Update*.
- **Before importing, change the following *Import Options*:**
  - Name: **lungs**
  - *Uncheck* the *First Row as Names*
  - Change *Delimiter* to *Whitespace*

#### **Our new data**

- Variables that were measured include:
  - Age in years.
  - Lung capacity, measured in liters.
  - The youth's heights, in inches
  - Genders; **"1"** for males, **"0"** for females.
  - Whether the participant was a smoker, **"1"**, or non-smoker **"0"**.

#### **About the data**

- The data come from the *Forced Expiratory Volume (FEV)* study that took place in the late 1970's.
- The observations come from a sample of 654 youths, aged 3 to 19, in/around East Boston.
- Researchers were interested in answering the *research question*:  
*What is the effect of childhood smoking on lung health?*

## Cleaning your data

- Now that we've got the data loaded, we need to clean it to get it ready for use (look at Lab 1F for help). Specifically:
  - We want to name the variables: "age", "lung\_cap", "height", "gender", "smoker", in that order.
  - (1) Write and run code changing the type of variable for gender and smoker from numeric to character.
- After changing the variable types for gender and smoker:
  - (2) For gender, write and run code using recode to change "1" to "Male" and "0" to "Female".
  - (3) For smoker, write and run code using recode to change "1" to "Yes" and "0" to "No".

## Analyzing our data

- Our lungs data is from an observational study.
- (4) Write down a reason the researchers couldn't use an experiment to test the effects of smoking on children's lungs.
- Observational studies are often helpful for analyzing how variables are related.
- (5) Do you think that a person's age affects their lung capacity? Make a sketch of what you think a scatterplot of the two variables would look like and explain.
- (6) Write and run code using the lungs data to create an xyplot of age and lung\_cap. Interpret the plot and describe why the relationship between the two variables makes sense.

## Smoking and lung capacity

- (7) Write and run code making a plot that can be used to answer the statistical investigative question:  
*Do people who smoke tend to have lower lung capacity than those who do not smoke?*
- (8) Use your plot to answer the question.
  - (9) Were you surprised by the answer? Why?
  - (10) Can you suggest a possible confounding factor that might be affecting the result?

## Let's compare

- (11) Write and run code creating three subsets of the data:
  - one that includes only 13-year-olds ...
  - one that includes only 15-year-olds ...
  - and one that includes only 17-year-olds.
- (12) Write and run code making a plot that compares the lung capacity of smokers and non-smokers for each subset.
- (13) How does the relationship between smoking and lung capacity change as we increase the age from 13 to 15 to 17?

## Sum it up!

- (14) Does smoking affect lung capacity? If so, how?
  - Support your answers with appropriate plots.
  - Explain why you included the variables you used in your plots.

# Are You Asking Me?

Instructional Days: 9

## Enduring Understandings

A survey is a data collection method that is administered to a sample. The sample is fraction of the target population. The sample must be representative of the population and random sampling is used to ensure an equal chance of being selected. A census is a special survey that collects data from the entire population. Sampling error and bias cause problems in analysis made from survey data.

## Engagement

Students will view a video clip from the game show *Family Feud* to begin to think about survey components. The video can be found at:

<https://www.youtube.com/watch?v=-3Nk9t7-rCs> (or <https://www.youtube.com/watch?v=ofQkOfeg60g>)

## Learning Objectives

### Statistical/Mathematical:

- S-IC B-4: Use data from a sample survey to estimate a population mean or proportion; develop a margin of error through the use of simulation models for random sampling.
- S-IC 1: Understand statistics as a process for making inferences about population parameters based on a random sample from that population.
- S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.
- S-IC 6: Evaluate reports based on data.

### Data Science:

Understand that bias and sampling error should be minimized when conducting surveys. The wording of questions, as well as who is asked to participate in a survey, can lead to bias. Learn that sampling error can be minimized when larger random samples are collected from a population.

### Applied Computational Thinking Using RStudio:

- Create random samples of different sizes to make estimates about a population.
- Create informal confidence intervals based on sample medians.

## Language Objectives

1. Students will describe the attributes of a survey.
2. Students will engage in partner and whole group discussions to express their understanding of sampling techniques.
3. Students will connect informal confidence intervals to formal confidence intervals in RStudio/Posit Cloud.

## Data File or Data Collection Method

### Data File:

1. American Time-Use Survey (ATUS): data(atus)

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 9: Survey Says...

### Objective:

Students will learn that a survey is another data collection method. They will learn what a survey is, what types of questions are used in a survey, and how a survey is conducted.

### Materials:

1. Video: *Family Feud*'s "Shocking Fast Money" found at:  
<https://www.youtube.com/watch?v=-3Nk9t7-rCs> (good quality, but sad ending)

OR

Video: *Family Feud* video clip titled "Family Feud – Comeback of the Century" found at:  
<https://www.youtube.com/watch?v=ofQkOfeg60g> (bad quality, but happy ending)

2. *Designing a Survey* handout (LMR\_U3\_L9)

### Vocabulary:

survey, self-reported, open-ended questions, closed-ended questions

**Essential Concepts:** Surveys ask simple, straightforward questions in order to collect data that can be used to answer statistical questions. Writing such questions can be hard (but fun)!

### Lesson:



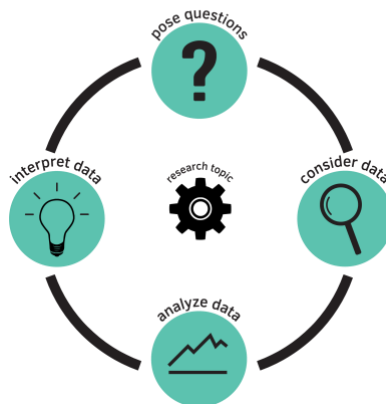
1. Introduce one of the videos listed above by informing students that they will be watching a clip from the television game show *Family Feud*. This segment of the show is called *Fast Money*, where the winning family plays for an additional \$20,000. Two family members are chosen to play and must reach a combined score of 200 points to win the money. The goal is to guess the most common responses to five questions. For example, if the question "What animal is a common pet?" were asked, each family member might answer with "dog" or "cat" since these are popular household pets. The first person accumulates as many points as possible during the 20-second first round. The second person is given 25 seconds to earn points with different answers.
2. As students watch the video, have them answer the following questions in their DS journals:
  - a. How many people were surveyed? **Answer: 100**
  - b. Who was represented in the survey? **Answer: Single men**
  - c. How many survey questions were asked? **Answer: 5**
  - d. When the host says "survey said" and we see the response, what does it mean? **Answer: It means that X number of people out of the 100 gave that response to the survey question.**
3. *Family Feud* uses surveys as its main data collection tool. In their DS journals, students should write down what they know about surveys individually.
4. Then, with a partner, students will share what each one of them knows about surveys.
5. Select a couple of students to either share their response or their partner's response with the whole class.
6. Inform students that a **survey** is a data collection method where the data are **self-reported**, meaning that participants answer questions themselves. Surveys are composed of:
  - a. Questionnaires or a series of questions
  - b. A representative sample of the population of interest
  - c. Carefully worded questions
7. Surveys rely on questions. There are two types of questions that can be asked in a survey: **open-ended questions** and **closed-ended questions**. Open-ended questions offer a free-response/text approach, whereas closed-ended questions give a fixed set of choices.





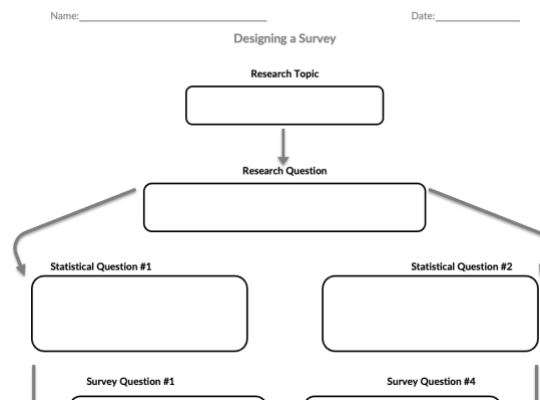
8. Display the following list to the class. With a partner, have the students categorize the following types of questions as either open-ended or closed-ended:
  - a. Multiple choice *Answer: closed*
  - b. Write a paragraph *Answer: open*
  - c. Yes/No *Answer: closed*
  - d. Comments *Answer: open*
  - e. Essays *Answer: open*
  - f. On a scale from 1 to 5 *Answer: closed*
  - g. Choose from a list *Answer: closed*
  - h. Write a sentence *Answer: open*
  - i. Check a box *Answer: closed*
9. Do a quick *Whip Around* to share the categorization for each type of question. Be sure that students make corrections to the list if any items were miscategorized.
10. Quickly review the Data Cycle.

## The Data Cycle



11. To give students an introduction to conducting surveys, they will first go through a practice scenario as a class to try to answer the following research question:
 

**What are ‘families’ in the United States?**
12. Distribute the *Designing a Survey* handout (LMR\_U3\_L9) and let students fill in the boxes for “Research Topic” (*Families*) and “Research Question” (*What are ‘families’ in the United States?*).



LMR\_U3\_L9



13. Inform students that the left side of the handout will be completed as a class, and then student teams will work together to complete the right side.

14. Using the Data Cycle as a guide, students should brainstorm a statistical question that is related to the research question. *A possible question might be: What is the typical family size in the United States?*

**Note:** This requires a definition of “family,” which can have a variety of meanings to different people. Different definitions will likely guide the discussion of possible survey questions in the following step.

15. Next, students need to determine 3 survey questions to help answer the statistical question. The goal in creating survey questions is to make sure they (1) are unambiguous, and (2) address the statistical question. Some examples are listed below (which come from different definitions of “family”):

**Note:** Survey questions **MUST** match the statistical question.

- a. How many siblings do you have?
- b. How many people live with you?

16. It may help to actually collect data once the first survey question has been created. For example: “How many siblings do you have?” – each student would give a response and the values could be recorded in a dotplot (if desired). If the question is too vague (do we include half-siblings, step-siblings, etc.), students can revise the question.

17. Once the class has agreed upon 3 survey questions for the first statistical question, allow students to join their student teams for the remainder of the activity.

18. Each team should come up with a statistical question that might answer the research question, then determine 3 survey questions that match their statistical question. Have the students create both open- and closed-ended questions in the handout. Each survey question should be a different type (see step 8 from above).

19. Have student teams share out their statistical questions and related survey questions with the rest of the class.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework & Next Day

For homework, students should choose one of their team’s survey questions and rewrite it 3 ways, using 3 different question types (see step 8 in lesson). Example rewrites for the survey question “How many siblings do you have?” are given below for reference.

- a. **Multiple choice:**

*How many siblings do you have? Select one option.*

- (a) 0 siblings
- (b) 1 sibling
- (c) 2 siblings
- (d) 3 siblings
- (e) more than 3 siblings

- b. **Write a paragraph:**

*In your own words, describe your siblings.*

- c. **Yes/No:**

*Do you have any siblings?*

## Lesson 10: We're So Random

### Objective:

Students will learn how to collect random samples from a population in order to estimate a parameter.

### Materials:

1. *Populations & Samples* handout (LMR\_U3\_L10\_A)
2. RStudio
3. Projector for RStudio functions
4. Dotplot titled "Percent of Students Who Have Met Friends Online"
5. *Parameters & Statistics* handout (LMR\_U3\_L10\_B)

### Vocabulary:

population, sample, representative, random sample, parameter, statistic

**Essential Concepts:** Another popular data collection method involves collecting data from a random sample of people or objects. Percentages based on random samples tend to "center" on the population parameter value.

### Lesson:

1. Introduce today's lesson by displaying the following statement made by the Pew Research Center in their August 2015 report titled *Teens, Technology & Friendships*:

"For today's teens, friendships can start digitally: 57% of teens have met new friends online. The margin of error is plus or minus 3.7 percentage points. Social media and online gameplay are the most common digital venues for meeting friends."

**Note:** The data for this report were collected via interviews of 1,060 teenagers between the ages of 13 and 17.



2. Discuss the results of the Pew poll with the following prompting questions:
  - a. The report says that 57% of teens have met new friends online. Since the report was based on a sample of 1,060 teens, how many of the teens reported that they have met friends online? *Answer:  $0.57(1060) = 604.2$ . This means that approximately 604 teens in the sample have met friends online.*
  - b. Do you think 57% of students in our class have met friends online? Why or why not? *Answers will vary by class. The discussion should include points about how similar and different samples can be. The sample of students in the Pew poll may not represent the students in our class.*
  - c. What percent of students in our class have met friends online while a teenager? *Answers will vary by class. Calculate the percentage by dividing the number of students who have met friends online by the total number of students who came to class today.*
  - d. What if [absent person] were in class today? Would that change our percentage? What effect would it have on the percentage if [absent person] answered "yes?" What effect would it have if [absent person] answered "no?" *Answers will vary by class. If a student who was absent for today's lesson had actually come to class, we would have a different sample of students. It would change the percentage because our sample size now includes 1 more person. If the person answered "yes," the values in the numerator and denominator of the percentage would change. If the person answered "no," the value in the denominator would change. (Students should calculate these values.)*
  - e. If we were able to interview every single teenager in the United States, would exactly 57% of them say they have met friends online? *Answer: Probably not because there are many more teenagers in the US than the 1,060 they interviewed. It would be unlikely for a group of 1,060 teenagers to exactly represent all teenagers in the entire country.*

- f. Why do you think the Pew Research Center only interviewed 1,060 teenagers, and not all teenagers in the US? *Answer: It would be impossible to talk to all teenagers in the US in a short period of time, or even a fairly long period of time.*



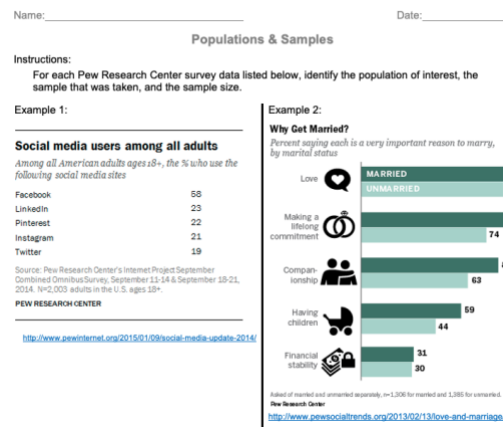
3. Introduce students to the terms **population** and **sample**. Explain that a population consists of all of the people we want to learn something about. A sample consists of people (or objects) that are selected *from* the population. In pairs, ask students to discuss and record answers to the following two questions:

- What was the population of interest to the researchers for the Pew poll above? *Answer: All teenagers in the US right now.*
- Based on your answer in (a), what characteristics should people have in order to be included in a sample for this poll? *Answer: People would need to be in the US and be between the ages of 13 and 17. People could be from many states, but we would not want to sample only people from California, or only people from Los Angeles. It would be impossible to survey every single person in the US; this is why we create a random representative sample of the population.*

**Note:** Steer the discussion so that students see that a sample has to be “like” or “similar to” or “**representative**” of the population.

- Select pairs to share their responses to the questions and let students revise their responses.
- Distribute the *Populations & Samples* handout (LMR\_U3\_L10\_A), which contains survey results from other Pew Research Center reports. Give students time to determine the population and sample for each scenario and then have them verify their results with a partner.

**Note:** Page 2 of the handout is an answer key for teacher reference.



LMR\_U3\_L10\_A

- State that we want to know the percentage of students in our class that have made friends online, but we don't want to ask every single student. Instead, we would like to ask only 10 students and then make some guesses about our class from those 10. Ask:
  - What is the population of interest? *Answer: The students in our classroom.*
  - How can we select 10 students to be part of our sample? *Answers will vary by class. There may be a variety of suggestions; here are some examples of what may be given: (1) put every student's name in a hat and pick out 10; (2) select the 10 students sitting closest to the teacher's desk; (3) have 10 students volunteer to be in the sample.*



- Inform students that, in general, we want samples to “look like” the population. One way to get a representative sample like this is to take a **random sample** – a sample that is chosen randomly.

Ask the students:

- a. Would the selection techniques we came up with in step 6 above result in random samples of our class? *Answers will vary by class. Using the examples from step 6 above: (1) putting each student's name in a hat and then picking out 10 would be a random sample as long as each piece of paper is the same size; (2) selecting the 10 students sitting closest to the teacher's desk would not be a random sample because those students might not represent the whole class; (3) if 10 students volunteer, we would not have a random sample because those students selected themselves to be part of the group and may not represent everyone in the class.*
8. Next, assign each student in the class a number by having them count off from 1 to  $N$  ( $N$  being the total number of students in the class). Show students that we can use RStudio to create random samples with the following function:  

```
> sample(1:N, size = 10, replace = FALSE)
```

**Note:** We use `replace = FALSE` because we only want each student to be selected once.
9. Using the results given in the output of RStudio, ask the students whose numbers were chosen to stand. Inform them that they are “in” the sample. Then, determine what percent of the sample (these 10 students) have made friends online. How does this percentage compare to the overall class percentage we found in step 2(c) above? *Answers will vary by class.*
10. Create a dotplot on the board titled “Percent of Students Who Have Met Friends Online.” Record the sample percentage from step 9 above on this dotplot.
11. Have the students return to their seats so that we can select a new sample. Before we do this, ask:
  - a. What do you predict the percentage to be for the next sample of 10 students? *Answers will vary by class. They might say the expected percentage will be close to the class's overall percentage.*
12. Using RStudio, create a new sample, calculate the percentage of those 10 students who have met friends online, and record the value in the dotplot.
13. Repeat step 11 from above for a few more rounds (at least 5 samples should be taken). Be sure that the students give a prediction before finding each new sample.
14. Display the following questions. Refer to the dotplot of sample percentages. Ask students to discuss the questions in teams:
  - a. What do you notice about the sample percentages? What is the “typical” value? *Answers will vary by class. The typical value should be close to the class's overall percent calculated in step 2(c) above since it is the population percent.*
  - b. What is the smallest value? What is the biggest value? *Answers will vary by class. There might be a lot of variability in the dotplot based on the selected samples. Most sample percentages will be within 30% of the population value, so that really gives a wide variety of possible sample values.*
  - c. If we took a larger sample – maybe of size 15 or 20 – would there be more or less variability in the dotplot? *Answer: There will be less variability because adding more people gets us closer to the population size. Also, point out that if the sample were exactly the same as the population, then there would be no variability in the plot.*
15. Select teams at random to share their responses to the questions above with the whole class. The rest of the teams should be in full agreement with the responses before moving on to the next question.
16. Explain that the population percentage (the percentage of all students in the class who met friends online) we have been using is called a **parameter**. A parameter is any number that summarizes a population. So, our class has been the population, and the percentage of students that have met friends online is the parameter.
17. Similarly, the term **statistics** is used for numbers that summarize a sample. Ask students what sample statistics they have seen today? *Each value we included in the dotplot is a statistic.*

18. Be sure to point out that there can be multiple values for a sample statistic (i.e. “We had 5 sample percentages in our dotplot.”), but there is always only one parameter value.
19. Using these new definitions, ask students to consider the original Pew poll data, which found that 57% of teens have met friends online. Ask the students:
  - a. Is 57% a parameter or statistic? *Answer: This is a statistic because it is based on a sample. Remind them that the population was ALL US teenagers and the sample included 1,060 teens.*
  - b. What is the population parameter then? *Answer: We don't know! We would have to interview every teenager in the US to determine the parameter, and that is not possible.*
20. Conclude the lesson by telling the students: although we cannot determine the actual population parameter for the percent of teens that have made friends online, we can estimate it using random samples.

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Homework



Students should complete the *Parameters & Statistics* handout (LMR\_U3\_L10\_B) for homework.

**Note:** Page 2 of the handout is an answer key for teacher reference.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Parameters & Statistics

Instructions:  
For each study below, identify the population, sample, parameter of interest, and any statistics.

1. A poll is a type of survey that is used to make statistical inference, a conclusion about a population based on a sample. In 2013, Gallup, a polling agency, surveyed 2,048 adults to find out Americans' main source of news. 55% of adults responded that they get their news from television.
 

|                   |                  |
|-------------------|------------------|
| Population: _____ | Parameter: _____ |
| Sample: _____     | Statistic: _____ |
2. In 2009, Time Magazine conducted an Internet poll of affluent adults (people whose income is \$150,000 per year or more). A total of 603 affluent adults over the age of 18 were interviewed. They found that 95% of affluent Americans made online purchases in the last year.
 

|                   |                  |
|-------------------|------------------|
| Population: _____ | Parameter: _____ |
| Sample: _____     | Statistic: _____ |
3. In a 2013 article published by The Guardian, an English newspaper, a survey found that 62% of 16–24-year-olds prefer print books over digital books. In this survey, 1,420 young adults aged 16–24 were interviewed.

LMR\_U3\_L10\_B

## Lesson 11: The Gettysburg Address

### Objective:

Students will learn the definition of sampling bias and will learn that random samples reduce bias when estimating a population parameter. They will gain practice collecting a random sample from a small population and estimating the population parameter.

### Materials:

1. *Gettysburg Address* handout (LMR\_U3\_L11\_A)
  2. *Sampling the Gettysburg Address* handout (LMR\_U3\_L11\_B)
  3. Dotplot titled “Mean Word Length, Self-Selected Sample, Size = 10” – on board or poster paper
  4. *Gettysburg Address – Word Length Histogram* file (LMR\_U3\_L11\_C)
  5. *Gettysburg Word Lengths* handout (LMR\_U3\_L11\_D)
  6. RStudio
  7. Projector for RStudio functions
  8. Dotplot titled “Mean Word Length, Random Sample, Size = 10” – on board or poster paper
- Note:** This dotplot will be used again during Lesson 13, so the results need to be kept in some way (this can be either on poster paper or by simply taking a photo)

### Vocabulary:

sampling bias

**Essential Concepts:** Statistics vary from sample to sample. If the typical value across many samples is equal to the population parameter, the statistic is “unbiased”. Bias means that we tend to “miss the mark.” If we don’t do random sampling, we can get biased estimates.

### Lesson:

1. Pose the following question to students: “What makes a speech truly great and memorable?” Have students discuss in their groups then invite the reporters to share out.  
Tell students that technical components, such as word length, matter. Ask: “What might be wrong with using long words in a speech?” *Answers will vary. Ann Wylie, a writing trainer and PR specialist, explains the importance of word length and why we should care about it. In her blog post, she says that using long words in a speech should be avoided in order to increase readability/comprehension and make your ideas more clear/compelling.*  
*For example, she cites The Gettysburg Address that has 272 words, of which 174 of them have only one syllable. You can read her blog post here: <https://www.wyliecomm.com/2021/11/what-is-word-length-and-why-should-you-care/>*
2. Ann Wylie states that The Gettysburg Address is a “small wonder” in that it uses small words to express big ideas. Share facts about the Gettysburg Address:
  - a. President Lincoln delivered the Gettysburg Address in November 1863.
  - b. It is one of the most famous speeches in the United States.
  - c. In it, Lincoln invoked the principles of human equality contained in the U.S. Constitution and Declaration of Independence.
3. Read the Gettysburg Address aloud to the class OR have students read it aloud. The text of the speech can be found in the *Gettysburg Address* handout (LMR\_U3\_L11\_A).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

The Gettysburg Address  
By President Abraham Lincoln

Four score and seven years ago our fathers brought forth on this continent, a new nation, conceived in Liberty, and dedicated to the proposition that all men are created equal.

Now we are engaged in a great civil war, testing whether that nation, or any nation so conceived and so dedicated, can long endure. We are met on a great battle-field of that war. We have come to dedicate a portion of that field, as a final resting place for those who here gave their lives that that nation might live. It is altogether fitting and proper that we should do this.

But, in a larger sense, we can not dedicate -- we can not consecrate -- we can not hallow -- this ground. The brave men, living and dead, who struggled here, have consecrated it, far above our poor power to add or detract. The world will little note, nor long remember what we say here, but it can never forget what they did here. It is for us the living, rather, to be dedicated here to the unfinished work which they who fought here have thus far so nobly advanced. It is rather for us to be here dedicated to the great task remaining before us -- that from these honored dead we take increased devotion to that cause for which they gave the last full measure of devotion -- that we here highly resolve that these dead shall not have died in vain -- that this nation, under God, shall have a new birth of freedom -- and that government of the people, by the people, for the people, shall not perish from the earth.

LMR\_U3\_L11\_A

4. Today we will use the Gettysburg Address to learn about different sampling techniques.
5. Inform students that the Gettysburg Address contains 272 words. We can consider these 272 words to be our population because it includes all words in the entire speech. From the population, we can sample 10 words that we think represent the speech. It is ok for this step to be vague -- students can come up with their own concept for what they think "representative" means in this case.
6. Distribute the *Sampling the Gettysburg Address* handout (LMR\_U3\_L11\_B), which includes the actual speech, as well as 2 sampling activities. For this part of the lesson, we will only be looking at Sampling Activity 1 on page 1 of the handout.

**Note:** This activity was originally created by Allan Rossman and Beth Chance, and has been modified for the IDS curriculum.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Sampling the Gettysburg Address

The Gettysburg Address  
By President Abraham Lincoln

Four score and seven years ago our fathers brought forth on this continent, a new nation, conceived in Liberty, and dedicated to the proposition that all men are created equal.

Now we are engaged in a great civil war, testing whether that nation, or any nation so conceived and so dedicated, can long endure. We are met on a great battle-field of that war. We have come to dedicate a portion of that field, as a final resting place for those who here gave their lives that that nation might live. It is altogether fitting and proper that we should do this.

But, in a larger sense, we can not dedicate -- we can not consecrate -- we can not hallow -- this ground. The brave men, living and dead, who struggled here, have consecrated it, far above our poor power to add or detract. The world will little note, nor long remember what we say here, but it can never forget what they did here. It is for us the living, rather, to be dedicated here to the unfinished work which they who fought here have thus far so nobly advanced. It is rather for us to be here dedicated to the great task remaining before us -- that from these honored dead we take increased devotion to that cause for which they gave the last full measure of devotion -- that we here highly resolve that these dead shall not have died in vain -- that this nation, under God, shall have a new birth of freedom -- and that government of the people, by the people, for the people, shall not perish from the earth.

**Sampling Activity 1**

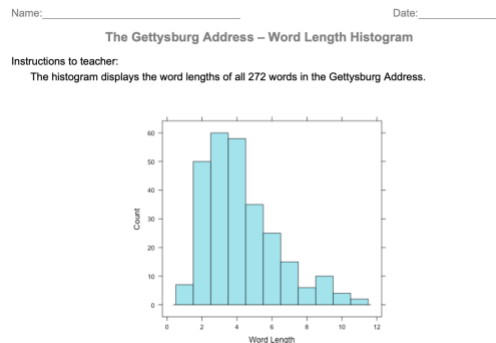
1. Circle 10 words that you think might be representative of all words in the speech.
2. Record your self-selected words and their corresponding word lengths in the table.

| Word # | Word | Word Length | Word # | Word | Word Length |
|--------|------|-------------|--------|------|-------------|
|--------|------|-------------|--------|------|-------------|

LMR\_U3\_L11\_B

7. Inform students that they will get 30 seconds to select 10 words that they think are representative of all words in the speech.  
**Note:** It is important that students work fast so they are forced to choose based on first impressions and don't have time to reflect. Also, this activity tends to not work well if students are informed of the punch line (that random samples are unbiased) before they begin.
8. At this point, explain to students that we are actually interested in answering a specific question:  
**What is the typical word length in the Gettysburg Address?**
9. Next, students should record each circled word, as well as the number of letters each word has (this is the word length) in the table on the handout. Then, they should summarize the data in a dotplot and calculate the mean word length of the sample.

10. On the board, create a class dotplot (may also be done on poster paper) titled "Mean Word Length, Self-Selected Sample, Size = 10." Once all students have completed the first page of the *Gettysburg Address* handout (LMR\_U3\_L11\_B), ask each student to record the mean word length of his or her sample on the class's dotplot.
11. When all students have recorded their sample statistics in the dotplot, conduct a class discussion based on the questions listed below.
- Note:** You might need to do a reality check. Students will often make mistakes when adding the word lengths and when dividing. Be suspicious (and double-check) extreme values.
- What does each point on the plot represent? *Answer: Each point represents one student's estimate of the mean length of all of words in the Gettysburg address.*
  - What is the typical value represented in the dotplot? *Answers will vary by class. You should indicate the approximate location of the mean of the distribution (the balancing point) on the dotplot. Remind students that when we ask for the 'typical' value we mean the value in the center of the distribution.*
  - How much variability is there in the distribution? *Answers will vary by class. One reasonable approach is for students to give the range (the difference between the largest and smallest values).*
  - What is the shape of the distribution? *Answers will vary by class. Often, the shape is right-skewed, but it might not be for you. Note that outliers here will often be arithmetic errors, but not always.*
12. Next, display the histogram from the *Gettysburg Address - Word Length Histogram* file (LMR\_U3\_L11\_C), which shows the distribution of word lengths for the entire population of words in the *Gettysburg Address*.



LMR\_U3\_L11\_C

13. Remind students that the population is the 272 words from the speech and inform them that the mean word length of the population, or the population parameter, is 4.22. Using *Think-Pair-Share*, ask:
- How does the typical value of our class's sample means compare to the actual population mean of 4.22? *Answer: Almost always, the class's typical mean will be higher (sometimes much higher) than 4.22. Some students will be close to 4.22. But point out that we are talking about the "trend" or typical outcome. The typical outcome is usually too high.*
14. Explain that sampling plan was a self-selected sample, which often produces biased results. **Sampling bias** occurs when the resulting samples tend to produce results that are influenced in one particular direction.
15. Refer back to the dotplot of sample means and point out how it is biased. Ask:
- Why was our original sampling procedure biased? *Answer: When we look for 'representative' words, and do so quickly, our eyes are drawn by the larger, more unusual words, and we tend to overlook small words such as "in," "a," "we," etc.*

16. Go back to the *Gettysburg Address* handout (LMR\_U3\_L11\_B), and direct students to page 2 for Sampling Activity 2. Inform students that they will now do a sampling procedure that results in a better representation of the population of words in the speech.

17. Explain that a random sample tends to produce unbiased sample results.

18. Before students begin the activity, demonstrate how to generate 10 random numbers from a possible 272 using RStudio.

```
> sample((1:272), size = 10, replace = FALSE)
```

19. Each student should generate his or her own set of 10 random numbers. Once students have created their random numbers, distribute the *Gettysburg Address Word Lengths* handout (LMR\_U3\_L11\_D).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

The Gettysburg Address - Word Lengths

| Number | Word       | Length | Number | Word       | Length | Number | Word        | Length |
|--------|------------|--------|--------|------------|--------|--------|-------------|--------|
| 001    | Four       | 4      | 046    | nation     | 6      | 091    | night       | 5      |
| 002    | score      | 5      | 047    | so         | 2      | 092    | live,       | 4      |
| 003    | and        | 3      | 048    | conceived  | 9      | 093    | It          | 2      |
| 004    | seven      | 5      | 049    | and        | 3      | 094    | is          | 2      |
| 005    | years      | 5      | 050    | so         | 2      | 095    | altogether  | 10     |
| 006    | ago        | 3      | 051    | dedicated, | 9      | 096    | fitting     | 7      |
| 007    | our        | 3      | 052    | can        | 3      | 097    | and         | 3      |
| 008    | fathers    | 7      | 053    | long       | 4      | 098    | proper      | 6      |
| 009    | brought    | 7      | 054    | endure.    | 6      | 099    | that        | 4      |
| 010    | forth      | 5      | 055    | We         | 2      | 100    | we          | 2      |
| 011    | on         | 2      | 056    | are        | 3      | 101    | should      | 6      |
| 012    | this       | 4      | 057    | met        | 3      | 102    | do          | 2      |
| 013    | continent, | 9      | 058    | on         | 2      | 103    | this.       | 4      |
| 014    | a          | 1      | 059    | a          | 1      | 104    | But,        | 3      |
| 015    | new        | 3      | 060    | great      | 5      | 105    | in          | 2      |
| 016    | nation,    | 6      | 061    | battle-    | 6      | 106    | a           | 1      |
| 017    | conceived  | 9      | 062    | field      | 5      | 107    | larger      | 6      |
| 018    | in         | 2      | 063    | of         | 2      | 108    | sense,      | 5      |
| 019    | liberty,   | 7      | 064    | that       | 4      | 109    | we          | 2      |
| 020    | and        | 3      | 065    | war.       | 3      | 110    | can         | 3      |
| 021    | dedicated  | 9      | 066    | We         | 2      | 111    | not         | 3      |
| 022    | to         | 2      | 067    | have       | 4      | 112    | dedicate -- | 8      |
| 023    | the        | 3      | 068    | come       | 4      | 113    | we          | 2      |

LMR\_U3\_L11\_D

20. Explain that the table contains the word number, word, and length of each word in the *Gettysburg Address*. Demonstrate how to find a word that corresponds to one of the random numbers generated by RStudio and explain that this word is now part of our random sample.

21. Then, each student will complete the handout by creating a dotplot and calculating the mean of their random sample.

22. On the board, near the first dotplot, create another class dotplot (may also be done on poster paper) titled "Mean Word Length, Random Sample, Size = 10." Once all students have completed the second page of the *Gettysburg Address* handout (LMR\_U3\_L11\_B), ask each student to record the mean word length of his or her random sample on the class's dotplot.

**Note:** As in step 9 from above, be sure to check arithmetic for outliers!



23. When all students have recorded their sample statistics in the dotplot, conduct a class discussion based on the following questions:

- What does each point in the dotplot represent? *Answer: Each dot represents one student's estimate of the mean word length. But this time, the estimates are based on a random sample of 10 words.*
- What do you notice about the typical value in this distribution? *Answers will vary by class. They should notice that the means of the random samples are fairly close to the population mean of 4.22. (Again, you might have to discard or correct outliers.)*
- What shape does this distribution have? What does that tell us? *Answer: Typically, the distribution of sample means for random samples will be symmetric and unimodal.*
- What does this distribution tell us about the benefits of random samples? *Answer: We can reduce bias by collecting random samples instead of self-selected samples.*
- Why do we need sampling? Why can't we just get the information from the actual population? *Answer: It is usually impossible to include every person or object from a population. Even for the population of size 272 words in the Gettysburg Address, it would take a long time to calculate the word lengths of every single word.*

24. Conclusions and takeaways:

- a. It turns out that there are approximately  $5.17 \cdot 10^{17}$  different possible samples of 10 words from the Gettysburg Address.
- b. If we could determine the mean for each of these samples and produce a dotplot for all of these means, then the center of the dotplot would lie exactly at 4.22.
- c. The resulting distribution of the means from all possible samples is called the sampling distribution for the sample mean (for samples of size 10 from this population).
- d. The above dotplot is an approximation of the actual sampling distribution.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

**Homework & Next Day**

Students should write a reflection about why random sampling is better at reducing bias than other sampling procedures.

**Lab 3C: Random Sampling**

Complete Lab 3C prior to Lesson 12.

### Lab 3C: Random Sampling

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

#### Learning by sampling

- In many circumstances, there's simply no feasible way to gather data about everyone in a *population*.
  - For example, the Department of Water & Power (DWP) wants to determine how much water people in Los Angeles use to take a shower. They've created a survey to pass out to collect this information.
  - **(1) Write down two reasons why getting everyone in Los Angeles to fill out the survey would be difficult. Also, write a sentence why the DWP might consider using a sample of households instead.**
- In this lab, we'll learn how *sampling methods* affect how *representative* a sample is of a *population*.

#### Loading a population

- In previous labs, we used the **cdc** data as a sample for young people in the United States.
  - In this lab, we'll consider these survey respondents to be our population.
- **(2) Load the **cdc** data into R and fill in the blanks to take a *convenience* sample of the first 50 people in the data:**

```
s1 <- slice(____, 1:____)
```

- **(3) Why do you think we call this method a *convenience* sample?**

#### Comparing your convenience sample

- A convenience sample is a sample from a population where we collect data on subjects because they're easy-to-find.
- **(4) Using your convenience sample, write and run code creating a **bargraph** for the number of people in each **grade**.**
  - **(5) Do you think the distribution of **grade** for your sample would look similar when compared to the whole **cdc** data?**
  - **(6) Which groups of people do you think are over- or under-represented in your convenience sample? Why?**
- **(7) Write and run code creating a **bargraph** for **grade** using the **cdc** data.**
  - **(8) Compare the distributions of the **cdc** data and your convenience sample and write down how they differ.**

#### Using randomness

- **(9) Fill in the blanks below to create a sample by randomly selecting 50 people in the **cdc** data, without replacement. Call this new sample **s2**:**

```
____ <- sample(____, size = ____, replace = ____)
```

- **(10) Write a sentence that explains why you think the distribution of **grade** for this *random sample* will look more or less similar to the distribution from the whole **cdc** data.**
- **(11) Write and run code creating a **bargraph** for **grade** based on this *random sample* to check your prediction.**

## Increasing sample size

- (12) Write and run code creating `bargraphs` for `grade` based on each of the following sample sizes: 10, 100, 1,000, 10,000.
  - Compare each distribution to that of the population.
- (13) How do the distributions change as the size of the sample increases? Why do you think this occurs?
- (14) Write and run code to `tally()` the proportion of `grades` for your *convenience* sample and all your *random* samples.
  - (15) Which set of proportions looks most similar to the proportions of the population?

## Lessons learned

- The proportion (or mean) from a *random* sample might not always be closer to that of the true population when compared to a *convenience* sample.
- However, as sample sizes get larger:
  - *Random* samples will tend to be better estimates for the population.
  - With *convenience* samples, this might not be the case.
- (16) Write down a reason why estimates based on *convenience* samples might not improve even as sample size increases.

## Lesson 12: Bias in Survey Sampling

### Objective:

Students will learn about bias in relation to survey sampling. More specifically, they will learn what types of sampling methods could result in a biased sample, who might be over/under-represented in the sample, and how to select a better sample.

### Materials:

1. *Identifying Biased Samples* handout (LMR\_U3\_L12\_A)
2. Poster paper
3. *Survey Sampling* handout (LMR\_U3\_L12\_B)

### Vocabulary:

survey sample, over-represented, under-represented, random sampling

**Essential Concepts:** Bias concerning survey sampling includes identifying sampling methods that may lead to biased samples, recognizing potential over- or under-representation in samples, and acquiring skills to choose more reliable sampling techniques.

### Lesson:

1. Remind students that they learned about biased samples during the last few lessons. Today, they will continue with this topic and discuss how people are selected to be in a sample.
2. The people who are asked to participate in a survey are known as the **survey sample**. Ideally, the people who are included in the survey sample are a representative group of the target population, or the population we would like to make inferences about.
3. Propose the following scenario to the class: “An elementary school is going to start serving ice cream in the cafeteria every Friday during lunch and needs to know the favorite flavor of its students.”
4. In pairs, ask students to come up with two examples of samples that might be biased. For instance, one biased sample might include only the four 3<sup>rd</sup> grade classes at the school. For each biased sample, the students should answer the following questions in their DS journals:
  - a. Who is the target population? *Answer: All students at the elementary school.*
  - b. Who is included in your biased sample? *Answer: Only 3<sup>rd</sup> grade students. These students are over-represented in the sample.*
  - c. Who is not included in your biased sample? *Answer: All other students in the school (Kindergartners, 1<sup>st</sup>, 2<sup>nd</sup>, 4<sup>th</sup>, and 5<sup>th</sup> graders). These students are under-represented in the sample.*
  - d. Is your sample representative of the target population? *Answer: No! We’re only including 3<sup>rd</sup> graders, and they may not have the same preferences as other students.*
5. Once pairs have come up with their biased samples and answered the questions in step 4 above, they should share out with their student teams and answer the following questions:
  - a. How is your biased sample different from the samples created by other pairs in your team? *Answers will vary by class. An example might be that one pair sampled only 3<sup>rd</sup> graders and the other pair sampled only girls.*
  - b. Which do you think is more representative of the target population? Why? *Answers will vary by class. Using the example above, we could argue that either one is more representative. We could maybe say that, since 3<sup>rd</sup> graders include both boys and girls, we have a more representative sample than if we just sampled girls. Or, we could say that since girls come from all grade levels, they’re more representative of the entire school than just 3<sup>rd</sup> graders.*

6. After the teams have discussed their samples, they should select one pair's biased sample to share with the rest of the class. Record each team's biased sample on a sheet of poster paper with the following layout:

| Biased Sample                      | Who is overrepresented?        | Who is underrepresented?                                                                                               |
|------------------------------------|--------------------------------|------------------------------------------------------------------------------------------------------------------------|
| All 3 <sup>rd</sup> grade students | 3 <sup>rd</sup> grade students | All other students (kindergartners, 1 <sup>st</sup> , 2 <sup>nd</sup> , 4 <sup>th</sup> , and 5 <sup>th</sup> graders) |

7. Distribute the *Identifying Biased Samples* handout (LMR\_U3\_L12\_A). In this activity, students will explain why a particular sampling method might result in a biased sample – a sample that is not representative of the population of interest.

**Note:** It is NOT enough for students to say that the “sample is not random.” They need to explain how the sample is biased.

**Note:** Page 2 of the handout provides sample answers for teacher reference. Do NOT distribute page 2 to students.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Identifying Biased Samples**

Instructions:  
For each example given below, explain why the resulting sample might be biased.

1. A researcher sends out 500 questionnaires about pollution in Los Angeles to local residents by mail. She receives 340 responses.

Population of interest: \_\_\_\_\_

Why might the sample be biased? Explain. \_\_\_\_\_

\_\_\_\_\_

2. A researcher is interested in learning the typical number people per household who own cell phones. He conducts a survey by randomly calling phones that have land-lines.

Population of interest: \_\_\_\_\_

Why might the sample be biased? Explain. \_\_\_\_\_

\_\_\_\_\_

3. A researcher has concluded that dolphins are nice animals by surveying people who were assisted by one in a shark attack.

Population of interest: \_\_\_\_\_

Why might the sample be biased? Explain. \_\_\_\_\_

LMR\_U3\_L12\_A



8. Each student should complete the handout independently. Afterwards, conduct a whole-class discussion to compare and contrast different students' explanations of how the samples might be biased. For each example given in the handout, discuss who is most likely over-represented and who is most likely under-represented in the sample.

9. Ask the students:

- a. Now that we have learned about sampling biases, how can we eliminate this type of bias?  
*Answers will vary by class.*

**Note:** Allow students to collaborate and come up with a few ideas on their own of how to eliminate sampling bias. If desired, ideas can be written on the board for discussion and comparison.

10. Conclude with the actual answer: **random sampling.**

- a. If we randomly sample people from our population of interest, we can reduce the bias of any sample statistics obtained from the survey responses.
- b. If we have a biased sample, we can only give descriptions about that particular sample; we CANNOT generalize to the population of interest.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Students will complete the *Survey Sampling* handout (LMR\_U3\_L12\_B) for homework.

**Note:** Page 2 of the handout provides sample answers for teacher reference.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

### Survey Sampling

#### Instructions:

Read the survey sampling scenario below. Then, read the questions that follow. Re-read the scenario focused on the questions. Finally, write your response to each question.

#### Scenario

A researcher was asked to design an investigation based on the following research question:

**Do American adult males spend more of their prime time hours online or watching regular television?**  
(Prime time is defined as weekday evenings from 7pm to 9pm)

The researcher decided to conduct a survey sample. He determined that he would choose 1,000 American males over the age of 18 through a landline telephone survey. The researcher would randomly select the telephone numbers of adult males that live on the east coast of the United States.

#### Answer the following questions:

1. What is the target population and what is the sample?
2. Explain why the sampling plan might be biased.

LMR\_U3\_L12\_B

## Lesson 13: The Confidence Game

### Objective:

Students will learn about informal confidence intervals for making estimates about population parameters based on statistics from random samples.

### Materials:

1. *The Confidence Game* handout (LMR\_U3\_L13)
2. Dotplot titled “Number Correct” displayed on the board (or on poster paper)

### Vocabulary:

inferences, interval, confidence interval

**Essential Concepts:** There is uncertainty when we estimate population parameters. Because of this, it is better to give a range of plausible values, rather than a single value.

### Lesson:

1. Remind students that they have been learning about why sampling allows us to make **inferences** about a population. Some methods of sampling produce biased sample statistics, which does not allow us to generalize the results from a sample to the population of interest. To obtain unbiased statistics, random sampling methods need to be used.
2. Conduct a class brainstorm about what it means to “estimate” something. Have students come up with possible synonyms for the word “estimate.” *Some example synonyms include: guess, approximation, projection, opinion, impression, etc.*
3. Inform students that, in statistics, to provide an estimate means that we can give a range of values that we are confident include the population parameter value.
4. In today’s lesson, explain that the students will be playing a game, called *The Confidence Game*, in which they will be asked a series of questions that each have one numerical answer. However, instead of guessing what the exact answer is, the students will create a range of possible values that they think might include the real answer. They should be 90% confident that the true value is within their interval.
5. Introduce *The Confidence Game* to students by first going through an example using the question:  
**How tall is the Empire State Building, in feet (including the spire at the very top)?**
  - a. Ask the students to write down an **interval**, or range, of values that they think contains the true height of the building.
  - b. Have a few students share their intervals with the class and discuss any obvious similarities or differences between them.  
**For example:** If Student A gives an interval from 500 to 2000 feet and Student B gives an interval from 1100 to 1400 feet, one discussion could stem from asking Student A why he or she isn’t as sure of the answer as Student B is (since Student B gave a narrower interval). Then see if Student A wants to change his or her interval.
  - c. After the discussion, tell the students that the actual height of the Empire State Building is 1,454 feet tall. Take a poll to see how many students’ intervals contained this value. We will learn what it means to have the true value in our intervals after we play the game.
6. Now we can actually play the game! Distribute *The Confidence Game* handout (LMR\_U3\_L13) and explain the rules. Students will have about 5 minutes to complete the handout, which gives them approximately 30 seconds per question.

**Note:** The rules are printed at the beginning of the handout. They are included here for your convenience.

- Each question must be answered WITH AN INTERVAL.
- You should choose your interval so that you are “90% confident” (whatever that means to you).
- You CANNOT use any reference tools (i.e. no cell phones or computers to find answers).
- A question is “correct” if the true answer is inside your interval.
- The winner is determined by who got the most questions correct. In the case of a tie, the winner is chosen by whose intervals were narrower.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**The Confidence Game**

Rules of the game:

- Each question must be answered WITH AN INTERVAL.
- You should choose your interval so that you are “90% confident” (whatever that means to you).
- You CANNOT use any reference tools (i.e., no cell phones or computers to find answers).
- A question is “correct” if the true answer is inside your interval.
- The winner is determined by who got the most questions correct. In the case of a tie, the person whose intervals were narrower is the winner.

---

- In what year did Mickey Mouse make his film debut?
- What is the lowest temperature (in degrees Fahrenheit) ever recorded in California?
- During the year 2014, how many television series were aired?
- How far away, in miles, is the Earth from the Moon?

LMR\_U3\_L13

- Once each student has completed *The Confidence Game* handout (LMR\_U3\_L13), have students choose partners and exchange handouts so that they can grade each other. Remind them that a question is marked as “correct” if the actual value (see answers in lesson step 8 below) falls within the interval.
- Display the answers for each of the 10 questions from the handout found below:
  - In what year did Mickey Mouse make his film debut? **Answer: 1928.**
  - What is the lowest temperature (in degrees Fahrenheit) ever recorded in California? **Answer: -45 degrees Fahrenheit (in Boca, CA on January 20, 1937).**
  - During the year 2014, how many television series were aired? **Answer: 1,715 TV shows.**
  - How far away, in miles, is Earth from the moon? **Answer: 238,900 miles.**
  - What is the greatest number of children officially recorded that were all born to one mother? **Answer: 69 children.**
  - In what year did Orville and Wilbur Wright, more commonly known as the Wright brothers, make the first-ever powered flight in a plane? **Answer: 1903.**
  - As of June 2015, how many of Rihanna’s songs have reached the Number 1 spot on *Billboard’s* “Dance Club Hits” chart? **Answer: 23 songs.**
  - How many years have actors Will Smith and Jada Pinkett-Smith been married? **Answer: As of June 2025, it was 27 years – they were married December 31, 1997.**
  - How many hours will it take to complete a cross-country road trip from Los Angeles to New York City according to Google Maps? **Answer: 41 hours (2,789.5 miles).**
  - How many baseball fans can attend game at Dodger Stadium during any given day? **Answer: 56,000 fans.**
- Each student should write the total number of “correct” responses at the top of his or her partner’s handout and then return it.



10. Engage the students in a discussion about how well they did at estimating the true values with their intervals. The following questions can be used to steer the discussion:
- Remember that we were aiming to be 90% confident for each question. Based on this, how many of the 10 questions should we each have gotten correct? *Answer: If we are 90% confident, then we would expect 90% of the 10 intervals to include the true value, which is 9 intervals.*
  - Did anyone in the class get exactly 9 correct? Did anyone get all 10 correct? *Answers will vary by class. However, it is very unlikely that many students will have gotten 9 or 10 correct responses on this first round.*



11. Create a dotplot on the board (or on poster paper) titled “Number Correct” and have each student record his or her value. Then, ask:
- How many students got 9 correct? In other words, how many students were actually 90% confident of their intervals? *Answers will vary by class.*
  - What is the typical number of correct responses for our class? Does it seem too high or too low? Explain. *Answers will vary by class. Most likely, the typical number of correct responses will be fairly low (maybe even 4 or less).*
  - Why is our typical score so much lower than 9? *Answer: We tend to be more confident than we should be, so we create narrower intervals.*
  - It looks like, even though we thought we were 90% confident, most of us (or all of us) did not succeed 90% of the time. How could we increase our level of confidence? *Answer: We could use wider intervals.*

12. Recall from step 3 above that, in statistics, to estimate something means that we can give a range of values that we are confident include the population parameter value. This range of values, like the ones the students created during *The Confidence Game* activity, is known as a **confidence interval**.

13. Students will continue to learn about confidence intervals during the next lesson.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



In your own words, write a description of what a confidence interval is and why it is used in statistics.

## Lesson 14: How Confident Are You?

### Objective:

Students will learn about informal confidence intervals and estimates for the margin of error.

### Materials:

1. Dotplot titled “Mean Word Length, Random Sample, Size = 10” – from Lesson 11

### Vocabulary:

margin of error, bootstrapping

**Essential Concepts:** The margin of error expresses our uncertainty in an estimate. The estimate, plus or minus the margin of error, gives us an interval in which we are very confident the true value lies.

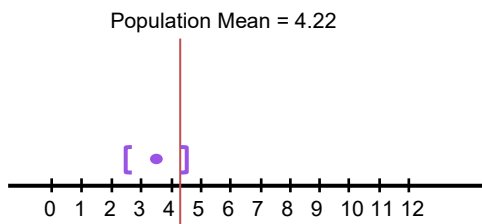
### Lesson:

1. In this lesson, students will learn about confidence intervals in more detail.
2. Display the dotplot the class created during Lesson 11 (The Gettysburg Address) titled “Mean Word Length, Random Sample, Size = 10.”
3. Have students recall that each dot represents one student’s calculation of the mean word length of a sample of 10 randomly selected words from the Gettysburg Address.
4. Also remind them that the population parameter, which is the mean word length of all words in the speech, was 4.22. There should already be a vertical line on the dotplot to indicate this value, but if it is not present, please add it during this step. Ask:
  - a. What vocabulary word was used to describe each of the sample means we each created during Lesson 11? *Answer: The sample statistic. Every dot on the graph represents one sample statistic, more specifically each dot corresponds to a different sample mean.*
  - b. How many of us got exactly the right value? *Answer: Probably none.*
  - c. Thinking back on The Confidence Game we played yesterday, what approach could we do so that 90% of us would be correct? *Answer: We could give an interval.*
5. Show students that they can give an interval in the form:

**Your sample statistic plus or minus AMOUNT**

6. Ask them to calculate what their AMOUNT must be so that their interval includes the parameter value. Ask them to write this as an interval.
7. Choose one student to illustrate what is to be done. Ask them for their AMOUNT. On the dotplot, find their value, and use bars to go out plus and minus the AMOUNT. Confirm that it includes the parameter value.

**For example:** The purple dot represents a sample mean of 3.5. The AMOUNT we have chosen for this particular case is 0.8, so the lower bracket is 0.8 below the sample mean, and the upper bracket is 0.8 above the sample mean. Notice that the population parameter is included within the brackets.



8. Now, convert this to an interval of the form [lowest value, highest value] by subtracting the amount from the sample statistic to get the lowest value, and adding to get the highest.

9. Inform students that this AMOUNT is called the **margin of error**. Explain that the students all now have different margins of error because in this unusual 'game' they know the population value. But in real life we do not, and so we have to choose one single margin of error that will work 90% of the time.
  10. Ask the students what margin of error they should use so that 90% of the estimates will have a 'successful' interval. You might want to tell them how many estimates that is for your class. A ballpark figure for the margin of error is 1.3.
  11. Explain: If we were to start all over, we could imagine picking one of these sample statistics at random.
    - a. What's the probability that the sample statistic plus or minus the margin of error would include the parameter value? **Answer: 90%.**
  12. Because of this, we call these 'confidence intervals.' When we report an interval, for example 2.7 to 4.3, we say "We are 90% confident that the population parameter value is between 2.7 and 4.3." This is another way of saying "We don't know what the exact true value is, but we're confident it is somewhere in this interval."
 

**Note:** Correct definition: 90% confidence means that if you were to compute the interval many times, 90% of them would include the population parameter.

Misconception: 90% confidence means there is a 90% chance that the population parameter is in the interval.
  13. Remind students of the Pew Poll they discussed during Lesson 10. For reference, the Pew Research Center made the following statement in their August 2015 report titled *Teens, Technology & Friendships*: **Pew Poll**

"For today's teens, friendships can start digitally: 57% of teens have met new friends online. The margin of error is plus or minus 3.7 percentage points. Social media and online gameplay are the most common digital venues for meeting friends."

**Note:** The data for this report were collected via interviews of 1,060 teenagers between the ages of 13 and 17.
- ☒ 14. Now that students have learned about the margin of error, have them write an *Exit Slip* about what the margin of error means in context of the Pew Poll.
15. Conclusions and takeaways:
  - a. Estimates that are based on random samples vary.
  - b. We can measure this variation.
  - c. The margin of error can tell us how much estimates vary.
  - d. We can use the estimate from our random sample, along with the margin of error, to give us a range of plausible values for the population parameter. This is called a confidence interval.
16. In preparation for the lab, introduce students to the idea of **bootstrapping**. Bootstrapping is a method statisticians use to learn more about a population of interest by studying a sample that they have already collected. Instead of getting more data, they create new samples by picking from the existing data, repeating this process many times. This helps them figure out more about the original population of interest without needing additional data. Show the following video for an example: [https://www.youtube.com/watch\\_popup?v=Xz0x-8-cgaQ&t=1s](https://www.youtube.com/watch_popup?v=Xz0x-8-cgaQ&t=1s)

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Next Day

### **Lab 3D: Are you sure about that?**

Complete Lab 3D prior to the Let's Build a Survey! Practicum.

### Lab 3D: Are you sure about that?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

#### Confidence and intervals

- Throughout the year, we've seen that:
  - Means are used for describing the typical value in a sample or population, but we usually don't know what they are, because we can't see the entire population.
  - Means of samples can be used to *estimate* means of populations.
  - By including a margin of error with our estimate, we create an interval that increases our confidence that we've located the correct value of the population mean.
- Today, we'll learn how we can calculate *margins of error* by using a method called the *bootstrap*.
  - Which comes from the phrase, *Picking yourself up by your own bootstraps*.

#### In this lab

- **Load the built-in `atus` (American Time Use Survey) dataset, which is a survey of how a sample of Americans spent their day.**
  - **(1) The United States has an estimated population of 336,302,171 (as of April 15, 2024 9:10 a.m. PDT). How many people were surveyed for this particular dataset?**
  - Note: If you want to know the US population or world population in real time, click on this link: <https://www.census.gov/popclock/>
- The statistical investigative question we wish to answer is:  
*What is the mean age of people older than 15 living in the United States?*
- **(2) Why is it important that the ATUS is a random sample?**
- **(3) Use our `atus` data to calculate an estimate for the average age of people older than 15 living in the U.S.**

#### One bootstrap

- A *bootstrapped* sample is when we take a random `sample()` of our original data (`atus`) **WITH** replacement.
  - The `size` of the sample should be the same size as the original data.
- We can create a single *bootstrapped* sample for the `mean` in 3 steps:
  1. Sample the number of the rows to use in our *bootstrap*.
  2. `slice` those rows from our original data into our *bootstrap* data.
  3. Calculate the mean of our *bootstrapped* data.

#### Our first bootstrap

- **(4) Fill in the blanks to `sample` the row numbers we'll use in our *bootstrapped* sample.**
  - Be sure to re-read what a *bootstrapped* sample is from the previous slide to help you fill in the blanks.
  - Use `set.seed(123)` before taking the sample.

```
bs_rows <- ____ (1:____, size = ____, replace = ____)
```

- **(5) Write and run code using the `slice` function to create a new dataset that includes each row from our sample.**

```
bs_atus <- slice(atus, bs_rows)
```

## Take a look

- Look at the values of `bs_rows` and `bs_atus`.
- (6) Write a paragraph that explains to someone who's not familiar with R how you created `bs_rows` and `bs_atus`. Be sure to include an explanation of what the values of `bs_rows` mean and how those values are used to create `bs_atus`. Also, be sure to explain what each argument of each function does.

## One strap, two strap

- (7) Write and run code calculating the mean of the `age` variable in your *bootstrapped* data, then use a different value of `set.seed()` to create your own, personal *bootstrapped* sample. Then calculate its mean.
- (8) Compare this second *bootstrapped* sample with three other classmates and write a sentence about how similar or different the *bootstrapped* sample means were.

## Many bootstraps

- To use *bootstrapped* samples to create *confidence intervals*, we need to create many *bootstrapped* samples.
  - Normally, the more *bootstrapped* samples we use, the better the *confidence interval*.
  - In this lab, we'll `do()` 500 *bootstrapped* samples.
- To make `do()`-ing 500 *bootstraps* easier, we'll code our 3-step bootstrap method into a function.
  - Open a new R Script (File -> New File -> R Script) to write your function into.

## Bootstrap function

- (9) Fill in the blank space below with the 3-steps needed to create a *bootstrapped* sample mean for our `atus` data.
  - Each step should be written on its own line between the curly braces.

```
bs_func <- function() {

}
```

- Highlight and Run the code you write.

## Visualizing our bootstraps

- (10) Once your function is created, fill in the blanks to create 500 *bootstrapped* sample means:  

```
bs_means <- do(____) * bs_func()
```
- (11) Create a histogram for your *bootstrapped* samples and describe the center, shape and spread of its distribution.
  - These bootstrapped estimates no longer estimate the average age of people in the U.S.
  - Instead, they estimate how much the estimate of the average age of people in the U.S. varies.
- In the next slide, we'll look at how we can use these *bootstrapped* means to create 90% *confidence intervals*.

## Bootstrapped confidence intervals

- To create a 90% confidence interval, we need to decide between which two *ages* the middle 90% of our bootstrapped estimates are contained.
- (12) Using your **histogram**, fill in the statement below:  
The lowest 5% of our estimates are below \_\_\_\_\_ years and the highest 5% of our estimates are above \_\_\_\_\_ years.
- (13) Write and run code using the **quantile()** function to check your estimates.
- (14) Based on your *bootstrapped* estimates, between which two ages are we 90% confident the actual **mean** age of people living in the U.S. is contained?

## On your own

- (15) Using your *bootstrapped* sample means, write and run code creating a 95% confidence interval for the **mean** age of people living in the U.S.
- (16) Why is the 95% confidence interval wider than the 90% interval?
- (17) Write down how you would explain what a 95% confidence interval means to someone not taking *Introduction to Data Science*.

### **Practicum: Let's Build a Survey!**

#### **Objective:**

Students will design a survey that has non-leading questions.

#### **Materials:**

1. Practicum: *Let's Build a Survey!* (LMR\_U3\_Practicum\_Build\_a\_Survey)

### **Practicum Let's Build a Survey!**

Based on what you have learned in Lessons 9 through 14, you will now design a survey. You and your team members must do all of the following:

1. Select a topic from the list below:
  - a. Social Media
  - b. Entertainment
  - c. Sports
  - d. The Environment
  - e. Health
  - f. Education
  - g. Other topic of interest
2. Create a research question about your topic of interest.
3. Create a statistical question that is related to the research question.
4. Identify the population of interest.
5. Describe how you will select your sample from the population so that you'll be able to make generalizations about your population of interest.
6. Identify the number of people who will be in your sample.
7. Create five survey questions that will try to answer your statistical question and describe how you have made sure that they are non-leading questions.
8. Identify a statistic that can be used to summarize the responses from this survey. Can you identify a parameter?
9. Submit a typed paper that details the survey you just designed.

# What's the Trigger?

Instructional Days: 5

## Enduring Understandings

Sensors are data collection devices that collect data either continuously or whenever they are triggered. A sensor is a converter that measures a physical quantity and converts it into a signal, which can be read by an observer or by an instrument. Participatory Sensing is a specific data collection method that uses sensor technology. This method emphasizes the involvement of citizens and community groups in the process of sensing and documenting where they live, work, and play. Triggers play an important role in the Participatory Sensing data collection process. The response to the triggers may or may not be the same each time.

## Engagement

Students will view and discuss a video clip called *Play Like Nadal With a Smart Tennis Racket* to begin to think about the sensors as data collection devices found ubiquitously in today's world. The video can be found at: <https://youtu.be/lcBnzddQECc>

## Learning Objectives

### Statistical/Mathematical:

S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.

S-IC 6: Evaluate reports based on data.

### Data Science:

Understand that sensors provide a continuous stream of data. Participatory Sensing provides real-time data from a user who is willingly providing the data. What differentiates a sensor as a data gathering method is the use of a trigger that signals a data collection session.

### Applied Computational Thinking:

- Create a Participatory Sensing campaign using a campaign Authoring Tool.

### Real-World Connections:

Sensors are found everywhere in today's world. They can provide data about environmental conditions as well as personal habits. More and more, sensors are being used for personal tracking, especially in the medical field, to inform people about what they do.

## Language Objectives

1. Students will engage in a "Gallery Walk" activity of sensors and how they are triggered.
2. Students will determine when it is appropriate to use surveys or sensors.
3. Students will engage in partner and whole group discussions to create a class Participatory Sensing campaign.

## Data File or Data Collection Method

### Data File:

1. Students' Participatory Sensing campaign data

### Data Collection Method:

**Class-generated Participatory Sensing Campaign:** Students will collect data on a class-selected topic.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## **Lesson 15: Ready, Sense, Go!**

### **Objective:**

Students will learn what sensors are and how they are used to collect data.

### **Materials:**

1. Video: *Play Like Nadal With a Smart Tennis Racket*  
<https://youtu.be/lcBnzddQECc>
2. Computers (see step 5 in lesson)
3. Poster paper
4. Flags in 3 different colors

**Advance preparation required:** Need to set up groups/activity (see step 10 in lesson)

### **Vocabulary:**

sensor, trigger, algorithm

**Essential Concepts:** Sensors are another data collection method. Unlike what we have seen so far, sensors do not involve humans (much). They collect data according to an algorithm.

### **Lesson:**

1. *Entrance Ticket:* What are some of the data collection methods we have learned about so far in this unit? *Answer: We have learned about experiments, observational studies, surveys, and getting data from a URL (in Lab 3B).*
  2. Inform students that, in this lesson, they will be introduced to another data collection method known as sensors.
  3. With a partner, ask students to discuss what they think a sensor is. Ask each pair to write down their ideas.
  4.  Show the *Play Like Nadal With a Smart Tennis Racket* video found at: <https://youtu.be/lcBnzddQECc>. As students watch the video, they should think about other sensors they may have come across, particularly ones used with smartphones. After watching the video, ask students to add to their definition of a sensor.
  5. Now, inform students that they will work in teams to compile a list of data-collecting sensors. They may use computers to conduct online research for this part of the lesson. Challenge each team to generate the longest list in the class.
  6. After students have had time to research and create their lists, ask students in each team to number off one through four (or five, depending on team sizes).
  7. Share out in rounds. First, ask students in each team whose number is one to share one sensor from their list. On the poster paper, create a class list of sensors as shared by the students. Repeat with the rest of the numbers.
  8. Score keeping: Each person gets five seconds to respond. You may hold up your hand with the palm facing the students and count down. The rules for teams are as follows:
    - a. add a sensor to the list, get 1 point
    - b. repeat an answer, lose 1 point
    - c. do not answer in five seconds, lose a turn
    - d. do not have an answer to contribute, may pass
- Note:** You may reward the winning team with extra credit points, if desired.
9. Next, students will engage in an activity to see sensors in action.

10. Create 3 groups of students:

a. Group 1 – Triggers (3 students)

Provide each Trigger a different colored flag (for example: **Pink**, **Purple**, **Green**). The teacher will call out a color, at random, and the Trigger assigned to that color will raise his or her flag. Each flag corresponds to a research question of interest.

**Pink** – Who is in our class?

**Purple** – What is on our classroom walls?

**Green** – What do we like to do after school?

b. Group 2 – Sensors (2 students)

Each Sensor should be assigned to one Trigger, or colored flag (**Pink** = Sensor A, **Purple** = Sensor B, **Green** = no sensor assigned to it). When the Sensors see their assigned Trigger, they send a signal to the Collector (see below) telling him or her to collect data from another student in the class. The Sensors are basically go-betweens for the Triggers and the Collector.

c. Group 3 – Collector (1 student)

One student is the Collector of all the data. The Collector is in charge of asking survey questions related to the research question of the original Trigger. Survey questions are provided here:

Survey questions related to the **Pink** trigger:

- (1) How did you get to school today (bus, car, walking, etc.)?
- (2) What size shoe do you wear?
- (3) What is your favorite pizza topping?

Survey questions related to the **Purple** trigger:

- (1) What is your favorite wall decoration?
- (2) What type of poster is it (motivational, reference, class work, etc.)?
- (3) What color is most prevalent in the poster?

Survey questions related to the **Green** trigger:

- (1) Do you have a sports team practice or club meeting today?
- (2) Are you hanging out with friends today after school?
- (3) Will you be working today after school?

Each time a sensor is active, the Collector must ask a new student in the class the appropriate survey questions.



11. Explain the activity to your students. Then, call out a flag color at random. Repeat several times. Make sure you call out the flag that has no assignment at least once so that students see that no action took place. Reflect on the activity with the following discussion questions:

- a. What data were missed? Why? *Answer: Data about what our class likes to do after school. They were missed because there was no Sensor connected to the Green trigger, so the Collector never knew to collect this type of data.*
- b. Grocery stores keep track of customer data when purchases are made with a loyalty card. What is the trigger in this case? What data are being collected? *Answer: The trigger is checking out at a grocery store. There are lots of data that are collected, including: items bought, cost of items, number of items on sale, etc.*

12. After they engage in the sensor activity, ask students to revisit their definition of a sensor (see step 3 above). Have them revise their definition based on the following concepts:

- a. A **sensor** is a converter that measures a physical quantity and converts it into a signal, which can be read by an observer or by an instrument.
- b. A sensor collects data continuously, or whenever a **trigger** is activated. A trigger is a something that responds to an event so that an action can occur.

- c. Sensors collect data according to an **algorithm**. An algorithm is a process or set of rules that are followed (just like the rules followed during the activity).
- d. Sensors may also collect data automatically, without anyone's knowledge or input. Examples include GPS location, time, and date.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework



Now that students learned what sensors are, ask them what data they would like to see collected on a sensor that they couldn't collect in an experiment or survey. They must explain why it is difficult to collect that data in an experiment or survey, and how a sensor would make it easier to collect that data.

## Lesson 16: Does It Have a Trigger?

### Objective:

Students will learn to identify and categorize survey questions versus sensor questions and will practice writing sensor questions.

### Materials:

1. Poster paper (one per student team)
2. Sticky notes
3. *Sensor or Survey?* handout (LMR\_U3\_L16)

### Vocabulary:

Participatory Sensing

**Essential Concepts:** A key feature that distinguishes the way sensors collect data from more traditional approaches is that sensors collect data when a 'trigger' event occurs. In Participatory Sensing, this event is something we humans agree upon beforehand. Every time that trigger happens, we collect data.

### Lesson:

1. Refer back to the list of sensors the class created during the previous lesson. Distribute a piece of poster paper to each student team, and have them create the following table:

| Sensor | How is it triggered? |
|--------|----------------------|
| 1.     | 1.                   |
| 2.     | 2.                   |
| 3.     | 3.                   |

2. Assign each student team 3 sensors from the class's list.
3. Then, each team should complete the table using their knowledge of triggers discussed during the previous lesson. Remind students that when a trigger occurs, a sensor reacts to it and sends a signal to a data collector.
4. Conduct a *Gallery Walk* of the posters. Each team will get to write one reaction or question about what they see on each poster.
5. After the Gallery Walk, ask each team to return to their posters. If the posters include questions, have teams take turns responding to the questions.
6. *Quickwrite:* In their DS Journals, ask student to respond to the following questions. They will have two minutes to write as much as they can:
  - a. When you learned about survey questions, what were the two categories of questions you learned about? *Answer: Open-ended and Closed-ended are the categories.*
  - b. What are some examples of these types of questions? *Possible answers: Open-ended: write a paragraph/sentence, comments, essays, single answer. Closed-ended: multiple/single choice, yes/no, scales (e.g. 1-5), choose from a list, check a box.*
7. In teams, ask students to share their responses using the *Give One/Get One* strategy. You may use a timer to keep track of time.
8. Remind students that one of the most important things they learned about sensors is that there is a trigger that reminds either a device or a person to answer a question or to collect data.
9. For this class, students have already had experience with using sensors as a data collection tool – all the **Participatory Sensing** campaigns. Participatory Sensing is an approach to data collection and interpretation in which individuals, acting alone or in groups, use their personal mobile devices and web services to systematically explore interesting aspects of their worlds ranging from health to culture.



10. Explain that survey questions are asked in Participatory Sensing (PS) campaigns. There is no difference in the type of questions that are asked when collecting data via surveys and when collecting data via PS campaigns.
11. When deciding whether to use a survey or a PS campaign for data collection, we have to look at the research question of interest. Some questions are better answered with survey data, while others with PS campaigns. Research questions that include variation across time or across locations are good candidates for PS. Some questions might be answered by both.

**For example:**

Consider the research question: How does my sense of safety and security change as I go about my daily routine? This question would best be answered via a PS campaign because students could collect data in real time about their sense of security. A possible trigger could be “when you change locations,” “once at the start of every hour,” or perhaps when a random alarm goes off.

Consider the research question: What proportion of high school students are superstitious? This question could be done with a survey based on a random sample from the population of all high school students.



12. Distribute the *Sensor or Survey?* (LMR\_U3\_L16) handout. In teams, students will determine whether a sensor or survey is better for a given research scenario.

**Note:** Page 2 of the handout is an answer key for teacher reference.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Sensor or Survey?**

Instructions:  
For each scenario in the table below, identify which data collection method is more appropriate: a sensor (Participatory Sensing campaign) or a survey. Include your reasoning in the appropriate column.

| Research Scenario                                                                                                 | Sensor or Survey | Why did you choose this method? |
|-------------------------------------------------------------------------------------------------------------------|------------------|---------------------------------|
| You want to know the percentage of students in the school district who complete all of their homework each night. |                  |                                 |
| You want to know how many times per day, and at what times, students in the district play sports.                 |                  |                                 |
| You want to know what foods your class eats most often.                                                           |                  |                                 |

LMR\_U3\_L16

13. Once the teams have completed the handout, assign each team one research scenario from the *Sensor or Survey* activity.
14. Conduct a *Whip Around* and have each team share their responses with the class. Allow students time to revise any incorrect responses.
15. Summarize the lesson by highlighting that PS campaigns and surveys use similar questions. However, depending on the research topic of interest, the decision to use one or the other relies on whether or not a trigger is involved.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Homework



- Suppose we wish to know more about whether people behave superstitiously. Write two research scenarios, using the following questions as a guide:
- a. How would you collect data to address this using PS? Include the trigger event you would use, and the data you would like to collect when the trigger happens.
  - b. How would you collect data to collect this using a survey based on a random sample of people in California?
  - c. Describe the differences between these two approaches. What can you learn in one approach that you can't in the other?

## Lesson 17: Creating Our Own Participatory Sensing Campaign

### Objective:

Students will be guided through the creation of a new Participatory Sensing campaign and survey on a topic of interest chosen by the class.

### Materials:

1. *Food Habits Campaign Questions* handout (LMR\_U3\_L17\_A)
2. *Campaign Creation Brainstorm* handout (LMR\_U3\_L17\_B)

**Essential Concepts:** Creating a Participatory Sensing Campaign requires that survey questions must be completed whenever they are “triggered.” Research questions provide an overall direction in a Participatory Sensing Campaign.

### Lesson:



1. Review homework questions by asking a couple of students to share their responses. The rest of the students will engage in *Agree/Disagree* as the questions are shared.
2. Display the following definition of Participatory Sensing that some computer scientists have agreed to and ask students to read and record this definition in their DS journals:

At its heart, Participatory Sensing is data collection and interpretation. Participatory Sensing emphasizes the involvement of citizens and community groups in the process of sensing and documenting where they live, work, and play. It can range from private personal observations to the combination of data from hundreds, or even thousands, of individuals that reveals patterns across an entire city. Most important, Participatory Sensing begins and ends with people, both as individuals and members of communities. The type of information collected, how it is organized, and how it is ultimately used, may be determined in a traditional manner by a centrally organized body, or in a deliberative manner by the collection of participants themselves. The latter case, in particular, emphasizes the novelty of Participatory Sensing as an approach and underscores the importance of using widely available and familiar technology. [Source: “Participatory Sensing: A citizen-powered approach to illuminating the patterns that shape our world.”]

3. Activate prior knowledge: Based on this definition, ask students to recall the Participatory Sensing campaigns in which they have engaged thus far. **Answer:** *Food Habits, Time Use, Stress/Chill*.

**Note:** *Personality Color* and *Time Perception* were surveys, not Participatory Sensing campaigns because they were only completed once. Their data was not collected over time.

4. Inform students that they will be creating a new, whole class Participatory Sensing campaign, but before they do that, they will analyze the *Food Habits Campaign Questions* handout (LMR\_U3\_L17\_A).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Food Habits Campaign Questions

| Prompt                                                              | Variable      | Data Type                                                     |
|---------------------------------------------------------------------|---------------|---------------------------------------------------------------|
| What is the name of your snack?                                     | name          | text                                                          |
| When did you eat the snack?                                         | when          | categorical<br>morning<br>afternoon<br>evening<br>night       |
| Is your snack salty or sweet?                                       | salty_sweet   | categorical<br>Salty<br>Sweet                                 |
| How healthy is the snack?<br>(1 = Very unhealthy, 5 = Very healthy) | healthy_level | categorical<br>1<br>2<br>3<br>4<br>5                          |
| How many calories per serving?                                      | calories      | numerical                                                     |
| How many grams of protein per serving?                              | protein       | numerical                                                     |
| How many grams of sugar per serving?                                | sugar         | numerical                                                     |
| How many milligrams of sodium per serving?                          | sodium        | numerical                                                     |
| How many ingredients are in your snack                              | ingredients   | numerical                                                     |
| Why are you eating this snack?                                      | why           | categorical<br>availability<br>craving<br>emotional<br>energy |

LMR\_U3\_L17\_A



5. In teams, allow students two minutes to discuss the following as they analyze the Food Habits Campaign questions:
  - a. How many questions does the campaign have and what do they notice about the questions? *Answers will vary. Students may notice that they are survey type of questions and may identify the type of questions such as open-ended, single-choice, etc.*
  - b. When do these questions need to be answered? *Answer: Each time they eat a snack.*
  - c. Who collects the data for this campaign? *Answer: The participants collect their own data.*
6. Ask a few teams to share their insights about the discussion. In the share-out, guide students to see that the questions are in fact survey questions. Although survey questions are answered once, when we collect data every time a 'trigger' event occurs, then we are engaging in Participatory Sensing.
7. Ensure that team roles have defined duties to keep teams on task for the rest of this lesson. Creating this class campaign will follow a process in which consensus (or a majority rule) will be reached in each step of the campaign development within each team. Inform students that they will be creating a Participatory Sensing campaign on a topic of their interest using LMR\_U3\_L17\_B.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Campaign Creation Brainstorm**

**Instructions:**  
In teams, work together to fill in the information in this handout. You will be deciding, as a class, what information will be used in your class campaign during each round.

**Round 1: Topic**  
*This is a hobby, area of interest, or place or process that you want to know more about.*

Team Ideas of Topics:

\_\_\_\_\_

\_\_\_\_\_

\_\_\_\_\_

\_\_\_\_\_

Class Decided Topic:

\_\_\_\_\_

**Round 2: Research Question**  
*This is the main question you want to answer about the topic and will be the focus of the Campaign.*

LMR\_U3\_L17\_B

8. **Round 1:** First, teams will discuss their hobbies, areas of interest, or places or processes they want to know more about. Prompt students to think about whether they want to learn about "where they live, work, or play." All students within the group must agree on a hobby or area of interest to be their topic of interest to create a campaign for. *An example of a hobby is practicing cello. An area of interest might be 'the environment.' A place of interest might be "our school" or "my church" or "Disneyland."*
9. Once teams have decided on a topic for their group, have teams share out their topic of interest. As a class, decide on one topic that will be used for creating a new Campaign.
10. **Round 2:** Now have teams consider what research questions you might ask about this topic of interest. *An example of a research question for practicing cello is "How can I improve my playing?" or "How can I practice more effectively?"*
11. Once teams have decided on a research question for their group, have teams share out their research question. As a class, decide on one research question that will be used for creating a new Campaign.
12. **Round 3:** Next, they will examine what kind of data needs to be collected in order to answer this research question. They should discuss possible triggers that will determine when data should be collected. Allow teams to engage in a discussion about when is the best time to trigger the data collection/completion of the survey. For example: every day at 8am; whenever they practice the cello; whenever they see an advertisement; etc. They should record it in their DS Journals. *An example of a trigger for practicing cello is whenever you play the cello. In this case, it could be any time of day or even multiple times of the day.*

13. Once teams have decided on a trigger for their group, have teams share out their possible trigger. As a class, decide on one trigger that will be used for creating a new Campaign.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework



Using the class topic, research question, trigger event, and discussion of the data they plan to collect. Classify our class campaign under the appropriate category with your justification:

(A) Individual; (B) Groups of people; (C) Community

**Note to teacher:** To determine which category a campaign should be placed under, consider the question “Who or what will we learn about?” If the answer is “only one person,” then place in the Individual category. The cello campaign is an example of this. If we might learn about lots of people, put it in the Groups of People category. The Food Habits, Stress/Chill, and Time Use campaigns fit into this category (they learn about the students in the class). A campaign that wanted to know where all of the churches in the neighborhood were located, or wanted to try to keep people from littering, or wasting water, these should go into the “community” category.

## Lesson 18: Evaluating Our Own Participatory Sensing Campaign

### Objective:

Students will create statistical questions and evaluate their Participatory Sensing Campaign.

### Materials:

1. *Campaign Creation Brainstorm* handout (LMR\_U3\_L17\_B) from previous lesson
2. Class Campaign Information from Lesson 16

**Essential Concepts:** Statistical questions guide a Participatory Sensing Campaign so that we can learn about a community or ourselves. These campaigns should be evaluated before implementing to make sure they are reasonable and ethically sound.

### Lesson:

1. Review homework by giving students about five minutes to share their classifications in their teams. They will decide as a team which classification is the most fitting.
2. Once the five minutes have passed, have a class discussion of classifications and their justifications. Explain to the class that the campaign must be carried out by the whole class so if it has been classified in the Individual category, it must be revised. Also discuss whether the campaign is feasible. (For example, is the trigger so rare that no one will collect data? Are the questions too intrusive?)
3. Inform students that one of the promises of PS is its potential for helping people bring about social and civic change. Ask teams to consider the following questions and report back:
  - a. Does our campaign try to do this?
  - b. Could it be changed or modified to do this?

**Note:** Feasible campaigns fall under the groups of people or community categories. If a campaign is in the individual category, it should be modified to fall under the other categories before moving to round 4.

4. Display the campaign information students generated (and selected as a class) the previous day or revised today: Topic, Research question, Trigger, and Type of Data needed.
5. Now they will continue the rounds using the Campaign Creation Brainstorm handout (LMR\_U\_L17\_B) from the previous lesson.
6. Round 4: Now that the class has decided on a trigger and the type of data needed, they will create survey questions to ask when the trigger is set. The questions should consider all of the possible data they might collect at this trigger event. It's ok if the list is long; the goal is to be creative and think of lots of different ideas.

*Examples of survey questions for practicing cello are:*

*"How long did you practice?"*

*"What did you play?"*

*"How would you rate your practice session: 1 to 5?"*

*"Any thoughts or comments about your practice?"*

7. Once teams have created 4 survey questions for their group, have teams share out their survey questions. As a class, decide on no more than 10 survey questions that will be used for creating a new Campaign.
  - a. Then, evaluate each survey question. For each question they should consider:
    - i. What type of data will this question collect? (Numerical, discrete numerical, text, categories, photos, location).
    - ii. How does this question help address the research question?
    - iii. Does the question need to be reworded? (Is it clear what is being asked for? Do they know how to answer it?)

- b. If the survey questions need to be rewritten, assign teams to rewrite survey questions. Then, as a class, decide on the changes.
  - c. Once finalized, write the survey question that goes along with that data variable, being cognizant of question bias.
8. **Round 5:** In teams, now generate two to three statistical questions that they might answer with these data. Make sure your statistical questions are interesting and relevant to the class topic of interest. They may keep a record in their DS Journals. Remind students that they will also have data about the date, time, and place of data collection.

*Examples of statistical questions that can be answered for practicing cello are:*

*"How frequently do I practice?"*

*"When I practice more frequently, do I rate my sessions higher?"*

*"Are higher-rated sessions associated with time of day?"*

9. Once teams have generated their statistical questions, have them share out with the class. Confirm that the questions are statistical and that they can be answered with the data the students propose to collect.



10. Now that they have all the pieces of the campaign, evaluate whether it's a reasonable and ethically sound campaign. Engage the class in a whole group discussion on the following questions:
- a. Are answers to your survey questions likely to vary when the trigger occurs? (If not, you'll get bored entering the same data again and again.)
  - b. Can the entire class carry out the campaign?
  - c. Do triggers occur so rarely that you'll have very little data? Do they occur so often that you'll get frustrated entering too much data?
  - d. Ethics: Would sharing these data with strangers or friends be embarrassing or undermine someone's privacy?
  - e. Can you change your trigger or survey questions to improve your evaluation?
  - f. Will you be able to gather enough relevant data from your survey questions to be able to answer your statistical questions?
11. Students have collaboratively created their first Participatory Sensing campaign. Inform them that you will be demonstrating one tool used to create the campaigns that they see on their smart devices or the computer. Students should take notes in their DS journals, as they will be using the tool later.

#### **Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 19: Implementing Our Own Participatory Sensing Campaign

### Objective:

Students will mock-implement, create their Participatory Sensing campaign, survey on their topic of interest, then begin data collection.

### Materials:

1. Campaign Creation Brainstorm handout (LMR\_U3\_L17\_B) from previous lesson
2. Campaign Authoring Tool (<https://portal.thinkdataed.org>)

**Essential Concepts:** Practicing data collection prior to implementation allows optimization of a Participatory Sensing Campaign.

### Lesson:

- 
1. Display the class generated campaign information for the class to clearly see.
  2. In teams, have students mock-implement the campaign they have created. They can do this by asking each other the survey questions to make sure they make sense/will generate relevant data to their research question and statistical questions. They can use the evaluative questions from Lesson 18 step 10.
  3. If there are suggestions for improvement, have teams propose them to the class and make final changes to the campaign.
  4. Inform students that you will now demonstrate the tool used to create the campaigns that is displayed on their mobile devices or computers.
  5. Login to the **IDS Home Page** found at <https://portal.thinkdataed.org>. Click on the **Campaigns** tab on the navigation bar at the top of the page. Then, follow the steps in the tool:

**Note:** If you would like a video tutorial to assist with the written instructions, it can be found here: <https://youtu.be/PzwMCHOghnl>

#### a. Campaign Info:

- i. **Campaign Name:** Give your campaign a name. A name related to the topic is recommended.
- ii. **Select your class/period.**
- iii. **Description:** Provide a one-sentence description of your campaign.
- iv. **Campaign Status:** Select Running.
- v. **Data Sharing:** Select Disabled in order to monitor for improper responses.
- vi. **Editable Responses:** Select Disabled.
- vii. Click the **+Add Survey** button.

#### b. Survey Window:

- i. **Title:** Give the survey a title (again, it may or may not be the same as the campaign name). Users see the title and all the prompts that follow.
- ii. **ID:** Give the survey a name (it may or may not be the same as the campaign name). Users do not see the survey ID.
- iii. **Description:** Provide a short description of the survey for display (optional – may be the same as the **Description** in **Campaign Info**).
- iv. **Submission Message:** Provide a brief message to be displayed after survey submission. *Note: it is helpful to include a reminder to click the green button to submit the survey.*
- v. Click the **+Add Prompt** button and select the prompt type for your first survey question.

- c. **Prompt Information:**
  - i. **Prompt ID:** This will be your first variable. A short one-word name or short two-word name separated by an underscore is recommended.
  - ii. **Prompt Label:** This is the variable name that will be displayed (it may be the same as the prompt ID without the underscore, if used).
  - iii. **Question Text:** Type the survey question about which you want to collect data.
  - iv. **Additional Prompt Information:** Depending on the prompt type, you will be asked to enter additional information. For example, if your prompt is Text, you will be asked a minimum and a maximum value for the number of characters the participant can enter.
  - v. **Skippable:** Select the checkbox if you would like the prompt to be skipped. It is recommended that photo prompts be skippable, since some users will submit their responses via a browser.
- d. Repeat step c for the remaining survey questions by clicking the **+Add Prompt** button.
- e. **XML Code:** As you create the campaign, the code that creates it will be displayed. You may select the checkbox titled **Highlight XML** so that you can keep track of where the information you are adding is embedded in the code. You will learn more about XML syntax in the next section of Unit 3.
- f. Click the **Submit Campaign** button on the top, right hand side of the page once all prompts have been added. This action will send the campaign to the server. Users will now be able to submit surveys.
6. Once all prompts have been created, students may use their smart devices or login to the IDS Home Page to view the new campaign. Remember to **Refresh Campaigns**.
7. Students should go through the entire Participatory Sensing survey to see how their questions are displayed. They do not have to upload the survey.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework

For the next 5 days, students will collect data using their newly created Participatory Sensing campaign.

# Webpages

Instructional Days: 6

## Enduring Understandings

Data takes on a variety of forms online and requires a different style of representation.

## Engagement

Students will view a video clip about a data farm, specifically, Google's Street View Data Center to begin thinking about data formats and accessing data online. The video can be found at:

<https://www.engadget.com/2012-10-17-google-inside-data-centers.html>

## Learning Objectives

### Statistical/Mathematical:

S-IC 3: Recognize the purposes of and differences among sample surveys, experiments, and observational studies; explain how randomization relates to each.

S-IC 6: Evaluate reports based on data.

DS: Use different techniques to access data from the web and understand why different data representations are useful for different software platforms.

### Applied Computational Thinking using RStudio:

- Read data from xml and html table and convert to R data frames.
- Use latitude and longitude coordinates of mountain data and overlay it on a map.

### Real-World Connections:

Data from the web has been used to predict outbreaks of the flu and is a source of extremely rich datasets.

## Language Objectives

1. Students will compare and contrast appropriateness of CSV, HTML, or XML for different data needs.
2. Students will engage in partner and whole group discussions and presentations to express their understanding of how data is stored and viewed online.
3. Students will engage in discussions regarding internet research as it applies to data science.

## Data File or Data Collection Method

### Data File:

1. Students' Participatory Sensing campaign data

### Data Collection Method:

1. Students will scrape data from online HTML sources.
2. Students will gather data generated through a class-generated Participatory Sensing campaign.

## Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 20: Online Data-ing

### Objective:

Students will discover that data exists on the Internet in a variety of areas, formats, and for a variety of purposes.

### Materials:

1. Video: *Explore a Google Data Center with Street View* found at: <https://www.engadget.com/2012-10-17-google-inside-data-centers.html>
2. Wikipedia – Video Games handout (LMR\_U3\_L20\_A)
3. Wikipedia – Video Games – CSV Format handout (LMR\_U3\_L20\_B)
4. Online Data-ing handout (LMR\_U3\_L20\_C)

### Vocabulary:

data farm, HTML, tags

**Essential Concepts:** We stretch students' conception of data, to help them see that many web pages present information that can be turned into data.

### Lesson:

1. By a show of hands, ask students if they have ever heard of the term **data farm**. If any of them have, ask him or her to share what they know about it.
2. Inform students that a data farm is a physical space where high-capacity servers are placed to store large amounts of data.
3. Introduce the video titled *Explore a Google data center with Street View* found at <https://www.engadget.com/2012-10-17-google-inside-data-centers.html> by explaining that the data center they are about to see is one of these large data farms used to store vast amounts of data.
4. After students watch the video, have a class discussion using the following questions:
  - a. We have been talking about data for a few months now. How would you respond if someone asked you, "What are data?" *Answers will vary by class.*
  - b. What are some ways that we have stored data? *Possible answers are data frames in R, Excel spreadsheets, .csv files.*
5. Explain that one of the main ways data are distributed is through the Internet. Storing and sharing data on the Internet requires a different format than what we have seen. For example, Wikipedia has a page dedicated to the top video games.
6. Distribute the *Wikipedia – Video Games* handout (LMR\_U3\_L20\_A), and have students explain the information that the data table provides.



Name: \_\_\_\_\_ Date: \_\_\_\_\_

Wikipedia – Video Games

Background:  
The Wikipedia website contains many informative webpages. One such page is dedicated to the top video games of all time, which can be found at [https://en.wikipedia.org/wiki/List\\_of\\_video\\_games\\_considered\\_the\\_best](https://en.wikipedia.org/wiki/List_of_video_games_considered_the_best).

This screenshot from the website shows the first five rows of the "List of Best Games" data table.

| Year | Game                    | Genre            | Publisher   | Original platform(s) <sup>[d]</sup> | Ref. |
|------|-------------------------|------------------|-------------|-------------------------------------|------|
| 1971 | <i>The Oregon Trail</i> | Strategy         | MECC        | HP 2100                             | [1]  |
| 1972 | <i>Pong</i>             | Sports           | Atari, Inc. | Arcade                              | [2]  |
| 1977 | <i>Combat</i>           | Top-down shooter | Atari, Inc. | Atari 2600                          | [3]  |
| 1977 | <i>Zork</i>             | Adventure        | Infocom     | PDP-10                              | [4]  |
| 1978 | <i>Space Invaders</i>   | Shoot 'em up     | Taito       | Arcade                              | [5]  |

The data table did not look like this automatically though. Behind the scenes, HTML source code is used to create the table in a reader-friendly format.

The beginning lines of the HTML source code (**bolded**) give us information about the structure of the data table. The following lines are the source code for each row of the table. For example, the first video game listed on page 2 is *Pong*.

```
<table class="wikitable sortable" width="auto">
<tbody><tr>
<th scope="col">Year
</th>
<th scope="col">Game
```

LMR\_U3\_L20\_A

7. Once the students understand what the data table describes, walk them through the first portion of the **HTML**, or Hypertext Markup Language, source code (on page 1). Notice that the first header on the table is denoted as “Game”. Ask:
  - a. How is “Game” represented in the source code? *Answer: <th>Game</th>*
  - b. What do you think the <th> code represents? *Answer: The <th> is a tag for “table header.”*
  - c. If this were in RStudio, what would we call this header? *Answer: A variable.*
8. Assign each student team one video game from the Wikipedia data table. Each team will compare how the information is stored in the table with its corresponding HTML source code. Each team should answer the following questions in the DS journals.
  - a. Where are the variable names stored? *Answer: The variable names are stored at the beginning of the code, in between <th> and </th>. <th> and </th> are called **tags** – they tell the browser to represent the information between them as a header in the table.*
  - b. How are different values of the variables stored? *Answer: Values are stored between the <td> and </td> tags.*
  - c. Why do you think the data are stored in such a complex way? Why can’t we just put them in a spreadsheet? *Answers may vary by class. One reason is that the data must be displayed in a way that allows a browser to make it look pretty (and readable) on a computer screen.*
  - d. How could we get this into an R dataframe so we can analyze it? *Answer: In its current form, this would be very difficult. We would need to represent the data in a different format in order for R to understand it.*
9. Distribute the *Wikipedia – Video Games – CSV Format* handout (LMR\_U3\_L20\_B) and explain that this is yet another way to represent the same video game data.  
**Note:** The handout only provides information on the first 5 rows of the Wikipedia table. A full version of the file (including all video games in the table) is located on the server with the title bestgames.csv.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Wikipedia – Video Games – CSV Format

Background:  
The data below represent the same 5 video games from the Wikipedia webpage, but in CSV format.  
CSV stands for “comma separated values.”

```
"Year","Game","Genre","Publisher","Original platform(s)","Ref."
1971,"The Oregon Trail","Strategy","MECC","HP 2100","[A]"
1972,"Pong","Sports","Atari, Inc.","Arcade","[B]"
1977,"Combat","Top-down shooter","Atari, Inc.","Atari 2600","[C]"
1977,"Zork","Adventure","Infocom","PDP-10","[D]"
1978,"Space Invaders","Shoot 'em up","Taito","Arcade","[E]"
```

LMR\_U3\_L20\_B

10. Inform students that a file with the CSV format is easily readable by R. Then ask:
  - a. Where are the variable names stored? *Answer: The variable names are stored in the first row.*
  - b. How are values of the variables separated? *Answer: The values are separated by commas.*
  - c. If we were interested in using the online data, how would we obtain it? *Answer: This is a challenging problem – one which students may not know how to answer at this point. The objective is for them to struggle with how they would obtain data and recognize that it is not always as simple as “export, upload, import.”*



11. Split the class into their student teams and distribute the *Online Data-ing* handout (LMR\_U3\_L20\_C). Assign each team a different website (each page of the handout lists a different site) and have them use this site to complete the questions in the handout.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Online Data-ing

Instructions:  
Visit the assigned website and, with your team, answer each question below. If you do not see data at the top of the page, explore the website a bit to find some.

Assigned Website: Yelp – [www.yelp.com](http://www.yelp.com)

- Describe the data on this website.  
\_\_\_\_\_  
\_\_\_\_\_
- What variables are present?  
\_\_\_\_\_  
\_\_\_\_\_
- What type of values do the variables have (i.e., words, numbers, dates, places, categories, etc.)?  
\_\_\_\_\_  
\_\_\_\_\_
- Where do the values of the data come from?  
\_\_\_\_\_

LMR\_U3\_L20\_C



- Have each student team share their findings with one other team. They should have their website displayed while discussing their results.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework & Next 2 Days

For the next 4 days, students will collect data using their newly created Participatory Sensing campaign.

## **Lab 3E: Scraping web data**

## **Lab 3F: Maps**

Complete Labs 3E and 3F prior to Lesson 21.

### Lab 3E: Scraping web data

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

#### The web as a data source

- The internet contains huge amounts of information.
  - Using computers to gather this information in an automated fashion is referred to as *scraping web data*.
  - Scraping data from the web can be difficult because each website displays & stores data differently.
- In this lab, we'll learn how to scrape data in two steps:
  - Step 1: Gather information from the web.
  - Step 2: Clean it up and turn it into a usable data frame for Lab 3F.

#### Our first web scraper

- **Copy and paste the link below into a web browser to view the website of data we'd like to scrape and analyze.**  
<https://labs.thinkdataed.org/extras/webdata/mountains.html>
- **(1) Briefly describe what the data on the website is about.**
- **(2) Write down 3 questions you'd be interested in answering by analyzing this data.**

#### HTML

- **HTML** is the code that's used to render every website you've ever visited.
- The following slide shows the **HTML** code used to create the first two rows of the web data.
  - **(3) How is the data table in HTML different than the data tables we're used to seeing in R, for example, when we use the `View()` function?**
  - **(4) What do you think the tags `<TABLE>`, `<TR>`, `<TH>`, `<TD>` mean? How does HTML use these tags to display the table?**

```
<TABLE>
<TR>
 <TH>peak</TH>
 <TH>range</TH>
 <TH>state</TH>
 <TH>long</TH>
 <TH>lat</TH>
 <TH>elev_ft</TH>
 <TH>elev_m</TH>
 <TH>prominence_ft</TH>
 <TH>prominence_m</TH>
 <TH>rank</TH>
</TR>
<TR>
 <TD>Denali (Mount McKinley)</TD>
 <TD>Alaska Range</TD>
 <TD>Alaska</TD>
 <TD>-151.0063</TD>
 <TD>63.0690</TD>
 <TD>20236</TD>
 <TD>6168</TD>
 <TD>20174</TD>
 <TD>6149</TD>
 <TD>1</TD>
</TR>
</TABLE>
```

## Get to scraping!

- Use your browser to go back to the website with the data we're interested in scraping.
- Find the URL address for the site and assign it the name `data_url` in R.
- (5) Then fill in the blanks below to have R scrape every web table available on the site:

```
tables <- readHTMLTable(____)
```

## Find our data

- Since `readHTMLTable()` scrapes every table that is on a particular web URL, we need to find out which table has the data we're interested in.
  - For example, [wikipedia.org](http://wikipedia.org) often has articles with 3 or more tables.
  - This means we need to check all 3 tables to find the one we're interested in.
- (6) Write and run code using the `length()` function to find out how many tables of data were scraped in our set of `tables`.

## Saving tables

- Now that we know how many tables we've scraped, we can go back and scrape individual tables by adding the `which` argument to the `readHTMLTable()` function.
- (7) Write and run code using `readHTMLTable()` to re-scrape the data from the web but this time use the `which` argument to scrape just the individual table.
  - The `which` argument should be the integer denoting which table you want scraped.
  - Assign the scraped data the name `mtns`.

## Check, save and use!

- After scraping the data, the only thing left to do is to save it and use it.
- (8) Fill in the blanks to save the data and give it a file name.

```
save(____, file = "____.Rda")
```

- (9) What is the mean and standard deviation of `elev_ft`?
- (10) Which `state` has the most mountains in our data?

## Lab 3F: Maps

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Informative and Fun!

- Maps are some of the most interesting plots to make because the info represents:
  - Where we live.
  - Where we go.
  - Places that interest us.
- Maps are also helpful to display geographic information.
  - John Snow (the physician, not the character from *Game of Thrones*...) once famously used [a map to discover how cholera was transmitted](#).
- In this lab, we'll use **R** to create an interactive map of the **mtns** data we scraped in Lab 3E.

### Getting ready to map

- The map we'll be creating will end up in RStudio's *Viewer* pane.
  - Which means you'll need to alternate between building the map and loading the lab.
- You'll find it very helpful, for this lab, to write all of the commands, including the `load_lab(23)` command, as an **R** Script.
  - This way you can edit the code that builds the map and quickly reload the lab.

### Load your data!

- **In Lab 3E you created a dataset. Load it into RStudio now by filling in the blank with the file name of the data.**

```
load("____.Rda")
```

- Didn't finish the lab or save the data file? Ask a friend to share it!

### Build a Basic Map

- Let's start by building a basic map!
- **(1) Write and run code using the `leaflet()` function and the **mtns** data to create the **leaf** that we can use for mapping.**

```
mtns_leaf <- leaflet(____)
```
- **(2) Then, write and run code inserting **mtns\_leaf** into the `addTiles()` function and assign the output the name **mtns\_map**.**
- **Run **mtns\_map** in the console to look at your basic map with no data displayed.**
  - Be sure to try clicking on the map to pan and zoom.

### Including our data

- Now we can add markers for the locations of the mountains using the `addMarkers()` function.
- **(3) Fill in the blanks below with the basic map we've created and the values for latitude and longitude.**

```
addMarkers(map = _____, lng = ~_____, lat = ~_____)
```

- (4) Write and run code supplying the `peak` variable, in a similar way as we supplied the `lat` and `long` variables, to the `popup` argument and include it in the code above.
- (5) Click on a marker within your state of choice and write down the name of the mountain you clicked on.

## Colorize

- Our current map looks pretty good, but what if we wanted to add some colors to our plot?
- (6) Fill in the blanks below to create a new variable that assigns a color to each mountain based on the state it's located in.

```
mtns <- mutate(____, state_colors = colorize(____))
```

- Now that we've added a new variable, we need to re-build `mtns_leaf` and `mtns_map` to use it.
  - (7) Write and run code creating `mtns_leaf` and `mtns_map` as you did before.
  - Then change `addMarkers` to `addCircleMarkers` and keep all of the arguments the same.

## Showing off our colors

- To add the colors to our plot, use the `addCircleMarkers` like before but this time include `color = ~state_colors` as an argument.
- It's hard to know just what the different colors mean so let's add a legend.
  - First, assign the map with the circle markers as `mtns_map`.
  - (8) Then, fill in the blanks below to place a legend in the top-right hand corner.

```
addLegend(____, colors = ~unique(____), labels = ~unique(____))
```

## Lesson 21: Learning to Love XML

### Objective:

Students will understand the need for data to be stored in different ways – specifically, why it makes sense for web data to be formatted as XML.

### Materials:

1. *Online Data-ing* handout (LMR\_U3\_L20\_C)  
**Note:** This should have been completed during the previous class.
2. Mountain Peak data found at:  
<https://labs.thinkdataed.org/extras/webdata/mountains.html>  
**Note:** Open with Google Chrome or Firefox browsers, NOT with Safari.
3. *Mountains – HTML vs. XML* handout (LMR\_U3\_L21)  
**Note:** An electronic copy should be provided to students (see steps 7 and 13 in lesson).
4. XML Viewer found at:  
<https://jsonformatter.org/xml-viewer>
5. HTML Viewer found at:  
<https://jsonformatter.org/html-viewer>
6. Projector

### Vocabulary:

XML

**Essential Concepts:** XML is a programming language that we use with our campaigns. We create basic XML “tags” in the code, which help us store data in a format we understand.

### Lesson:

1. Allow time for student teams to present their findings from the *Online Data-ing* handout (LMR\_U3\_L20\_C) if there was not sufficient time during the previous lesson.
2. Remind students that in the previous lesson they learned about a variety of ways that data can be stored and presented online.
3. They’ve been working with comma separated (CSV) files and R data frames. In the previous lesson and lab they worked with HTML tables. Today they are going to learn how the data in tables built with HTML can be stored using XML.
4. **XML**, or Extensible Markup Language, is a popular format for storing data on the Internet. It allows programmers to easily update values in the data table if those values change.
5. In pairs, ask students to brainstorm ways in which data that is found online is different than the way we see data in RStudio. Then, create a class brainstorm from the student pair responses.
6. After the brainstorm, emphasize the following:
  - a. RStudio’s default way to work with data is as large data frames (tables) where rows represent observations and columns represent variables.
  - b. Data that is viewed online often has a different structure.
  - c. Data structures found on the web might be displayed in tables, such as those on Wikipedia, or streams, such as Twitter, and might even include data spread across multiple sections of a web page, such as Yelp.

Show students, on a projector, the Mountain Peak data found at:

<https://labs.thinkdataed.org/extras/webdata/mountains.html>

Ask students to look at the data and determine if they have seen it before. Hint: They have! It was the data they scraped during Lab 3E.



7. Once students figure out that the table is just the same data as the website they scraped during Lab 3E, distribute the *Mountains – HTML vs. XML* handout (LMR\_U3\_L21), which displays both HTML and XML versions of the data.

**Note:** Provide an electronic copy of LMR\_U3\_L21 for students to easily copy/paste the code later in the lesson. The handout only includes the first 3 mountains.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Mountains – HTML vs. XML**

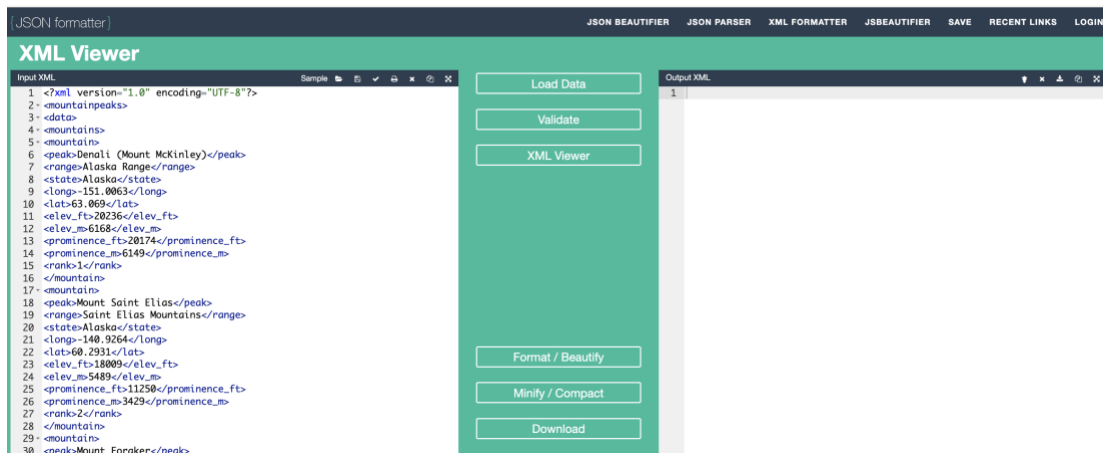
**Background:**  
The Mountain Peak data are displayed below in two different formats – the first is HTML and the second is XML.

<b>HTML Format</b>	<b>XML Format</b>
<pre>&lt;table border="1"&gt; &lt;tr&gt; &lt;th&gt;peak&lt;/th&gt; &lt;th&gt;range&lt;/th&gt; &lt;th&gt;state&lt;/th&gt; &lt;th&gt;long&lt;/th&gt; &lt;th&gt;lat&lt;/th&gt; &lt;th&gt;elev_ft&lt;/th&gt; &lt;th&gt;elev_m&lt;/th&gt; &lt;th&gt;prominence_ft&lt;/th&gt; &lt;th&gt;prominence_m&lt;/th&gt; &lt;th&gt;rank&lt;/th&gt; &lt;/tr&gt; &lt;tr&gt; &lt;td align="center"&gt;Denali (Mount McKinley)&lt;/td&gt; &lt;td align="center"&gt;Alaska Range&lt;/td&gt; &lt;td align="center"&gt;Alaska&lt;/td&gt; &lt;td align="center"&gt;-151.0063&lt;/td&gt; &lt;td align="center"&gt;63.0688&lt;/td&gt; &lt;td align="center"&gt;20326&lt;/td&gt; &lt;td align="center"&gt;6188&lt;/td&gt; &lt;td align="center"&gt;20324&lt;/td&gt; &lt;td align="center"&gt;6149&lt;/td&gt; &lt;td align="center"&gt;1&lt;/td&gt; &lt;/tr&gt; &lt;tr&gt; &lt;td align="center"&gt;Mount Saint Elias&lt;/td&gt; &lt;td align="center"&gt;Saint Elias Mountains&lt;/td&gt; &lt;td align="center"&gt;Alaska&lt;/td&gt; &lt;td align="center"&gt;-138.2431&lt;/td&gt; &lt;td align="center"&gt;58.3026&lt;/td&gt; &lt;td align="center"&gt;14992&lt;/td&gt; &lt;td align="center"&gt;4484&lt;/td&gt; &lt;td align="center"&gt;14960&lt;/td&gt; &lt;td align="center"&gt;2&lt;/td&gt; &lt;/tr&gt; &lt;tr&gt; &lt;td align="center"&gt;Saint Elias&lt;/td&gt; &lt;td align="center"&gt;Saint Elias Mountains&lt;/td&gt; &lt;td align="center"&gt;Alaska&lt;/td&gt; &lt;td align="center"&gt;-138.2431&lt;/td&gt; &lt;td align="center"&gt;58.3026&lt;/td&gt; &lt;td align="center"&gt;14992&lt;/td&gt; &lt;td align="center"&gt;4484&lt;/td&gt; &lt;td align="center"&gt;14960&lt;/td&gt; &lt;td align="center"&gt;2&lt;/td&gt; &lt;/tr&gt; &lt;/table&gt;</pre>	<pre>&lt;?xml version="1.0" encoding="UTF-8"&gt; &lt;mountainpeaks&gt; &lt;data&gt; &lt;mountain&gt; &lt;peak&gt;Denali (Mount McKinley)&lt;/peak&gt; &lt;range&gt;Alaska Range&lt;/range&gt; &lt;state&gt;Alaska&lt;/state&gt; &lt;long&gt;-151.0063&lt;/long&gt; &lt;lat&gt;63.0688&lt;/lat&gt; &lt;elev_ft&gt;20326&lt;/elev_ft&gt; &lt;elev_m&gt;6188&lt;/elev_m&gt; &lt;prominence_ft&gt;20324&lt;/prominence_ft&gt; &lt;prominence_m&gt;6149&lt;/prominence_m&gt; &lt;rank&gt;1&lt;/rank&gt; &lt;/mountain&gt; &lt;mountain&gt; &lt;peak&gt;Mount Saint Elias&lt;/peak&gt; &lt;range&gt;Saint Elias Mountains&lt;/range&gt; &lt;state&gt;Alaska&lt;/state&gt; &lt;long&gt;-138.2431&lt;/long&gt; &lt;lat&gt;58.3026&lt;/lat&gt; &lt;elev_ft&gt;14992&lt;/elev_ft&gt; &lt;elev_m&gt;4484&lt;/elev_m&gt; &lt;prominence_ft&gt;14960&lt;/prominence_ft&gt; &lt;prominence_m&gt;4429&lt;/prominence_m&gt; &lt;rank&gt;2&lt;/rank&gt; &lt;/mountain&gt; &lt;/data&gt; &lt;/mountainpeaks&gt;</pre>

LMR\_U3\_L21

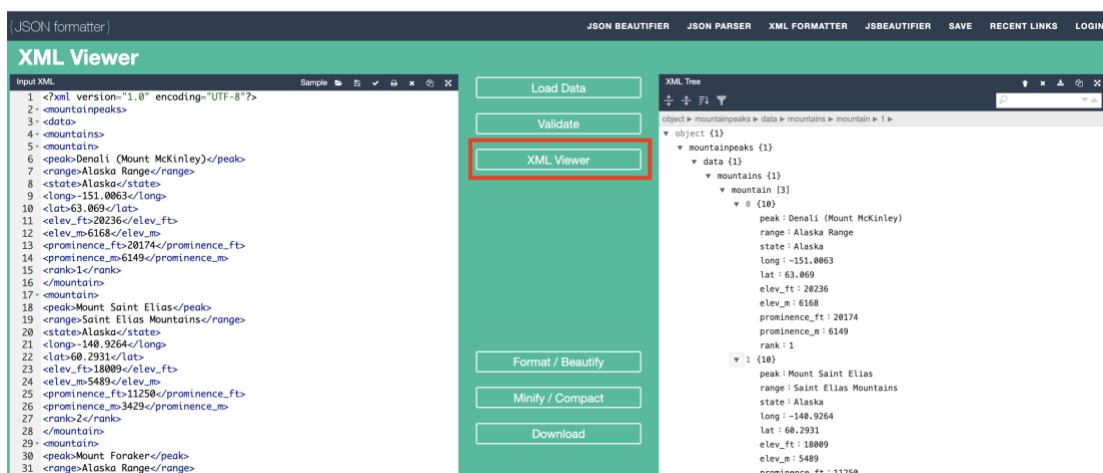
8. Ask student pairs to answer the following:
- Why are certain XML tags indented in the XML version of the data? *Answer: The indentations help us visually see the structure of the data. For example, all the mountains are contained in the <data> section and are further tagged by each particular mountain within the <mountain> and </mountain> tags. All information stored between those two tags tell us the data value connected to the variable name within the tag.*
  - What are the role of tags (ex. <state>) and end tags (ex. </state>) in the XML code? *Answer: Tags tell us when a certain type of data begins, and end tags tell us when the data should end. In other words, it tells us where to find the specific values of a variable (ex. Alaska would be the value of the "state" variable since it is between the <state> and </state> tags).*
  - Where are the variable names? *Answer: The variable names can be found between each <mountain> and </mountain> tags. Specifically, the first variable is "peak" and the last variable is "rank".*
  - Where are the observations? *Answer: The observations are located within each of the variable tags. For example, the observation "Denali (Mount McKinley)" is found between the <peak> and </peak> tags.*
9. Assign student pairs one of the above questions to share out with the class. Student pairs that did not receive an assignment must participate using the Agree/Disagree strategy.
10. As a class, discuss the answers to the questions above.
11. Inform students that they will be comparing how HTML and XML code display by default.
12. Demonstrate how to navigate to the two websites we will be using to display the HTML and XML code by projecting the following URLs:
- <https://jsonformatter.org/xml-viewer>
  - <https://jsonformatter.org/html-viewer>

13. Inform students that they will be copying and pasting their XML code from the electronic copy of LMR\_U3\_L21 into the left “XML Viewer” (Input XML) panel as shown below.

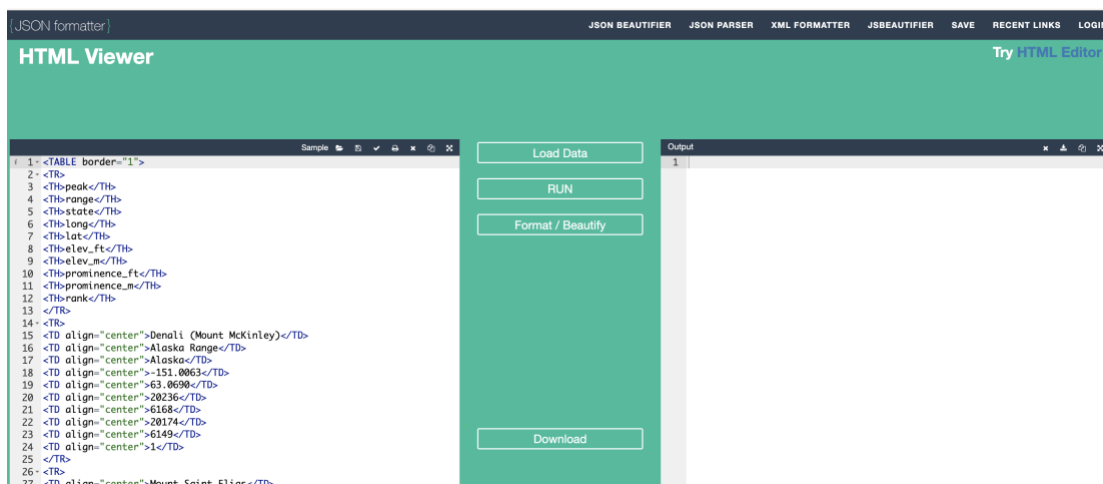


14. In order to view how the XML code displays, click “XML Viewer” in the center of the web page as shown below.

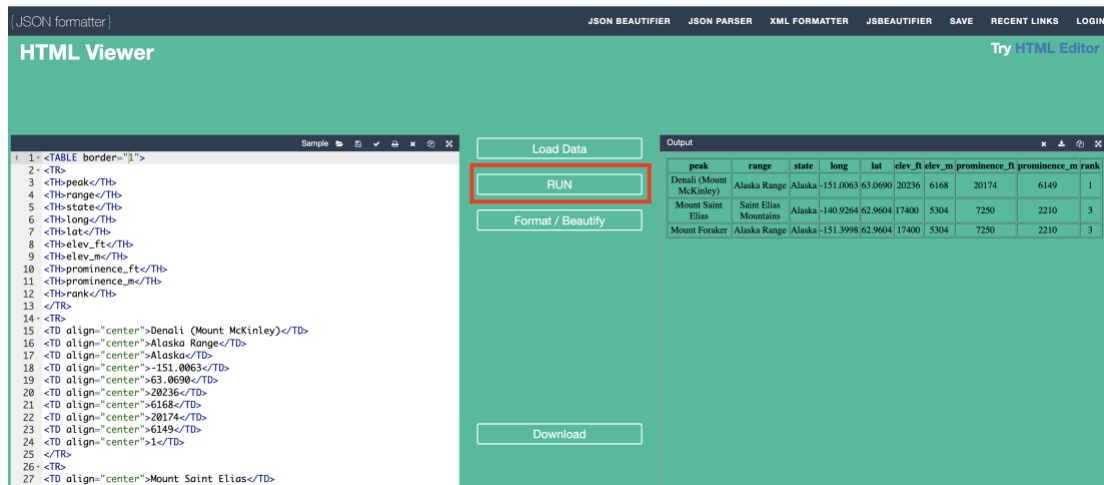
**Note:** Click on the dropdown arrows of the Output XML to explore how it is displayed.



15. We want to compare how the XML code displays with how the HTML code displays. Inform students that they will now be copying and pasting their HTML code into the left “HTML Viewer” (Input HTML) panel as shown below.

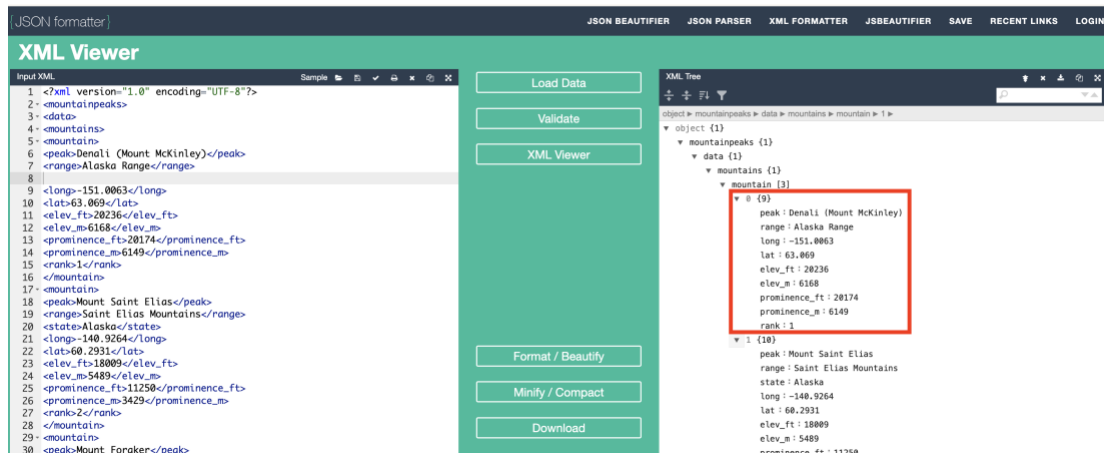


16. In order to view how the HTML code displays, click “RUN” in the center of the web page as shown below.



17. Examine the differences in how HTML and XML are displayed. Ask:
- What do you notice? *Answers will vary but students might notice that HTML is displayed as a table while XML displays the data in a list-like format. XML also has each data value tied to the variable name while HTML has one variable name for the three data values it represents.*
  - How does the Output XML compare to the Output HTML? *Possible answers might be that XML shows the data organized but not in any type of visualization/table while HTML is displayed as a table.*
18. One of the benefits of XML is that removing a data point effectively removes it from the data file but does not necessarily change the integrity of the data – the variable names are still tied to the data values. Demonstrate this by removing a data point from the Input XML and then click “XML Viewer” again to see the updated Output XML.

**Note:** In the image below, the <state>Alaska</state> data point has been removed.



19. Examine the updated Output XML. Ask:
- What changed when we removed a data point? *Answer: The first section of mountain data now has 9 data values instead of 10.*
  - Can we still analyze this data accurately? *Answer: Yes, each data value is still tied to the variable name that represents it.*
  - What do you think will happen when we remove a data point from our HTML code? *Answers will vary.*

20. Demonstrate removing a data point from the Input HTML and then click “RUN” again to see the updated Output HTML.

**Note:** In the image below, the `<TD align="center">Alaska</TD>` data point has been removed.

The screenshot shows the HTML Viewer interface. On the left, the input HTML code is displayed, showing a table with 8 columns: peak, range, state, long, lat, elev\_ft, elev\_m, prominence\_ft, prominence\_m, and rank. The first row is for Denali (Mount McKinley) in Alaska. The second row is for Mount Saint Elias in Alaska. The third row is for Mount Foraker in Alaska. The state 'Alaska' is highlighted in the second row. In the center, there are buttons for 'Load Data', 'RUN', 'Format / Beautify', and 'Download'. On the right, the output HTML is displayed, showing the same table but with the second row (Mount Saint Elias) removed. The first row (Denali) is now the only row in the table.

peak	range	state	long	lat	elev_ft	elev_m	prominence_ft	prominence_m	rank
Denali (Mount McKinley)	Alaska Range	Alaska	-151.0063	63.0690	20236	6168	20174	6149	1

21. Examine the Output HTML. Ask:
- What changed when we removed a data point? *Answer: The table has a blank value at “rank”, and it appears that all the data values in that row have shifted over – HTML skipped the missing data value and filled in the rest of the table with the data values given.*
  - Can we still analyze this data accurately? *Answer: No, starting with the “state” value that is missing, all of the values in the row shifted to the left ending up in the wrong columns.*
22. Summarize: XML formats make it easier to store data on the web in a pleasing manner and make it easier for programmers to find and alter data if the values change or if, for example, they wish to add a new row to a table. Depending on how the web page presents the XML (table, list, etc.), the consequences of changing the data vary on how the program is designed to display it.

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework

For the next 3 days, students will collect data using the class’s newly created Participatory Sensing campaign (see Lessons 17-19).



For homework, students should reflect on the purpose of XML and HTML as it pertains to data.

## Lesson 22: Changing Format

### Objective:

Students will learn how to convert XML files to the more familiar data table format and vice versa.

### Materials:

1. RStudio
2. *There and Back Again: From XML to Data Tables* handout (LMR\_U3\_L22\_A)  
**Note:** An electronic copy should be provided to students (see step 3 in lesson)
3. *There and Back Again: From Data Tables to XML* handout (LMR\_U3\_L22\_B)

**Essential Concepts:** Converting XML to spreadsheet format helps us better understand and view our data.

### Lesson:

1. Take a few minutes to compare the default views of XML code to HTML code (refer to steps 14 and 16 from Lesson 21).
2. Inform students that in today's lesson, they will learn how to translate information from XML code into a data table.
3. Distribute the *There and Back Again: From XML to Data Tables* handout (LMR\_U3\_L22\_A) to students.

**Note:** Provide an electronic copy of LMR\_U3\_L22\_A for students to easily copy/paste the code later in the lesson.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

There and Back Again:  
From XML to Data Tables

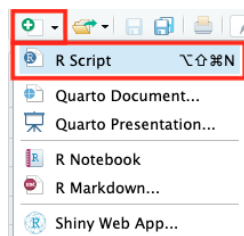
Instructions:  
Translate the XML data into an R data table and answer the questions on page 2.

```
<volunteers_data>
 <volunteer>
 <name>Dakota</name>
 <organization>No Kill LA (NKL)</organization>
 <time>3</time>
 </volunteer>
 <volunteer>
 <name>Huyden</name>
 <organization>Yosemite Foundation</organization>
 <time>3</time>
 </volunteer>
 <volunteer>
 <name>Charlie</name>
 <organization>Yosemite Foundation</organization>
 <time>2</time>
 </volunteer>
 <volunteer>
 <name>Emerson</name>
 <organization>City of Hope</organization>
 <time>1</time>
 </volunteer>
 <volunteer>
 <name>Jessie</name>
 <organization>Rounded Warrior Project</organization>
 <time>2</time>
 </volunteer>
 <volunteer>
 <name>Sawyer</name>
 <organization>City of Hope</organization>
 <time>2</time>
 </volunteer>
</volunteers_data>
```

LMR\_U3\_L22\_A

4. Inform the students that XML code is provided on page 1 of the handout. Their goal is to translate the XML code in RStudio into a readable table. They will then transfer all of the information to the empty data table in the handout.

**Note:** It may be helpful to create a new RScript for the process.

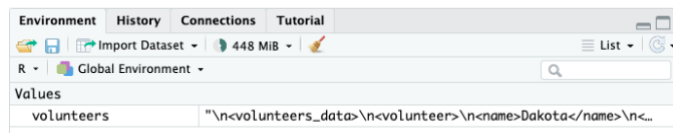


5. Load the XML package by running the following code:  
**library("XML")**

6. Next, assign the data the name **volunteers** and copy/paste the code within quotes. Verify that it is in your Environment.

```
volunteers <- "
<volunteers_data>
 <volunteer>
 <name>Dakota</name>
 <organization>No Kill LA (NKLA)</organization>
 <time>3</time>
 </volunteer>
 <volunteer>
 <name>Hayden</name>
 <organization>Yosemite Foundation</organization>
 <time>3</time>
 </volunteer>
 <volunteer>
 <name>Charlie</name>
 <organization>Yosemite Foundation</organization>
 <time>2</time>
 </volunteer>
 <volunteer>
 <name>Emerson</name>
 <organization>City of Hope</organization>
 <time>1</time>
 </volunteer>
 <volunteer>
 <name>Jessie</name>
 <organization>Wounded Warrior Project</organization>
 <time>2</time>
 </volunteer>
 <volunteer>
 <name>Sawyer</name>
 <organization>City of Hope</organization>
 <time>2</time>
 </volunteer>
 <volunteer>
 <name>Kamryn</name>
 <organization>No Kill LA (NKLA)</organization>
 <time>1</time>
 </volunteer>
 <volunteer>
 <name>London</name>
 <organization>LA Regional Food Bank</organization>
 <time>4</time>
 </volunteer>
</volunteers_data>"
```

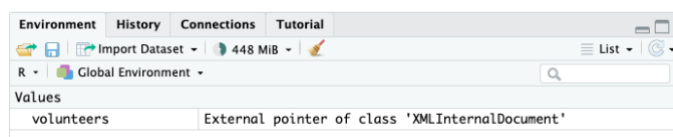
7. Ask students: “How does our data look? Can we analyze it?” *Answers will vary but students should notice that RStudio/Posit Cloud is displaying it in XML format (\n means new line).*



8. We need to translate our XML data into something that RStudio recognizes. The **xmlParse()** function will do this for us. Run the following line of code to translate our data into something RStudio/Posit Cloud will recognize.

```
volunteers <- xmlParse(volunteers)
```

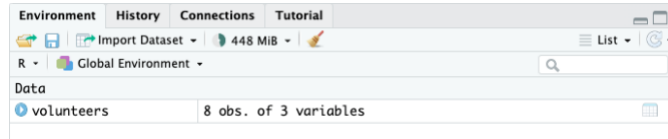
9. Ask students: “How is RStudio recognizing our **volunteers** data now?” *Answers will vary but students should see that RStudio/Posit Cloud recognizes our data as an XML file. Their Environment pane should show the data as an 'XMLInternalDocument'.*



10. We're getting closer – RStudio now recognizes our data as an XML file. Let's do one final translation by using the **xmlToDataFrame()** function to translate our xml file into an R data frame. Run the code below:

```
volunteers <- xmlToDataFrame(volunteers)
```

11. Ask students: “What happened when we translated our xml file into an R data frame?” *Answers will vary but the data should look familiar now – it displays in the Environment pane like datasets we have worked with throughout the year (cdc, atus, slasher, and titanic).*



12. Have students **View** the **volunteers** data and fill in their table in LMR\_U3\_L22\_A.

	name	organization	time
1	Dakota	No Kill LA (NKLA)	3
2	Hayden	Yosemite Foundation	3
3	Charlie	Yosemite Foundation	2
4	Emerson	City of Hope	1
5	Jessie	Wounded Warrior Project	2
6	Sawyer	City of Hope	2
7	Kamryn	No Kill LA (NKLA)	1
8	London	LA Regional Food Bank	4

13. Provide time for students to complete the handout individually.



14. Then, conduct a whole class discussion regarding student responses to the questions on page 2 of the handout.

15. Distribute the *There and Back Again: From Data Tables to XML* (LMR\_U3\_L22\_B) to student teams and allow them time to complete it.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**There and Back Again:  
From Data Tables to XML**

Instructions:  
Translate the data table into an XML data file using appropriate tags and end tags.

Name of Data: **Yosemite Hiking Trails**

trail_name	park_area	miles
North Rim	Yosemite Valley	27.4
South Rim	Yosemite Valley	21.6
Glen Aulin	Tuolumne Meadows	10.6
10 Lakes Basin	Tioga Road	12.4
Clouds Rest	Tuolumne Meadows	14.0

<hikingTrails\_data>  
 <trail>  
 < > </ >  
 < > </ >  
 < > </ >  
 </trail>

LMR\_U3\_L22\_B

16. Once teams have finished, teams will guide you to write the correct XML code.

17. Using a *Whip Around*, teams will tell you the first line of the XML code you need to write. Teams waiting their turn will check if the team is guiding you correctly. If not, they need to stop you and propose their line of code. Do not continue writing the lines of code until all teams agree.

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Homework & Next Day

Students will continue to collect data using the class's Participatory Sensing campaign (see Lessons 17-19). They will analyze the data the next day during the *What Does Our Campaign Data Say?* Practicum.

### **End of Unit 3 Project: Analyzing Our Own Campaign Data**

#### **Objective:**

Students will answer one of the statistical questions they generated at the beginning of the Participatory Sensing campaign creation lesson. They will use RStudio to make graphical representations or numerical summaries of their data to answer their question.

#### **Materials:**

1. *End of Unit 3 Project: Analyzing Our Own Campaign Data* (LMR\_U3\_End\_of\_Unit\_Project)

### **End of Unit 3 Project Analyzing Our Own Campaign Data**

At the start of the Participatory Sensing campaign creation in Lesson 16, the class developed a research question about your class's topic of interest.

It is now time to analyze and interpret your class campaign data. You will use the data from your class-created campaign only. Based on the analysis, you can also wonder about what other data would be necessary to better answer your question, if any.

Based on the class's campaign data collected:

1. Refer back to the statistical questions your team and class generated in Lessons 16-18 that address the research question.
2. Choose one of these statistical questions and determine which variables will answer this question.
3. Analyze the data to answer the question you've chosen. Your analysis should include graphs and numerical summaries. You should:
  - a. Provide the plot and numerical summary.
  - b. Describe what the plot shows.
  - c. Explain why you chose to make that particular plot.
  - d. Explain how the plot and numerical summary answers your statistical question.
  - e. Include the code you used in RStudio to make your plot.
4. After analyzing your data, determine if additional data would better answer your statistical question. If so, propose what that data would be. Different variables? Different data collection approach? Same variables, but more people? Same variables and people but more time?
5. Now, choose two more statistical questions that address the research question.
6. Analyze and interpret the data to answer these questions.
7. Sometimes, when analyzing data, we think of new statistical questions to ask, or we realize that the data need to be cleaned before we can answer. Explain whether this is the case with any of your statistical questions.
8. Write a one-page report and present it to another member of the class who is not in your team.



# Introduction to Data Science

## Unit 4

# Introduction to Data Science

## Daily Overview: Unit 4

Theme	Day	Lessons and Labs	Campaign	Topics	Page
Campaigns and Community (5 days)	1	Lesson 1: Trash		Modeling to answer real world problems, official datasets	307
	2	Lesson 2: Drought		Exploratory data analysis, campaign creation	310
	3	Lesson 3: Community Connection		Community topic research, campaign creation	312
	4	Lesson 4: Evaluate and Implement the Campaign		Evaluate & mock implement campaign	315
	5 <sup>^</sup>	Lesson 5: Refine and Create the Campaign	Team Campaign—data	Revise and edit campaign, data collection	317
Predictions and Models (14 days)	6	Lesson 6: Statistical Predictions in One Variable	Team Campaign—data	One-variable predictions using a rule	320
	7	Lesson 7: Statistical Predictions by Applying the Rule	Team Campaign—data	Predictions applying mean squared error, mean absolute error	322
	8	Lesson 8: Statistical Predictions Using Two Variables	Team Campaign—data	Two-variable statistical predictions, scatterplots	326
	9	Lesson 9: Spaghetti Line	Team Campaign—data	Estimate line of best fit, single linear regression	329
	10	Lab 4A: If the Line Fits...	Team Campaign—data	Estimate line of best fit	331
	11	Lesson 10: What's the Best Line?	Team Campaign—data	Predictions based on linear models	333
	12	Lab 4B: What's the Score?	Team Campaign—data	Comparing predictions to real data	336
	13	Lab 4C: Cross-Validation	Team Campaign—data	Use training and test data for predictions	338
	14	Lesson 11: What's the Trend?	Team Campaign—data	Trend, associations, linear model	341
	15	Lesson 12: How Strong Is It?	Team Campaign—data	Correlation coefficient, strength of trend	345
	16	Lab 4D: Interpreting Correlations	Team Campaign—data	Use correlation coefficient to determine best model	348
	17	Lesson 13: Improving Your Model	Team Campaign—data	Non-linear regression	351
	18	Lab 4E: Some Models Have Curves	Team Campaign—data	Non-linear regression	353
	19	Practicum: Predictions	Team Campaign—data	Linear regression	355
Piecing it Together (3 days)	20	Lesson 14: More Variables to Make Better Predictions	Team Campaign—data	Multiple linear regression	358
	21	Lesson 15: Combination of Variables	Team Campaign—data	Multiple linear regression	361
	22	Lab 4F: This Model Is Big Enough for All of Us	Team Campaign—data	Multiple linear regression	364
Decisions, Decisions! (3 days)	23	Lesson 16: Football or Futbol?	Team Campaign—data	Multiple predictors, classifying into groups, decision trees	366
	24	Lesson 17: Grow Your Own Decision Tree	Team Campaign—data	Decision trees based on training and test data	372
	25	Lab 4G: Growing Trees	Team Campaign—data	Decision trees to classify observations	376
Ties that Bind (3 days)	26	Lesson 18: Where Do I Belong?	Team Campaign—data	Clustering, k-means	380
	27	Lab 4H: Finding Clusters	Team Campaign—data	Clustering, k-means	386
	28+	Lesson 19: Our Class Network	Team Campaign—data	Clustering, networks	388
End of Unit Project (7 days)	29-36	End of Unit 4 Project: Modeling a Community Issue	Team Campaign	Synthesis of above	391

<sup>^</sup>=Data collection window begins.

<sup>+</sup>=Data collection window ends.

## IDS Unit 4: Essential Concepts

### **Lesson 1: Trash**

Exploring different datasets can give us insight about the same processes. Data from our Participatory Sensing campaigns rely on human sensors and limit the ability to generalize to the greater population.

### **Lesson 2: Drought**

Data can be used to make predictions. Official datasets rely on censuses or random samples and can be used to make generalizations.

### **Lesson 3: Community Connection**

Data collected through Participatory Sensing campaigns will be used to create models that answer real-world problems related to our community.

### **Lesson 4: Evaluate and Implement the Campaign**

Statistical questions guide a Participatory Sensing campaign so that we can learn about a community or ourselves. These campaigns should be evaluated before implementing to make sure they are reasonable and ethically sound.

### **Lesson 5: Refine and Create the Campaign**

Statistical questions guide a Participatory Sensing campaign so that we can learn about a community or ourselves. These campaigns should be tried before implementing to make sure they are collecting the data they are meant to collect and refined accordingly.

### **Lesson 6: Statistical Predictions in One Variable**

Anyone can make a prediction. But statisticians measure the success of their predictions. This lesson encourages the classroom to consider different measures of success.

### **Lesson 7: Statistical Predictions by Applying the Rule**

If we use the mean squared errors rule, then the mean of our current data is the best prediction of future values. If we use the mean absolute errors rule, then the median of the current data is the best prediction of future values.

### **Lesson 8: Statistical Predictions Using Two Variables**

When predicting values of a variable  $y$ , and if  $y$  is linearly associated with  $x$ , then we can get improved predictions by using our knowledge about  $x$ . For every value of  $x$ , find the mean of the  $y$  values for that value of  $x$ . If the resulting mean follows a trend, we can model this trend to generalize to unseen values of  $x$ .

### **Lesson 9: Spaghetti Line**

We can often use a straight line to summarize a trend. “Eyeballing” a straight line to a scatterplot is one way to do this.

### **Lesson 10: What’s the Best Line?**

The regression line can be used to make good predictions about values of  $y$  for any given value of  $x$ . This works for exactly the same reason the mean works well for one variable: the predictions will make your score on the mean squared errors as small as possible.

### **Lesson 11: What’s the Trend?**

A positive or negative association between variables provides valuable insights into increasing or decreasing trends, particularly in making predictions. By understanding these associations, we can anticipate future outcomes or behaviors more accurately.

### **Lesson 12: How Strong Is It?**

A high absolute value for correlation means a strong linear trend. A value close to 0 means a weak linear trend.

### **Lesson 13: Improving your Model**

If a linear model is fit to a non-linear trend, it will not do a good job of predicting. For this reason, we need to identify non-linear trends by looking at a scatterplot or the model needs to match the trend.

### **Lesson 14: More Variables to Make Better Predictions**

We can use scatterplots to assess which variables might lead to strong predictive models. Sometimes using several predictors in one model can produce stronger models.

### **Lesson 15: Combination of Variables**

If multiple predictors are associated with the response variable, a better predictive model will be produced, as measured by the mean squared error.

### **Lesson 16: Football or Futbol?**

Some trends are not linear, so the approaches we've done so far won't be helpful. We need to model such trends differently. Decision trees are a non-linear tool for classifying observations into groups when the trend is non-linear.

### **Lesson 17: Grow Your Own Decision Tree**

We can determine the usefulness of decision trees by comparing the number of misclassifications in each.

### **Lesson 18: Where Do I Belong?**

We can identify groups, or "clusters," in data based on a few characteristics. For example, it is easy to classify a group of people into football players and swimmers, but what if you only knew each person's arm span? How well could you classify them into football players and swimmers now?

### **Lesson 19: Our Class Network**

Networks are made when observations are interconnected. In a social setting, we can examine how different people are connected by finding relationships between other people in a network.

# Campaigns and Community

Instructional Days: 5

## Enduring Understandings

Modeling Activities are posed as open-ended problems that are designed to challenge students to build models that are used to solve complex, real-world problems. They may be used to engage students in statistical reasoning and provide a means to better understand students' thinking.

## Engagement

Students will watch a video called *Fighting Pollution Through Data*. This video will set the context of the real-world problem facing many cities around the world — trash. The video provides background information as well as a baseline topic to launch ideas for the modeling process. The video can be found at: <https://www.youtube.com/watch?v=xOYAIXjHveA>

## Learning Objectives

### Statistical/Mathematical:

According to the California Common Core State Standards-Mathematics (CCSS-M) Framework:

“Modeling links classroom mathematics and statistics to everyday life, work, and decision-making. Modeling is the process of choosing and using appropriate mathematics and statistics to analyze empirical situations, to understand them better, and to improve decisions. Quantities and their relationships in physical, economic, public policy, social, and everyday situations can be modeled using mathematical and statistical methods. When making mathematical models, technology is valuable for varying assumptions, exploring consequences, and comparing predictions with data.”

### Focus Standards for Mathematical Practice for All of Unit 4:

SMP-2: Reason abstractly and quantitatively.

SMP-4: Model with mathematics.

SMP-7: Look for and make use of structure.

### Data Science:

Students will apply the conceptual understandings learned up to this point in the curriculum.

### Applied Computational Thinking using RStudio:

- Create a Participatory Sensing campaign using a campaign Authoring Tool.

### Real-World Connections:

Engineers, data scientists, and statisticians, to name a few, use modeling in their everyday work. Whether it is for creating a scale model of a bridge or a mathematical model of force impact measures, modeling is an integral part of what they do in the real world.

## Language Objectives

1. Students will use complex sentences to construct summary statements about their understanding of data, how it is collected, how it used, and how to work with it.
2. Students will engage in partner and whole group discussions and presentations to express their understanding of data science concepts.
3. Students will read informative texts to evaluate claims based on data.

## Data File or Data Collection Method

### Data File:

1. Trash: data(trash)

Data Collection Method:

**Team-generated Participatory Sensing Campaign:** Students will collect data for a team-selected topic.

**Legend for Activity Icons**



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 1: Trash

### Objective:

Students will learn about reducing the burden of trash landfills.

### Materials:

1. Video: *Fighting Pollution Through Data*  
<https://www.youtube.com/watch?v=xOYAIXjHveA>
2. *Landfill Readiness Questions* handout (LMR\_U4\_L1\_A)
3. *Trash Campaign Exploration* handout (LMR\_U4\_L1\_B)
4. *Trash Campaign Creation* handout (LMR\_U4\_L1\_C)

**Essential Concepts:** Exploring different datasets can give us insight about the same processes. Data from our Participatory Sensing campaigns rely on human sensors and limit the ability to generalize to the greater population.

### Lesson:

1. Inform students that they will investigate a problem that faces many cities around the world today: trash.
2. Using the 5 Ws strategy, ask students to write down the 5 Ws in their DS journals as they watch a video about trash. The 5 Ws summarize the What, Who, Why, When, and Where of the resource.



3. Play the *Fighting Plastic Pollution Through Data* video, found at:  
<https://www.youtube.com/watch?v=xOYAIXjHveA>

**Note:** While the video is just over 13 minutes long, students should be able to answer the 5 Ws from the content of the first 10 minutes.



4. After they have finished watching the video, engage in a class discussion around the following question and discuss their insights, questions, and/or reactions to the video:
  - a. What types of data did they collect in the video? *Answer: Photos, location, brand names, weight (kg), waste categories (plastic bags, straws, toothpaste, etc.), and number of packaging.*
  - b. How are they useful in fighting plastic pollution? *Answer: Brand accountability, community involvement, cleaning rivers/dumping sites, and finding sources of pollution.*



5. Now we'll take inventory of our own understanding of landfills and how trash travels there. Distribute the *Landfill Readiness Questions* handout (LMR\_U4\_L1\_A). Allow students private think-time before having them discuss in their teams.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

#### Landfill Readiness Questions

##### Directions:

Answer the questions below and discuss your responses with your team members.

1. What types of items belong in a landfill? Give three examples of things you throw away that belong here.

2. What types of items are recyclable? Give three examples of recyclable items that you use every day.

3. How does an item that you throw away at school get to either the landfill or the recycling center? Why might items sometimes go to the wrong place?

LMR\_U4\_L1\_A

6. Let students know that they will be exploring data from a trash Participatory Sensing campaign, titled the "Trash Campaign," that was conducted at a number of high schools in the Los Angeles Unified School District (LAUSD).

- Concerned students in LAUSD engaged in a model eliciting activity and created a Trash Participatory Sensing campaign to investigate possible trash issues in their communities. Based on the data collected, they made recommendations to the Los Angeles County Sanitation District (LACSD) that would help reduce the use of the regional landfills.
- Distribute the *Trash Campaign Exploration* handout (LMR\_U4\_L1\_B) to assist in students' interaction with the IDS public dashboard.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

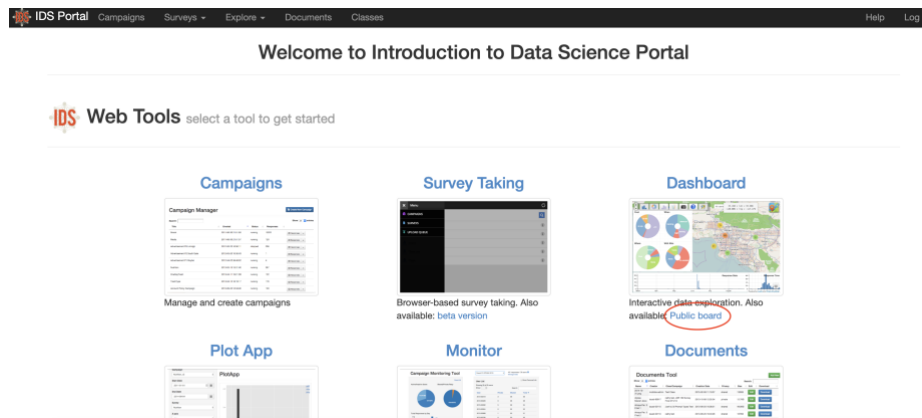
### Trash Campaign Exploration

**Instructions:**  
The survey questions/prompts for the Trash Campaign, a Participatory Sensing Campaign that was conducted at a number of high schools in the Los Angeles Unified School District (LAUSD), are provided below for your reference. Use the IDS public dashboard, <https://portal.idsucia.org/>, to answer the questions on the next page.

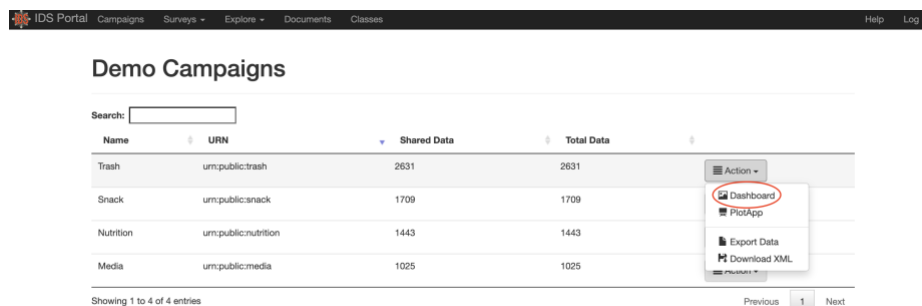
Survey Question/Prompt	Variable Name	Data Type
1. Please take a photo of your trash.	photo	photo
2. Please describe your trash.	trash_name	text
3. What type of trash? <input type="checkbox"/> recyclable <input type="checkbox"/> landfill <input type="checkbox"/> compost	trash_type	category
4. Where was this trash generated/located? <input type="checkbox"/> home <input type="checkbox"/> school <input type="checkbox"/> work <input type="checkbox"/> restaurants <input type="checkbox"/> stores/malls <input type="checkbox"/> in transit <input type="checkbox"/> others	where	category
5. What activity generated this trash? <input type="checkbox"/> eating/looking <input type="checkbox"/> drinking <input type="checkbox"/> school work <input type="checkbox"/> cleaning <input type="checkbox"/> shopping <input type="checkbox"/> I found it <input type="checkbox"/> other	activity	category
6. Where did you put this trash when you were done? <input type="checkbox"/> recyclable <input type="checkbox"/> trash <input type="checkbox"/> compost/green waste <input type="checkbox"/> litter	disposal_bin	category
7. How many recycling bins can you see from your location?	recycle_bins	number

LMR\_U4\_L1\_B

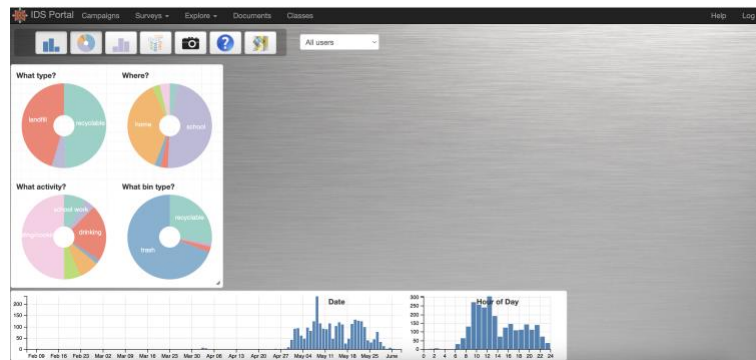
- Navigate students to the IDS public dashboard: <https://portal.thinkdataed.org/>



- They should use the “Trash” campaign data and select “Dashboard” from the “Action” button.



- The dashboard is a visual tool for exploring and analyzing data. An example screenshot of the Trash campaign in the dashboard is shown on the next page.



12. Students should “play” with the data and think about characteristics of the campaign. Answers to the questions in the *Trash Campaign Exploration* handout (LMR\_U4\_L1\_B) are included below:
- How many observations are in this dataset? *Answer: 2,631.*
  - Where was the majority of trash generated? *Answer: School (1,254).*
  - How many observations were generated at school? At work? *Answer: School has 1,254 and Work has 59.*
  - What material or item was most commonly thrown away? *Answer: Recyclable (1,302) or Paper (477).*
  - Between what hours is the largest percentage of trash generated at home? *Answer: Between 2100 and 2200, which is 9pm to 10pm (106).*
  - Which activity generates the largest percentage of landfill-destined trash? *Answer: Eating/cooking with 69.02% (822/1,191)*
  - Is eating or drinking more likely to generate a recyclable piece of trash? *Answer: Drinking, because it resulted in 480 recyclables versus Eating resulted in 399 recyclables, out of 1,302 total.*
  - When recycle bins were present, what percentage of time did a recyclable item end up in a trash bin? *Answer: 26.6% (225/846)*
  - When recycle bins are present, did a higher proportion of recyclable items end up in the trash bin when people were at home or at school? *Answer: School (134/225) had a higher proportion of recyclable items end up in the trash than Home (76/225).*
  - When someone littered, how many times was the person not around any type of waste receptacle? *Answer: 15 out of 58.*
13. Take time at the end of class to share out and discuss the components of the Trash Campaign. If needed, you may use the *Trash Campaign Creation* handout (LMR\_U4\_L1\_C) as an additional resource to help with the deconstruction of the Trash Campaign.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Trash Campaign Creation**

**Instructions:**  
As a class, work together to fill in the information in this handout. You will be deconstructing the information that was used for the Trash Campaign.

**Round 1: Topic**

What do you think was the area of interest that students wanted to know more about with the Trash Campaign?

**Topic:** \_\_\_\_\_

**Round 2: Research Question**

What do you think was the main question that students wanted to answer about the topic?  
This would have been the focus of the Trash Campaign.

**NOTE:** You should NOT be able to simply search the Internet to find the answer to this question; data collection is required.

**Research Question:** \_\_\_\_\_

LMR\_U4\_L1\_C

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 2: Drought

### Objective:

Students will learn about the cause of California droughts and its effect on other states.

### Materials:

1. Article: *California, 'America's garden', is drying out*  
<https://yaleclimateconnections.org/2021/06/california-americas-garden-is-drying-out/>
2. U.S. drought data interactive map  
<https://www.drought.gov/historical-information?dataset=0&selectedDateUSDM=20120710>
3. Team Campaign Creation handout (LMR\_U4\_L2)

**Essential Concepts:** Data can be used to make predictions. Official datasets rely on censuses or random samples and can be used to make generalizations.

### Lesson:

1. Begin the lesson with the quote: "The consequences of drought in California are felt well outside the state's borders. California is effectively America's garden – it produces two-thirds of all fruits and nuts grown in the U.S."
2. Using the K-L-W strategy in their DS journals, give students a couple of minutes to write what they Know about droughts.
3. Then, students will write what they Learned about the California drought as they read the article titled *California, 'America's garden', is drying up*. The article is found at:  
<https://yaleclimateconnections.org/2021/06/california-americas-garden-is-drying-out/>
4. Finally, they will write 2-3 questions about what they Want to know/learn about droughts.
5. Do a quick Whip Around to share some of the students' responses to the K-L-W.
6. Next, load an interactive map of U.S. drought data by visiting:  
<https://www.drought.gov/historical-information?dataset=0&selectedDateUSDM=20120710>

7. Lead a discussion about what is on the page. Ask:
  - a. What do the colors and percentages on the legend mean? *Answer: The color is the drought type (Abnormally Dry, Moderate Drought, etc.) and the percentage is the percentage of area of the U.S. that is that type of drought.*
  - b. What do the colors, percentages, and years on the graph mean? *Answer: The colors correspond to the drought type (as mentioned in the previous question), the percentages correspond to the percentage of area of the U.S. that is that type of drought (as mentioned in the previous question), and the years signify the dates for which the data is available.*
  - c. What information are the map and graph displaying? *Answer: The map tells us the areas of drought and the graph tells us percentages of drought over time.*
  - d. Which date is selected by default and what percent of the U.S. is in Exceptional Drought that day? *Answer: July 10, 2012, and 0.62% is in Exceptional Drought.*

**Note:** If you used a different link other than was listed here, you may have a different default date. In order to see the types of droughts and their percentages, hover your cursor over the graph to line up with the date.

8. Then click on a new date to update the map. Ask:
  - a. What date is displayed now? *Answers will vary.*
  - b. What information is the map displaying? *Answers will depend on the chosen date above.*
  - c. Which states are affected by drought? *Answers will depend on the chosen date above.*

**Note:** You can click on a state to display its name.



9. Then click on the Combine States option, choose a state, and click Combine. Ask:
  - a. What information is the map displaying now? *Answers will vary but students should see the chosen state as the new map.*
  - b. What else do you see? *Answers will depend on the chosen state above.*
  - c. What are some wonderings you have about the data? *Answers will vary.*

**Note:** There are multiple options here. You can choose to display multiple states by adding another state (or states) and choosing Combine. You can also designate a specific time period in the Time Series Options.

10. Now that they know the overall layout of the interactive map and its options, ask student teams to complete the *Team Campaign Creation* handout (LMR\_U4\_L2).

**Note:** The interactive map has a download data feature that allows customization/filtering of U.S. drought data. Keep this in mind as an option for students who may be interested in this topic.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Team Campaign Creation**

**Instructions:**  
As a team, work together to fill in the information in this handout. You will be deciding, as a team, what information will be used for your participatory sensing campaign.

**Article that addresses community concern or focus:**

\_\_\_\_\_

**Round 1: Topic**  
*This is a hobby, area of interest, or place or process that you want to know more about.*

**Team Ideas of Topics:**

\_\_\_\_\_

\_\_\_\_\_

\_\_\_\_\_

**Team Decided Topic:**

\_\_\_\_\_

LMR\_U4\_L2

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Lesson 3: Community Connection

#### Objective:

Students will research and review articles from accredited sources to guide the design of their community focused Participatory Sensing campaign.

#### Materials:

1. Resource: *Web Literacy for Student Fact-Checkers*  
<https://pressbooks.pub/webliteracy/>
2. Video: *Online Verification Skills - Investigate the Source* found at:  
[https://www.youtube.com/watch\\_popup?v=hB6qjlxKItA](https://www.youtube.com/watch_popup?v=hB6qjlxKItA)
3. Video: *Online Verification Skills - Verifying Images and Videos* found at:  
[https://www.youtube.com/watch\\_popup?v=7eKG9RuqUE4](https://www.youtube.com/watch_popup?v=7eKG9RuqUE4)
4. Video: *Online Verification Skills - Find the Original Source* found at:  
[https://www.youtube.com/watch\\_popup?v=tRZ-N3OvvUs](https://www.youtube.com/watch_popup?v=tRZ-N3OvvUs)
5. *Team Campaign Creation* handout (LMR\_U4\_L2)  
**Note:** This handout was used in Lesson 2 but a new copy should be used for this lesson.
6. Article: *Great Pacific Garbage Patch: The World's Biggest Landfill in the Pacific Ocean*  
<https://science.howstuffworks.com/environmental/earth/oceanography/great-pacific-garbage-patch.htm>
7. Data file: *US Landfills*  
<https://www.epa.gov/lmop/landfill-technical-data>
8. Article: *3 reasons why California's drought isn't really over, despite the rain*  
<https://www.npr.org/2023/03/23/1165378214/3-reasons-why-californias-drought-isnt-really-over-despite-all-the-rain/>
9. Data file: *US Drought*  
<https://www.drought.gov/historical-information?dataset=0&selectedDateUSDM=20120710>
10. Article: *The number of homeless people in America grew in 2023 as high cost of living took a toll*  
<https://www.usatoday.com/story/news/nation/2023/12/15/homelessness-in-america-grew-2023/71926354007/>
11. Data file: *Annual Homeless Assessment Report (AHAR)*:  
<https://www.huduser.gov/portal/datasets/ahar.html>
12. Article: *COVID-19 pandemic triggers 25% increase in prevalence of anxiety and depression worldwide*  
<https://www.who.int/news/item/02-03-2022-covid-19-pandemic-triggers-25-increase-in-prevalence-of-anxiety-and-depression-worldwide>
13. Data file: *National Health Interview Survey (NHIS)*:  
<https://www.cdc.gov/nchs/nhis/data-questionnaires-documentation.htm>

**Essential Concepts:** Data collected through Participatory Sensing campaigns will be used to create models that answer real-world problems related to our community.

#### Lesson:

1. In this lesson, students will decide on a topic and research question for their Team Participatory Sensing campaign. They will review articles and choose at least one to set the context for the real-world problem that they will be addressing with their campaign.
2. It is important to check the reliability of your sources so that you are not spreading misinformation or basing your research on untrustworthy or inaccurate information.
3. An option for checking the credibility of sources is the SIFT method (the "SIFT" method is adapted from <https://pressbooks.pub/webliteracy/> by Michael A. Caulfield):
  - a. **Stop** – First, ask yourself if you know and/or trust the source of the page. Don't read it or share it until you know what it is. Second, if you have started working through the moves and find yourself in a "click cycle," STOP and remind yourself what your goal is – don't lead yourself astray down a rabbit hole of facts/links.



- b. **Investigate the Source** – Take time to figure out where it is from before reading it. If it's from an unfamiliar source, do a Google or Wikipedia search. Here's a video to help with investigating a source:

[https://www.youtube.com/watch\\_popup?v=hB6qjlxKltA](https://www.youtube.com/watch_popup?v=hB6qjlxKltA)



- c. **Find Trusted Coverage** – This is where we determine if a claim is reliable. Do a Google News search for relevant stories and use known fact-checking sites like snopes.com, factcheck.org, or politifact.com. If there are images in your source, you can do a reverse image search via Google to check its validity or if it's used elsewhere. Here's a helpful video on how to do a reverse image search:

[https://www.youtube.com/watch\\_popup?v=7eKG9RuqUE4](https://www.youtube.com/watch_popup?v=7eKG9RuqUE4)



- d. **Trace claims, quotes, and media back to the original context** – A lot of what we see on the web is a re-reporting or commentary of a story. When you see phrases like "As reported by..." or "According to..." it can be an indication of this. Here's a video to help navigate finding an original source:

[https://www.youtube.com/watch\\_popup?v=tRZ-N3OvvUs](https://www.youtube.com/watch_popup?v=tRZ-N3OvvUs)

4. Now that you know how to check the credibility of sources, we will review characteristics of a Participatory Sensing campaign.



5. Refer back to the class created campaign from Unit 3 to review. Using a *Pair-Share* strategy, ask students to discuss when a Participatory Sensing campaign should be used rather than a survey. *Answers will vary. Research questions that include variation across time or across locations are good candidates for Participatory Sensing campaigns; therefore, a trigger is necessary in order to record observations at multiple time points and locations. If a question needs to be answered only once, then a survey is a better method.*

6. Remind students that in the last unit, they created one campaign for the entire class. In this unit, each student team will be creating and implementing a campaign on a topic that addresses a community concern or interest.



7. Distribute the *Team Campaign Creation* handout (LMR\_U4\_L2). Use the remainder of the class for students to find and review articles that will be the basis for their team Participatory Sensing campaigns. They will decide on a topic today and continue to create their campaigns in the next class period.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Team Campaign Creation**

**Instructions:**  
As a team, work together to fill in the information in this handout. You will be deciding, as a team, what information will be used for your participatory sensing campaign.

**Article that addresses community concern or focus:**

\_\_\_\_\_

**Round 1: Topic**  
*This is a hobby, area of interest, or place or process that you want to know more about.*

**Team Ideas of Topics:**

\_\_\_\_\_

\_\_\_\_\_

\_\_\_\_\_

**Team Decided Topic:**

\_\_\_\_\_

LMR\_U4\_L2



8. Facilitate the student teams' brainstorm session by circulating around the room to check for understanding. If teams need help with finding an article and choosing a topic, you may recommend one of the following:

- a. Article: *Great Pacific Garbage Patch: The World's Biggest Landfill in the Pacific Ocean*  
<https://science.howstuffworks.com/environmental/earth/oceanography/great-pacific-garbage-patch.htm>

- i. Data file: *US Landfills* <https://www.epa.gov/lmop/landfill-technical-data>

**Note:** This article and supporting data align with the *trash* topic introduced in lesson 1.

- b. Article: *3 reasons why California's drought isn't really over, despite the rain*  
<https://www.npr.org/2023/03/23/1165378214/3-reasons-why-californias-drought-isnt-really-over-despite-all-the-rain/>
    - i. Data file: *US Drought*  
<https://www.drought.gov/historical-information?dataset=0&selectedDateUSDM=20120710>  
**Note:** This article and supporting data align with the *drought* topic introduced in lesson 2.
  - c. Article: *The number of homeless people in America grew in 2023 as high cost of living took a toll*  
<https://www.usatoday.com/story/news/nation/2023/12/15/homelessness-in-america-grew-2023/71926354007/>
    - i. Data file: *Annual Homeless Assessment Report (AHAR)*:  
<https://www.huduser.gov/portal/datasets/ahar.html>  
**Note:** This article and supporting data align with the topic of homelessness.
  - d. Article: *COVID-19 pandemic triggers 25% increase in prevalence of anxiety and depression worldwide*  
<https://www.who.int/news/item/02-03-2022-covid-19-pandemic-triggers-25-increase-in-prevalence-of-anxiety-and-depression-worldwide>
    - i. Data file: *National Health Interview Survey (NHIS)*:  
<https://www.cdc.gov/nchs/nhis/data-questionnaires-documentation.htm>  
**Note:** This article and supporting data align with the topic of mental health.
9. Take time near the end of class to have student teams share out their Round 1 chosen topics.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

#### Homework

Students will come up with at least 2 possible research questions for their Participatory Sensing campaign. They will come to a consensus about their research question with their team in the next class period.

## Lesson 4: Evaluate and Implement the Campaign

### Objective:

Students will complete the design of their community focused Participatory Sensing campaign and implement a mock campaign to evaluate the feasibility of the campaign.

### Materials:

1. *Team Campaign Creation* handout (LMR\_U4\_L2) from previous lesson

**Essential Concepts:** Statistical questions guide a Participatory Sensing campaign so that we can learn about a community or ourselves. These campaigns should be evaluated before implementing to make sure they are reasonable and ethically sound.

### Lesson:

1. Student teams will continue designing their team Participatory Sensing campaign. Allow them some time to review their possible research questions with their team and to decide on a team research question for Round 2 before moving on to Round 3.
-  2. Round 3: Allow student teams a reasonable amount of time to engage in a brainstorm, in which they will discuss what kind of data needs to be collected in order to answer their research question and when is the best time to trigger the data collection/completion of the survey. Before they begin, ask students to keep the following question in mind: Which of these data will give us information that addresses our research question?
3. Once teams have decided on their types of data and trigger, they will create survey questions/prompts to collect this type of data with this trigger.
-  4. Round 4: Now that the teams have decided on a trigger and the type of data needed, they will discuss and create survey questions/prompts to ask when the trigger is set. The questions should consider all of the possible data they might collect at this trigger event.
5. Once teams have created their survey questions/prompts, they will evaluate each survey question. For each question they should consider:
  - a. What type of survey question/prompt will this be (e.g. number, text, photo, time, single choice)?
  - b. How does this question/prompt help address the research question?
  - c. Does the question/prompt need to be reworded? (Is it clear what is being asked for? Do they know how to answer it?) One way to do this is to pair teams and take turns asking each other prompts. The team that is being asked may explain what information they think the question is asking for.
6. If survey questions asking need to be rewritten, students will decide as a team on the changes.
7. Once finalized, they'll record the survey question/prompt that goes along with each data variable on their *Team Campaign Creation* handout (LMR\_U4\_L2), being cognizant of question bias.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Team Campaign Creation**

**Instructions:**  
As a team, work together to fill in the information in this handout. You will be deciding, as a team, what information will be used for your participatory sensing campaign.

**Article that addresses community concern or focus:**

\_\_\_\_\_

**Round 1: Topic**  
*This is a hobby, area of interest, or place or process that you want to know more about.*

**Team Ideas of Topics:**

\_\_\_\_\_

\_\_\_\_\_

\_\_\_\_\_

**Team Decided Topic:**

\_\_\_\_\_

LMR\_U4\_L2



8. **Round 5:** In teams, students will now generate three statistical questions that they might answer with the data they will collect and to guide their campaign. They need to make sure that their statistical questions are interesting and relevant to their chosen topic. Remind students that they will also have data about the date, time, and place of data collection.

9. Confirm that the questions are statistical and that they can be answered with the data the students propose to collect by circulating around the room to check on each team. Each team will decide on no more than 3 statistical questions to guide their campaign.



10. Now that they have all the pieces of the campaign, teams will evaluate whether their campaign is reasonable and ethically sound. Each team will hold a discussion on the following questions:

- a. Are answers to your survey questions likely to *vary* when the trigger occurs? (If not, you'll get bored entering the same data again and again)
- b. Can the team carry out the campaign?
- c. Do triggers occur so rarely that you'll have very little data? Do they occur so often that you'll get frustrated entering too much data?
- d. Ethics: Would sharing these data with strangers or friends be embarrassing or undermine someone's privacy?
- e. Can you change your trigger or survey questions to improve your evaluation?
- f. Will you be able to gather enough relevant data from your survey questions to be able to answer your statistical questions?

11. During their discussion about whether their campaign is reasonable and ethically sound, if teams discover that they need to make changes, they can make adjustments at this time.

12. Students have collaboratively created their Team Participatory Sensing campaign.

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework

Students will collect data by mock implementing their Team Participatory Sensing campaign.

## Lesson 5: Refine and Create the Campaign

### Objective:

Students will revise their community focused Participatory Sensing campaign according to the finding from the mock implementation of their campaign. Student teams will then create their campaigns using the Campaign Authoring tool.

### Materials:

1. Campaign Authoring handout (LMR\_U4\_L5)

**Essential Concepts:** Statistical questions guide a Participatory Sensing campaign so that we can learn about a community or ourselves. These campaigns should be tried before implementing to make sure they are collecting the data they are meant to collect and refined accordingly.

### Lesson:



1. Student teams will come together to discuss their findings regarding the mock implementation of their campaign. Allow them time to share their findings with their team members.
2. Next, ask student teams to discuss the revisions they need to make according to their findings.
3. Now they will use the Campaign Authoring tool to create a campaign like the ones they see on their smart devices or the computer.
4. Distribute the Campaign Authoring handout (LMR\_U4\_L5). Each team will select a member to type the information required to create their campaign. Then, they will follow the instructions on the handout.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

Campaign Authoring Instructions

Follow the instructions below to create your campaign. If you would like a video tutorial to accompany the written instructions, it can be found here: <https://youtu.be/PzwMCH0ghnI>

Go to the **IDS Home Page** found at <https://portal.ids.ucdla.org/> and click on **Campaign Manager**. Then, click on the **Create New Campaign** button on the top right-hand side of the page. Finally, follow the steps below:

- a. **Campaign Info:**
  - i. **Campaign Name:** Give your campaign a name. A name related to the topic is recommended (do not include commas, periods, etc.).
  - ii. **Select your class/period.**
  - iii. **Description:** Provide a one-sentence description of your campaign.
  - iv. **Campaign Status:** Select Running.
  - v. **Data Sharing:** Select Disabled in order to monitor for improper responses.
  - vi. **Editable Responses:** Select Disabled.
  - vii. Click the **+Add Survey** button.
- b. **Survey Window:**
  - i. **Title:** Give the survey a title (again, it may or may not be the same as the campaign name). Users see the title and the all the prompts that follow.
  - ii. **ID:** Give the survey a name (it may be the same as the campaign name). Users do not see the survey ID.
  - iii. **Description:** Provide a short description of the survey for display (optional – may be the same as the Description in Campaign Info).
  - iv. **Submission Message:** Provide a brief message to be displayed after survey submission. Note: It is helpful to include a reminder to click the green button to submit the survey.
  - v. Click the **+Add Prompt** button and select the prompt type for your first survey question.
- c. **Prompt Information:**

LMR\_U4\_L5

**Note:** To name their campaign, a naming convention is suggested. Otherwise, you will have multiple campaigns with the same name. For example, teams may include their team name or number in order to easily identify their campaigns (do not include commas, periods, etc.).

5. Ask teams to *refresh* their campaigns on their smartphones or their web browser to verify that their campaign appears as one of the choices.
6. Now that they are finished with their campaigns, student teams will use the remaining time to plan the expectations for their End of Unit 4 Modeling Activity Project.

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework

Students may begin collecting data by implementing their Team Participatory Sensing campaign. They will have until the end of Unit 4 to collect data. Since only the members of their team will be involved in gathering data for their campaign, this allows time to ensure they have a sufficient amount of data.

# Predictions and Models

Instructional Days: 13

## Enduring Understandings

The regression line is a prediction machine. We give it an x-value, it gives us a predicted y-value. The regression line summarizes the trend in the data, but there may still remain variability in the dependent variable that is not explained by the independent variable. Although the regression line provides optimal predictions when the association is linear, other models are needed for when it is not linear.

## Engagement

Students will explore and make predictions with a dataset consisting of arm span and height values from a group of Los Angeles high school students. The Arm Span vs. Height data allows for a real-world connection while learning about linear models and predictions. They will engage in multiple discussions as they build their understanding of linear models, refine how they make their predictions, and test the accuracy of those predictions.

## Learning Objectives

### Statistical/Mathematical:

- A-SSE A: Interpret the structure of expressions.
- A-REI D-10: Understand that the graph of an equation in two variables is the set of all its solutions plotted in the coordinate plane, often forming a curve (which could be a line).
- F-IF A-2: Use function notation, evaluate functions for inputs in their domains, and interpret statements that use function notation in terms of context.
- F-BF A-1: Write a function that describes a relationship between two quantities.
- S-ID 6: Represent data on two quantitative variables on a scatter plot, and describe how the variables are related.
  - a. Fit a function to the data; use functions fitted to data to solve problems in the context of the data. *Use given functions or choose a function suggested by the context. Emphasize linear models.*
  - b. Informally assess the fit of a function by plotting and analyzing residuals.
  - c. Fit a linear function for a scatter plot that suggests a linear association.
- S-ID 7: Interpret the slope (rate of change) and the intercept (constant term) of a linear model in the context of the data.
- S-ID 8: Compute (using technology) and interpret the correlation coefficient of a linear fit.
- S-IC 6: Evaluate reports based on data. \*

\*This standard is woven throughout the course. It is a recurring standard for every unit.

### Data Science:

Judge whether or not the linear model is appropriate. Learn to interpret a correlation coefficient in a linear model and interpret slope and intercept. Evaluate the strength of a linear association. Evaluate the potential error in a linear model.

### Applied Computational Thinking using RStudio:

- Use linear regression models to predict response values based on sets of predictors.
- Fit a regression line to data and predict outcomes.
- Compute the correlation coefficient of a linear model.

### Real-World Connections:

Many studies are published in which predictions are made, and media reports often cite data that make predictions. They involve one or more explanatory variables and a response variable, such as income vs. education, weight vs. exercise, and cost of insurance vs. age. Understanding linear regression helps evaluate these studies and reports.

### Language Objectives

1. Students will use complex sentences to construct summary statements about their understanding of Mean Squared Error and Mean Absolute Error.
2. Students will engage in partner and whole group discussions to express their understanding of predictions and how to measure their accuracy.
3. Students will use mathematical vocabulary to explain orally and in writing the attributes of various scatterplots.
4. Students will make connections, in writing, between predictions using different types of models (i.e., linear, quadratic, cubic).

### Data File or Data Collection Method

#### Data File:

1. Arm Spans vs. Heights: `data(arm_span)`
2. Movies: `data(movie)`

#### Data Collection Method:

**Team-generated Participatory Sensing Campaign:** Students will collect data for a team-selected topic.

### Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

## Lesson 6: Statistical Predictions in One Variable

### Objective:

Students will devise a rule to determine how to choose a winner when predicting the typical height of all students in a large high school and measure the success of their prediction. They will consider different measures of success.

### Materials:

1. *Heights of Students at a Large High School* handout (LMR\_U4\_L6)

### Vocabulary:

rule

**Essential Concepts:** Anyone can make a prediction. But statisticians measure the success of their predictions. This lesson encourages the classroom to consider different measures of success.

### Lesson:

1. Inform the class that for this lesson, our class will help judge a contest held at a particular high school. This school held a contest in which they selected students at random from a classroom and reported their height.
2. The information in the lesson's steps 3-7 is included in the *Heights of Students at a Large High School* handout (LMR\_U4\_L6).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Heights of Students at a Large High School**

**Background:**  
Our class will help judge a contest held at a particular high school to see who can make the best predictions.  
Height data for 40 randomly selected students were provided to three teams.  
Using this data, each team was asked to predict the heights of a random sample of 10 students. Here is the catch: **teams were allowed to give only ONE number that had to be used to predict all 10 heights.**  
As the judges of this contest, you will determine the winner.

**Instructions:**

1. Your job is to determine the winning team. You must come up with two things:
  - a. You must support your choice of a winner by using a **rule** for calculating a total score for each team.
  - b. The rule must be applied to each team's prediction, and you must be able to explain how your rule helped select the winner. For example, do you choose the team with the largest score? The smallest?
2. Each team's predictions are provided here for your reference.
3. Answer the questions that follow.

**Team Predictions:**

Team A: 67.9 inches	Team B: 68.1 inches	Team C: 70.9 inches
---------------------	---------------------	---------------------

LMR\_U4\_L6

3. Our class will help judge a contest held at a particular high school to see who can make the best predictions. Height data for 40 randomly selected students were provided to three teams. Using this data, each team was asked to predict the heights of a random sample of 10 students. Here is the catch: teams were allowed to give only ONE number that had to be used to predict all 10 heights. As the judges of this contest, you will determine the winner.
4. Your job is to determine the winning team. You must come up with two things:
  - a. You must support your choice of a winner by using a **rule** for calculating a total score for each team.
  - b. The rule must be applied to each team's prediction, and you must be able to explain how your rule helped select the winner. For example, do you choose the team with the largest score? The smallest?
5. Here are the predictions of the three teams:  
Team A: 67.9 inches  
Team B: 68.1 inches  
Team C: 70.9 inches

6. Display Dataset A, found on page 1 of the *Heights of Students at a High School* handout (LMR\_U4\_L6).

**Notes to teacher:**

- a. Students may have to be reminded that negative values with large absolute values are larger than positive values with small absolute values (e.g.,  $|-10|$  is larger than  $|3|$  because 10 is larger than 3).
- b. Let students struggle for a little bit. A prompt to get them started: Look at the difference between a team's prediction and the actual outcomes (e.g., for the first height, Team A predicted 67.9, actual outcome was 70.1, so  $67.9 - 70.1 = -2.2 = |-2.2| = 2.2$ ). They might also need to be nudged towards the *sum* of these differences – they need to produce a single score, not 10 separate scores.
- c. Here are some rules you can “feed” to the class to move them along. Ask them: (a) Describe this rule in words. (b) Is it better to get a high score or a low score or some other score? (c) Which teams win for each? (Note, some of these rules produce ties.)
  - i. Rule 1:  $\text{sum}(\text{heights-predicted.value} == 0)$   
*Translated into words: the number of exactly correct predictions*
  - ii. Rule 2:  $\text{sum}(\text{heights-predicted.value})$   
*Translated into words: the sum of the differences between predicted value and the actual heights*
  - iii. Rule 3:  $\text{sum}(\text{abs}(\text{heights-estimate}))$   
*Translated into words: the sum of the absolute values of the deviations/errors*
  - iv. Rule 4:  $\text{sum}((\text{heights-estimate})^2)$   
*Translated into words: the sum of the squared deviations/errors*

**Note:** It is unlikely that students will think of the last two. That's okay, because we will introduce them in a future lesson, but you might want to present one (or both) to see what they think about these rules.



7. Allow student teams time to discuss and complete the task for Dataset A.
8. Do not share their responses to Dataset A. Instead, display the following questions:
  - a. What if we had a different set of 10 randomly selected students?
  - b. Would the same team win?



9. Allow teams to discuss the questions, then share a couple of responses to the questions in the previous step.
10. Display Dataset B, found on page 2 of the *Heights of Students at a High School* handout (LMR\_U4\_L6), then have them find the winner using this new sample. Is it the same as they chose before?  
**Note:** We do NOT know the value of the true population mean/typical value. This is what we are really trying to predict.
11. Teams will take turns to share their work as follows:
  - a. Which team did you select as the winner using Dataset A?
  - b. Explain the method, or **rule**, your team used to declare the winner.
  - c. Which team did you select as the winner using Dataset B? Is the winner the same?
  - d. Did you use the same rule to select a winner or did it change? If it changed, explain.
12. During the share out, students will take notes about the other teams' rules in their DS journals.
13. Teams may continue to share at the start of the next lesson if they run out of time.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 7: Statistical Predictions by Applying the Rule

### Objective:

Students will apply the rule statisticians use to determine the best method for predicting heights for students at a high school.

### Materials:

1. Each team's rule for determining a winner (from previous lesson)
2. *Heights of Students at a Large High School* handout (LMR\_U4\_L6)
3. *A Tale of Two Rules* handout (LMR\_U4\_L7\_A)
4. *Prediction Games* handout (LMR\_U4\_L7\_B)

### Vocabulary:

training data, test data, mean squared error, mean absolute error, residual

**Essential Concepts:** If we use the mean squared errors rule, then the mean of our current data is the best prediction of future values. If we use the mean absolute errors rule, then the median of the current data is the best prediction of future values.

### Lesson:

1. Ask students to recall that in the previous lesson, each student team created a rule to determine a winner. Which team's rules worked well for determining a winner?
2. Remind them that in their DS Journals, they took notes about each team's rule as they presented. This time, they will be switching roles – instead of creating a rule to judge the given predictions, they will be given a rule and it is their job to find the best procedure to win the contest.
3. Have students refer back to the *Heights of Students at a Large High School* handout (LMR\_U4\_L6) from the previous lesson.
4. Recall that the student teams were provided with height data on 40 selected students to come up with their predictions for future observations. This is a common practice with statisticians and data scientists. The first dataset of 40 students is called the **training data** where we train a model to make predictions. Then we use the **test data** (Dataset A and Dataset B) to test those predictions. Using the training data, the teams used different statistics for their predictions:
  - a. Team A used the mean.
  - b. Team B used the median.
  - c. Team C used the third quartile.
5. In the previous lesson, you created your own rules to determine the winner. Today, you will learn rules that statisticians and data scientists use. The first is called the **mean squared error** rule.  
**Note to teacher:** acknowledge any groups who came up with MSE or MAE on their own in the previous lesson.
6. An “error” is the difference between our prediction and the actual outcome and is sometimes called a “**residual**.” The mean squared error is also called:
  - a. Mean squared deviation
  - b. Mean squared residual
  - c. Residual sum of squares

The formula looks like this, where  $\hat{x}$  stands for the predicted value:

$$MSE = \frac{\sum_{i=1}^n (x_i - \hat{x})^2}{n}$$

Using this formula, the teams' scores are determined by finding the average of the squared differences between their predictions and the actual values. The winner is the team with the lowest mean squared error.

7. Let's use R to do the heavy lifting for us. Demonstrate how to find the mean squared error by typing the following commands in the console (throughout the process, show what is happening in the Environment and the dataframe):

```
#First, let's create a vector of the heights in our first dataset:
heightA <- c(70.1, 61, 70.1, 68.1, 63, 66.1, 61, 70.1, 72.8, 70.9)

#Next, convert this vector into a dataframe:
datasetA <- data.frame(heightA)

#Now we find the residuals using one of the statistics.
#For this example, we'll use the first quartile from the training data (65 inches):
datasetA <- mutate(datasetA, residual = heightA-65)

#Next, we will square each residual:
datasetA <- mutate(datasetA, sq_res = residual^2)

#Finally, we use the mean function to:
#sum up the squared residuals and divide by 10 to find our mean squared deviation:
mean(~sq_res, data = datasetA)
```

This process gives us the mean squared error of 22.05.

**Note to teacher:** The value of the mean squared error will always be in square units. In order to convert back to the original units, simply take the square root of the mean squared error.

Interpretation: When using the  $Q_1$  height to make predictions about all heights, our predictions will typically be off by  $\sqrt{22.05} = 4.6957$  inches.

Here is the vector of heights for Dataset B:

```
heightB <- c(70.1, 72, 68.9, 61.8, 70.9, 59.8, 72, 65, 66.1, 68.9)
```

8. Distribute the *A Tale of Two Rules* handout (LMR\_U4\_L7\_A).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**A Tale of Two Rules**

**Background:**  
How consistent were our rules for the contest in determining a winner?

Statisticians and Data Scientists use rules, like mean squared error, for determining the accuracy of their predictions.

- The **mean squared error** rule says: the score is determined by finding the average of the squared differences between the prediction and the actual values.

**Instructions:**  
For each of the test datasets below, calculate the mean squared error and determine which prediction is best.

**Dataset A**

Here are the heights from Dataset A – 70.1, 61, 70.1, 68.1, 63, 66.1, 61, 70.1, 72.8, 70.9

Summary of 40 Students (heights in inches)					
	Team A	Team B	Team C		
Minimum					
1 <sup>st</sup> Quartile					
Mean					
Median					
3 <sup>rd</sup> Quartile					
Maximum					
MSE					

Which team/statistic made the best prediction using MSE?

**Dataset B**

Here are the heights from Dataset B – 70.1, 72, 68.9, 61.8, 70.9, 59.8, 72, 65, 66.1, 68.9

LMR\_U4\_L7\_A



9. Let's see how well our teams' predictions did on the heights of the test data. Students will work in teams to answer: Using the mean squared errors, which statistic is the winner? Discuss which statistic/team made the best predictions.

*See answers on next page:*

### Dataset A

Here are the heights from Dataset A – 70.1, 61, 70.1, 68.1, 63, 66.1, 61, 70.1, 72.8, 70.9

Summary of 40 Students (heights in inches)					
		Team A	Team B	Team C	
Minimum	1 <sup>st</sup> Quartile	Mean	Median	3 <sup>rd</sup> Quartile	Maximum
59.1	65	67.9	68.1	70.9	76
MSE	84.236	22.05	17.004	17.276	29.484

### Dataset B

Here are the heights from Dataset B – 70.1, 72, 68.9, 61.8, 70.9, 59.8, 72, 65, 66.1, 68.9

Summary of 40 Students (heights in inches)					
		Team A	Team B	Team C	
Minimum	1 <sup>st</sup> Quartile	Mean	Median	3 <sup>rd</sup> Quartile	Maximum
59.1	65	67.9	68.1	70.9	76
MSE	87.673	22.273	16.393	16.573	27.493

**Note to teacher:** Explain that the mean/Team A was the winner of this contest. Data scientists (and mathematicians) can prove that the mean will **always** work best (except in a few weird cases from time to time). So if you want to predict the future, the mean is the best single guess you can make.

- Ask: What if another data science class has a best rule that is different from ours?
- Another agreed upon method that data scientists and statisticians often use is the **mean absolute error**. It's unlikely that students will figure this out on their own. The reasons why we do it in statistics can be proven mathematically but it's beyond the scope of this course. The mean absolute error is expressed as (where  $\hat{x}$  stands for the predicted value):

$$MAE = \frac{\sum_{i=1}^n |x_i - \hat{x}|}{n}$$

- Explain that each team will now use the statisticians' method for declaring a winner. Display the mean absolute error formula and discuss what each symbol means.

Here is the script for MAE:

```
#Find the absolute value of the residuals using one of the statistics
#For this example, we'll use the first quartile from the training data (65 inches):
datasetA <- mutate(datasetA, residual = abs(heightA-65))
#Finally, we use the mean function to:
#sum up the residuals and divide by 10 to find our mean absolute error:
mean(~residual, data = datasetA)
```

This process gives us the mean absolute error of 4.32.

- Using our previous examples, recalculate your predictions using the MAE.
- Using the mean absolute error, which statistic/team is the winner? *See answers below:*

### Dataset A

Here are the heights from Dataset A – 70.1, 61, 70.1, 68.1, 63, 66.1, 61, 70.1, 72.8, 70.9

Summary of 40 Students (heights in inches)					
		Team A	Team B	Team C	
Minimum	1 <sup>st</sup> Quartile	Mean	Median	3 <sup>rd</sup> Quartile	Maximum
59.1	65	67.9	68.1	70.9	76
MAE	8.22	4.32	3.52	3.48	3.96

### Dataset B

Here are the heights from Dataset B – 70.1, 72, 68.9, 61.8, 70.9, 59.8, 72, 65, 66.1, 68.9

Summary of 40 Students (heights in inches)					
		Team A	Team B	Team C	
Minimum	1 <sup>st</sup> Quartile	Mean	Median	3 <sup>rd</sup> Quartile	Maximum
59.1	65	67.9	68.1	70.9	76
MAE	8.45	4.23	3.43	3.39	3.79

**Note to teacher:** Explain that in this instance, the median/Team B is the “winner.” This means that the way you play the game depends on the rules of the game. If we use the mean squared error (MSE), play with the mean. If we use the mean absolute error (MAE), play with the median.

15. Optional practice: Students can practice finding the mean squared error and mean absolute error using the mean and median with the *Prediction Games* handout (LMR\_U4\_L7\_B). The LMR includes the five number summary, if they were curious how the MSE and MAE for other statistics compare.

**Note:** Page 3 of the handout is an answer key for teacher reference.

Name: \_\_\_\_\_ Date: \_\_\_\_\_

**Prediction Games**

**Instructions:**  
For each of the games below, calculate the statistics and determine which one works best.

- The **mean squared error** rule says: the score is determined by finding the average of the squared differences between the prediction and the actual values.
- The **mean absolute error** rule says: the score is determined by finding the average of the absolute value of the differences between the prediction and the actual values.

**Game 1**

Given the following statistics from randomly selected height training data, determine the best predictor for the following test data.

Here are the heights from the test data – 66, 67, 73, 68, 68, 73, 69, 64, 66, 67

Training Data Summary (Heights in inches)					
Minimum	1 <sup>st</sup> Quartile	Median	Mean	3 <sup>rd</sup> Quartile	Maximum
64.20	66.40	67.76	68.22	69.13	73.15

MSE					
MAE					

Which statistic did best with MSE?

Which statistic did best with MAE?

LMR\_U4\_L7\_B

### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 8: Statistical Predictions Using Two Variables

### Objective:

Students will learn how to predict height using arm span data - and vice versa - visually on a scatterplot.

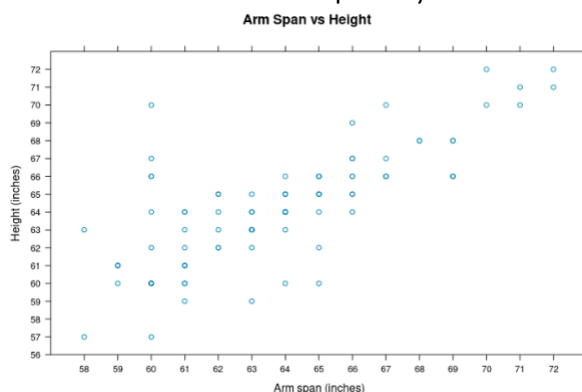
### Materials:

1. *Arm span vs. Height* Scatterplot (LMR\_U4\_L8)  
**Note:** This handout will be referenced in subsequent lessons
2. Assorted color markers (dry erase or overhead) – see step 3 of lesson
3. Overhead or LCD projector

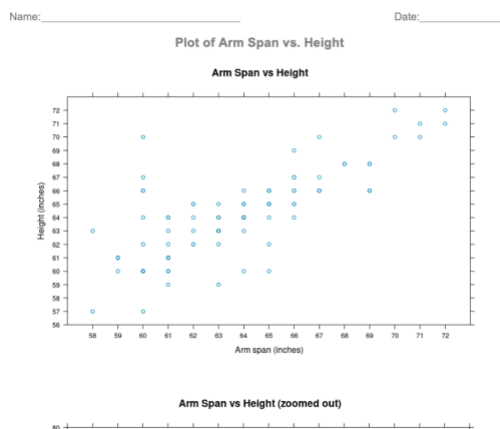
**Essential Concepts:** When predicting values of a variable  $y$ , and if  $y$  is linearly associated with  $x$ , then we can get improved predictions by using our knowledge about  $x$ . For every value of  $x$ , find the mean of the  $y$  values for that value of  $x$ . If the resulting mean follows a trend, we can model this trend to generalize to unseen values of  $x$ .

### Lesson:

1. Remind students that in the previous lessons they were working with height data to predict the typical height of all the students at a large high school, implementing a method used by statisticians to help them make good predictions.
2. In addition to the height data, it turns out that each student's arm span data was also collected and recorded.
3. Display the *Arm Span vs. Height* Scatterplot from LMR\_U4\_L8 on a white board or overhead projector (you will write on the board or the transparency later in the lesson – see step 9 below).



4. Distribute the *Arm Span vs. Height* handout (LMR\_U4\_L8). Students will refer to this handout again later in a subsequent lesson.



5. In teams, ask students to analyze the plot and discuss the following questions:
  - What kind of plot is this? *Answer: Scatterplot.*
  - How many variables are displayed in this plot? *Answer: Two variables.*
  - Which variable is shown on the x-axis? On the y-axis? *Answer: Arm span is shown on the x-axis and height is shown on the y-axis.*
  - What is this plot showing? *Answer: It is showing the relationship between a person's height and the person's corresponding arm span measurement.*
  - How can I find out the height of the person whose arm span measures 68 inches? *Answer: Find 68 on the x-axis. Then find the data point located at 68. Place finger on the data point and track its location on the y-axis. The height is also 68 inches.*



6. Using *Talk Moves*, conduct a class discussion of the questions in step 5 above.
7. Remind students that we've learned that the mean is the best way of predicting heights. The mean height of these people is 64 inches.
8. Ask students: Do you think we can do better? Is 64" a good prediction for someone whose arm span is 72"? What about 60"? How can you come up with a rule for determining the best predicted height if you know the person's arm span?  
**Note to teacher:** Lead students to realize that they can do this by "subsetting" the data for the fixed x value. For example, if arm span is 60", they should consider only the heights of people whose arm span is 60" and find the mean.
9. In teams, ask students to approximate the mean height for people whose arm span is 60", 64", 68", and 72".  
**Note:** Because the plot does not clearly show duplicate ordered pairs, an approximation is sufficient at this point. You may have students use RStudio to calculate the mean height for the specific arm spans. Refer to the OPTIONAL section at the end of this lesson.
10. Then plot these points on the graph. We will use this later – the points should be roughly along a straight line.  
**Note:** These arm spans have a range of height values associated with them. Students may take a mean of the heights, but answers may vary.
11. Ask students if they see any patterns or rules they can use from this to help with predictions. Because there were multiple height values associated with each arm span length, you will likely get multiple answers from students. The goal now is to come up with a rule that suggests a plausible height value for anyone with a particular arm span.
12. A sentence starter to guide students: If a person has a bigger arm span, then we should predict   [a bigger height]  . If time permits, you might push them to be more precise. Let's take someone who has a 60-inch arm span. You predicted a height of       . How much should we increase our prediction for people with a 62-inch arm span? Can you do this without subsetting the data and re-calculating?
13. Conceptually, students are wrestling with the notion of the slope of the regression line but there's no need to point this out just yet.

Important: The equation of the line of best fit will be revealed in Lesson 9.

**OPTIONAL FOR ITEM 9:** If you want to obtain the exact mean height for each arm span value in step 9 above, copy the code below and run it in an R Script.

```
xyplot(height~armspan, data = arm_span,
 scales = list(x = list(at = seq(58, 72, 1)),
 y = list(at = seq(52, 72, 1))),
 xlab = "Arm span (inches)", ylab = "Height (inches)")

armspan_60 <- filter(arm_span, armspan==60)
mean(~height, data = armspan_60)
#62.66667
```

```

armspan_64 <- filter(arm_span, armspan==64)
mean(~height, data = armspan_64)
#64

armspan_68 <- filter(arm_span, armspan==68)
mean(~height, data = armspan_68)
#68

armspan_72 <- filter(arm_span, armspan==72)
mean(~height, data = armspan_72)
#71.5

#Base R Code
#syntax to create a scatterplot using base R
plot(arm_span$height, arm_span$armspan)

#Points function in base R is more user friendly
points(60, 62.66667, col = "red", cex = 2)
points(64, 64, col = "red", cex = 2)
points(68, 68, col = "red", cex = 2)
points(72, 71.5, col = "red", cex = 2)

```

#### Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

## Lesson 9: The Spaghetti Line

### Objective:

Students will estimate the line of best fit for a height and arm span dataset using a strand of spaghetti as a modeling tool.

### Materials:

1. *The Spaghetti Line* (LMR\_U4\_L9)  
**Advance preparation required:** Cut out plots prior to beginning the lesson (see step 2 in lesson)
2. 1 lb. of Uncooked Spaghetti
3. Grid Paper
4. Tape or Glue
5. Poster paper

### Vocabulary:

line of best fit

**Essential Concepts:** We can often use a straight line to summarize a trend. "Eyeballing" a straight line to a scatterplot is one way to do this.

### Lesson:

**Note:** Lab 4A may be done in the same class period as this lesson.

1. Inform students that in this lesson, they will estimate the equation of the **line of best fit** for a height and arm span dataset.
2. Distribute *The Spaghetti Line* handout (LMR\_U4\_L9) to each student and a couple of spaghetti strands per team. Students will estimate the line of best fit as outlined in the handout. Team solutions should be recorded on poster paper. They will glue their assigned plot on the poster and record their responses to the questions on the poster paper.

**Advance preparation required:** Cut out the plots on pages 2-4, one for each group. If there are more than 6 groups, plots can be duplicated for remaining groups.

**Note to teacher:** If necessary, review how to find the slope of a line using two points and how to write an equation using the slope and y-intercept. Students may need extra help with finding the correct y-intercept because they cannot rely on the y-axis values in this case since the plots are not starting at the origin, (0, 0).

Name: \_\_\_\_\_ Date: \_\_\_\_\_

The Spaghetti Line:  
Estimating the Line of Best Fit

Background:

Arm span and height data of students at a large high school were collected.  
Your team will be assigned a plot of a subset of these data. Using the plot, investigate the statistical question:

Is there a relationship between a person's arm span and height?

Instructions:

1. Once your team has been assigned a plot, tape or glue it to poster paper.
2. Using a strand of spaghetti, position the spaghetti to simulate a line that best fits all the data points.
3. Tape or glue the spaghetti line to the plot.
4. Use the grid lines to find two points that go through the line. Identify the points using their coordinates.
5. Find the slope of the line.
6. Find the point in your line where the x-value equals zero. What is the y-value? This is your y-intercept.
7. Write the equation of your spaghetti line on your plot.
8. Use your equation to make a prediction.
9. Answer the statistical question based on your plot.

LMR\_U4\_L9

3. Ask teams to post their work around the room. Conduct a *Gallery Walk* so that teams can see each other's work.



4. Lead a discussion about the teams' lines. Ask: Which team has the best line? Why?

**Note to teacher:** Push the students a bit by adding an obviously bad line to the graph and asking why their line is better than this one. Push them to come to an understanding that the “best” line comes close to the *most* points and this line is called the **line of best fit**.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

### Homework & Next Day

## **Lab 4A: If the line fits...**

Complete Lab 4A prior to Lesson 10.

## Lab 4A: If the line fits...

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### How to make predictions

- Anyone can make predictions.
  - Data scientists use data to inform their predictions by using the information learned from the sample to make predictions for the whole population.
- In this lab, we'll learn how to make predictions by finding the *line of best fit*.
  - You will also learn how to use the information from one variable to make predictions about another variable.

### Predicting heights

- **(1) Write and run code using the `data()` function to load the `arm_span` data.**
- This data comes from a sample of 90 people in the Los Angeles area.
  - The measurements of `height` and `armspan` are in inches.
  - A person's `armspan` is the maximum distance between their fingertips when they spread their arms out wide.
- **(2) Write and run code making a plot of the `height` variable.**
  - **(3) If you had to predict the height of someone in the Los Angeles area, what single height would you choose and why?**
  - **(4) Would you describe this as a *good* guess? What might you try to improve your predictions?**

### Predicting heights knowing arm spans

- **(5) Write and run code creating two subsets of our `arm_span` data:**
  - One for `armspan >= 61` and `armspan <= 63`.
  - A second for `armspan >= 64` and `armspan <= 66`.
- **(6) Write and run code creating a histogram for the `height` of people in each subset.**
- **Answer the following based on the data:**
  - **(7) What `height` would you predict if you knew a person had an `armspan` around 62 inches?**
  - **(8) What `height` would you predict if you knew a person had an `armspan` around 65 inches?**
  - **(9) Does knowing someone's `armspan` help you predict their `height`? Why or why not?**

### Fitting lines

- Notice that there is a trend that people with a larger `armspan` also tend to have a larger mean `height`.
  - One way of describing this sort of trend is with a line.
- Data scientists often *fit* lines to their data to make predictions.
  - What we mean by *fit* is to come up with a line that's close to as many of the data points as possible.

- **(10) Write and run code creating a scatterplot for height and armspan. Then run the following code.**

```
add_line()
```

- On the Plot pane, click two data points to draw a line through.
- NOTE: Watch the following video if you are experiencing difficulties obtaining your line:  
<https://youtu.be/pGqXHGhhwJ8>
- If you are unsuccessful using the `add_line()` function, refer to the next slide to learn how to use the `get_line()` function.

### get\_line()

- The `get_line()` function does not rely on clicking on the scatterplot to choose points, but rather on you providing the points manually.
- For example, let's say you want to obtain the equation of the line that passes through the points (59,60) and (68,67). This is how you would use the `get_line()` function:

```
get_line(c(59,60), c(68,67))
intercept slope
14.1111111 0.7777778
```

- Notice the output is the y-intercept and slope of your line.
- Now you can use the `add_line()` function to include the line in your scatterplot.

```
add_line(intercept = 14.111111, slope = 0.7777778)
```

- If your line doesn't quite fit the way you want it, try another ordered pair or make modifications to the existing equation.

### Predicting with lines

- **(11) Draw a line that you think is a good fit and write down its equation.**
- **(12) Using your equation: Predict how tall a person with a 62-inch armspan and a person with a 65-inch armspan would be.**
- Using a line to make predictions also lets us make predictions for armspans that aren't in our data.
  - **(13) How tall would you predict a person with a 63.5-inch armspan to be?**
  - **(14) Compare your answers with a neighbor. Did both of you come up with the same equation for a line? If not, can you tell which line fits the data best?**

## Lesson 10: What's the Best Line?

### Objective:

Students will understand that the mean squared error (MSE) is a way to assess the fit of a linear model. The MSE measures the total squared distances between all the data values from the line of best fit and divides it by the number of observations in the dataset.

### Materials:

1. *Arm Span vs. Height* Scatterplot (LMR\_U4\_L8) from lesson 8
2. *Testing Line of Best Fit* handout (LMR\_U4\_L10)

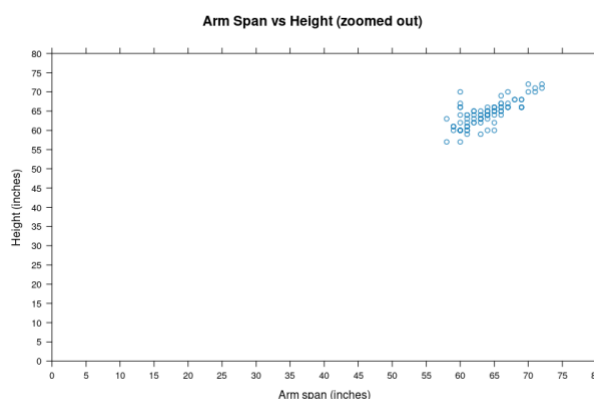
### Vocabulary:

regression line, observed value, predicted value

**Essential Concepts:** The regression line can be used to make good predictions about values of  $y$  for any given value of  $x$ . This works for exactly the same reason the mean works well for one variable: the predictions will make your score on the mean squared errors as small as possible.

### Lesson:

1. Ask student teams to refer back to the *Arm Span vs. Height* handout (LMR\_U4\_L8) but this time have them look at the zoomed out scatterplot.



2. Using their understanding of a line of best fit from The Spaghetti Line lesson and Lab 4A, have them draw (or use strands of spaghetti) what they believe to be the line of best fit for the data.  
**Note:** They can use their equations from Lab 4A as a guide but note that it will be difficult to plot decimals on this scatterplot.
3. Ask students: How does this line compare to the lines from the team posters in The Spaghetti Line lesson? *Answers will vary but students may notice that the y-intercepts may be similar or that the overall slope appears similar (they are not writing an equation for the line in step 2 above, but they may notice where their line intercepts the y-axis and/or the steepness of the line, i.e. slope).*
4. Reveal the equation of the line of best fit for the Arm Span vs. Height data and ask students to compare it to their equations from Lab 4A:

$$\widehat{\text{height}} = 0.7328(\text{armspan}) + 17.4957$$

**Note:** Any time a *hat* is on top of a variable, this means we are making “predicted values” of that variable.

5. Whose equation came closest to the equation of the regression line? Ask the student whose equation came closest to share how he/she came up with the equation.
6. Inform students that the equation of the line is a rule that predicts the height based on a second variable, in this case, arm span. Data points are **observed values** and points on the line are **predicted values**.



7. Team discussion question:

**Using the equation of the line of best fit provided, how can we predict the height of a student whose arm span is 67 inches?**

- What was the actual height for someone with an arm span of 67 inches? *Answer: There are three points on our Arm Span vs Height scatterplot at an arm span of 67 inches; 66-inch height, 67-inch height, and 70-inch height.*
- How close was our predicted height? *Answer: Our line of best fit predicted that someone with an arm span of 67 inches has a height of 66.5933 inches, which rounded to the nearest whole number, is 67 inches. This is pretty close as it matches one of the possible heights for someone with an arm span of 67 inches in our data.*

8. Remind students that lines of best fit are also known as **regression lines** and they are models that can be used to make predictions.

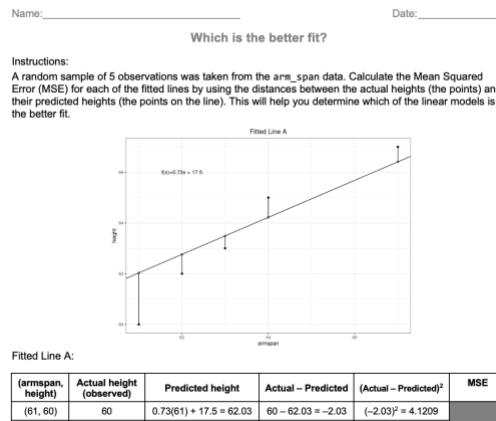


9. Inform students that data scientists have a way of finding the best line. They choose the line so that the mean squared distances between the points and the line is as small as possible. Discuss with students:

- What methods have we used so far? *Answer: We've used Mean Squared Error and Mean Absolute Error (Lesson 7).*
- When is it appropriate to use each method? *Answer: It is best to use Mean Squared Error when we were looking at mean and Mean Absolute Error when we were looking at median.*

10. Distribute the *Testing Line of Best Fit* handout (LMR\_U4\_L10). Students will calculate MSE by using the distances between the actual heights (the points) and their predicted heights (the points on the line) of two different lines. They do this so that they can understand what those distances mean – that together they form our “error” that help us determine the best fitting line.

**Note:** Page 3 of the handout is an answer key for teacher reference.



LMR\_U4\_L10

11. Discuss with students:

- What did you have to do to your MSE value to make it usable for interpretation? *Answer: We had to take the square root of our MSE value in order to convert it back to inches.*
- Which linear model was the better fit? How do you know? *Answers will vary but this is where students should compare the MSE values – a smaller MSE indicates a smaller error, and therefore a better fit.*

**Note:** Students may ask for an easier and/or faster way to calculate MSE. They will be using RStudio to calculate MSE in the next lab.

**Class Scribes:**



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

**Lab 4B: What's the score?**

**Lab 4C: Cross-Validation**

Complete Labs 4B and 4C prior to Lesson 11.

## Lab 4B: What's the score?

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

### Previously

- In the previous lab, we learned we could make predictions about one variable by utilizing the information of another.
- In this lab, we will learn how to measure the accuracy of our predictions.
  - This in turn will let us evaluate how well a model performs at making predictions.
  - We'll also use this information later to compare different models to find which model makes the best predictions.

### Predictions using a line

- Load the `arm_span` data again.
- (1) Write and run code creating an `xyplot` with `height` on the y-axis and `armspan` on the x-axis.
  - Use `add_line()` or `get_line` to graph a line that you think fits the data well.
- (2) Fill in the blanks below to create a function that will make predictions of people's heights based on their `armspan`:

```
predict_height <- function(armspan) {
 ____ * armspan + ____
}
```

### Make your predictions

- (3) Fill in the blanks to include your predictions in the `arm_span` data.  
`____ <- mutate(____, predicted_height = ____ (____))`
- Now that we've made our predictions, we'll need to figure out a way to decide how accurate our predictions are.
  - We'll want to compare our *predicted heights* to the *actual heights*.
  - At the end, we'll want to come up with a single number summary that describes our model's accuracy.

### Sums of differences

- A *residual* is the difference between the actual and predicted value of a quantity of interest.
- (4) Fill in the blanks below to add a column of residuals to `arm_span`.  
`____ <- mutate(____, residual = ____ - ____)`
- (5) What do residuals measure?
- One method we might consider to measure our model's accuracy is to sum the residuals.
- (6) Fill in the blanks below to calculate our accuracy summary.  
`summarize(____, sum(____))`
- Hint: Like `mutate`, the first argument of `summarize` is a dataframe, and the second argument is the action to perform on a column of the dataframe. Whereas the output of `mutate` is a column, the output of `summarize` is (usually) a single number summary.
- (7) Describe and interpret, in words, what the output of your accuracy summary means.
- (8) Write down why adding positive and negative errors together is problematic for assessing prediction accuracy.

## Mean squared error

- When adding residuals, the positive errors in our predictions (underestimates) are cancelled out by negative errors (overestimates) which lead to the impression that our model is making better predictions than it actually is.
- To solve this problem we calculate the squared values of the errors because squared values are always positive.
- The *mean squared error* (MSE) is calculated by squaring all of the residuals, and then taking the mean of the squared residuals.
- **(9) Fill in the blanks below to calculate the MSE of your line.**  

```
summarize(____, mean((____)^2)
```
- **(10) Compare your MSE with a neighbor. Whose line was more accurate and why?**

## Regression lines

- If you were to go around your class, each student would have created a different line that they feel *fit* the data best.
  - Which is a problem because everyone's line will make slightly different predictions.
- To avoid this variation in predictions, data scientists use *regression lines*.
  - We also refer to *regression lines* as *linear models*.
  - This line connects the mean *height* of people with similar *armspans*.
  - **(11) Fill in the blanks below to create a regression line using `lm`, which stands for linear model:**

```
best_fit <- lm(____ ~ ____, data = arm_span)
```

## Plotting regression lines

- **Type `best_fit` into the console to see the slope and intercept of the regression line.**
- **(12) Run the code to create the scatterplot of *armspan* vs. *height* again. Then fill in the blanks below to add the line of best fit.**

```
add_line(intercept = ____, slope = ____)
```

## Predicting with regression lines

- Making predictions with models `R` is familiar with is simpler than with lines, or models, we come up with ourselves.
    - **(13) Fill in the blanks to make predictions using `best_fit`:**
- ```
____ <- mutate(____, predicted_height = predict(____))
```
- Hint: the `predict` function takes a linear model as input, and outputs the predictions of that model.

The magic of `lm()`

- The `lm()` function creates the *line of best fit* equation by finding the line that minimizes the *mean squared error*. Meaning, it's the *best fitting line possible*.
- **(14) Calculate the MSE for the values predicted using the regression line.**
- **(15) Compare the MSE of the linear model you fitted to the MSE of the linear model obtained with `lm()`. Which linear model performed better?**
- **(16) Ask your neighbors if any of their lines beat the `lm()` line in terms of the MSE. Were any of them successful?**

Lab 4C: Cross-Validation

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

What is cross-validation?

- In the previous two labs, we learned how to:
 - Create a linear model predicting **height** from the **arm_span** data (4A).
 - See how well our model predicts **height** on the **arm_span** data by computing mean squared error (MSE) (4B).
- In this lab, we will see how our model predicts heights of *people we haven't yet measured*.
- To do this, we will use a method called *cross-validation*.
- Cross-validation consists of three steps:
 - Step 1: Split the data into **training** and **test** sets.
 - Step 2: Create a model using the **training** set.
 - Step 3: Use this model to make predictions on the **test** set.

Step 1: training-test split

- Waiting for new observations can take a long time. The U.S. takes a census of its population once every 10 years, for example.
- Instead of waiting for new observations, data scientists will take their current data and divide it into two distinct sets.
- **Split the **arm_span** data into training and test sets using the following two steps.**
- **(1) First, fill in the blanks below to randomly select which rows of **arm_span** will go into the training set.**

```
set.seed(123)
training_rows <- sample(1:____, size = 68)
```

- **(2) Second, use the **slice** function to create two dataframes: one called **training** consisting of the **training_rows**, and another called **test** consisting of the remaining rows of **arm_span**.**

```
training <- slice(arm_span, ____)  
test <- slice(____, - ____)
```

- **(3) Explain these lines of code and describe the **training** and **test** datasets.**

Aside: `set.seed()`

- When we split data, we're randomly separating our observations into *training* and *test* sets.
 - It's important to notice that no single observation will be placed in both sets.
- Because we're splitting the datasets randomly, our models can also vary slightly, person-to-person.
 - This is why it's important to use `set.seed`.
- By using `set.seed`, we're able to reproduce the random splitting so that each person's model outputs the same results.

Whenever you split data into training and test, always use `set.seed` first.

Aside: training-test ratio

- When splitting data into `training` and `test` sets, we need to have enough observations in our data so that we can build a good model.
 - This is why we kept 68 observations in our `training` data.
- As datasets grow larger, we can use a larger proportion of the data to `test` with.

Step 2: training the model

- Step 2 is to create a linear model relating `height` and `armspan` using the `training` data.
- **(4) Write and run code fitting a line of best fit model to our `training` data and assign it the name `best_training`.**
- Recall that the slope and intercept of our linear model are chosen to minimize MSE.
- Since the MSE being minimized is from the training data, we can call it *training MSE*.

Step 3: test the model

- Step 3 is to use the model we built on the `training` data to make predictions on the `test` data.
- Note that we are NOT recomputing the slope and intercept to fit the test data best. We use the same slope and intercept that were computed in step 2.
- Because we're using the *line of best fit*, we can use the `predict()` function we introduced in the last lab to make predictions.
 - **(5) Fill in the blanks below to add predicted heights to our `test` data:**
 - Hint: the `predict` function without the argument `newdata` will output predictions on the `training` data. To output predictions on the `test` data, supply the `test` data to the `newdata` argument

```
test <- mutate(test, _____ = predict(best_training, newdata = _____))
```

- **(6) Calculate the test MSE in the same way as you did in the previous lab (test MSE is simply MSE of the predictions on the test data).**

Recap

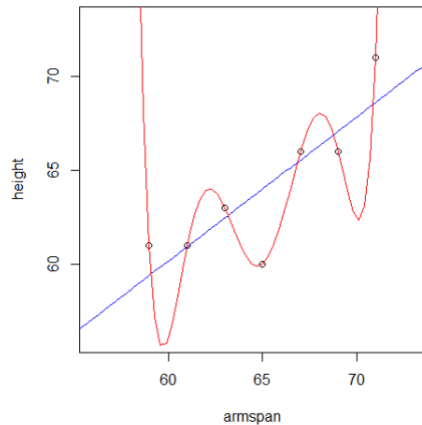
- Another way to describe the three steps is
 - Step 1: Split the data into `training` and `test` sets.
 - Step 2: Choose a slope and intercept that minimize training MSE.
 - Step 3: Using the same slope and intercept from step 2, make predictions on the `test` set, and use those predictions to compute test MSE.
- This begs the question, why do we care about test MSE?

Why cross-validate?

- Why go to all this trouble to compute test MSE when we could just compute MSE on the original dataset?
- When we compute MSE on the original dataset, we are measuring the ability of a model to make predictions on *the current batch of data*.
- Relying on a single dataset can lead to models that are so specific to the current batch of data that they're unable to make good predictions for future observations.
 - This phenomenon is known as *overfitting*.
- By splitting the data into a training and test set, we are *hiding a proportion of the data* from the model. This emulates future observations, which are unseen.
- Test MSE estimates the ability of a model to make predictions on *future observations*.

Example of overfitting

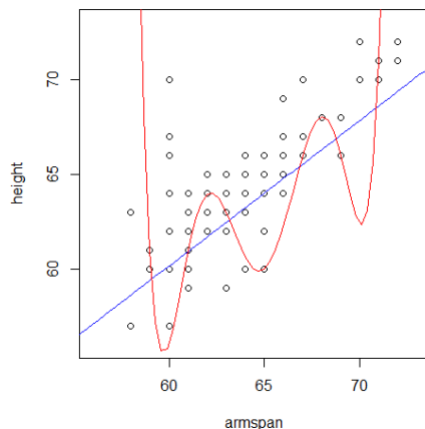
- The following example motivates cross-validation by illustrating the dangers of overfitting.
- We randomly select 7 points from the `arm_span` dataset and fit two models: a linear model, and a *polynomial model*.
 - You will learn how to fit a polynomial model in the next lab.
- Below is a plot of these 7 **training** points, and two curves representing the value of height each model would predict given a value of armspan.



- **(7) Which model does a better job of predicting the 7 training points?**
- **(8) Which model do you think will do a better job of predicting the rest of the data?**

Example of overfitting, continued

- Below is a plot of the rest of the `arm_span` dataset, along with the predictions each model would make.



- **(9) Which model does a better job of generalizing to the rest of the `arm_span` dataset?**

Lesson 11: What's the Trend?

Objective:

Students will understand that the regression line is a model for a linear association (trend). They will learn to identify the direction of trends and interpret the slope and the intercept of a linear model in the context of the data.

Materials:

1. *What's the Trend?* handout (LMR_U4_L11_A)
2. *Predicting Values* handout (LMR_U4_L11_B)

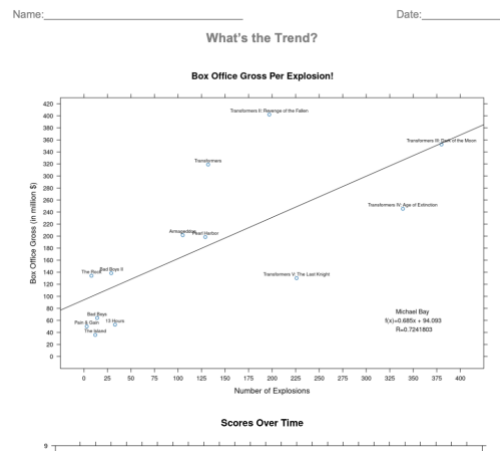
Vocabulary:

trend, positive association, negative association, no association, linear, model

Essential Concepts: A positive or negative association between variables provides valuable insights into increasing or decreasing trends, particularly in making predictions. By understanding these associations, we can anticipate future outcomes or behaviors more accurately.

Lesson:

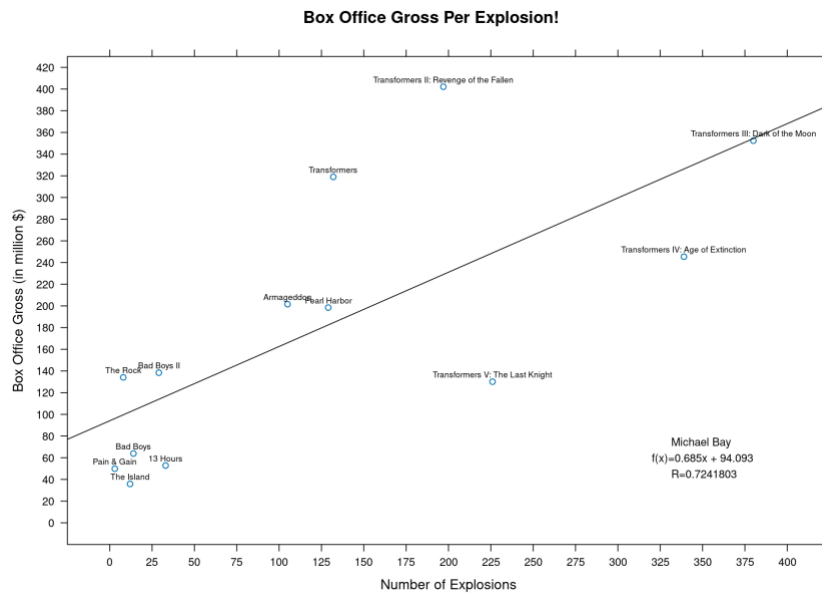
1. Distribute the *What's the Trend?* handout (LMR_U4_L11_A). Students will analyze the two scatterplots on the handout. The *Profits per Explosion* plot shows the relationship between the number of explosions in Michael Bay's movies and the box office gross earned by each movie. The *Scores Over Time* plot shows the relationship between M. Night Shyamalan movies made since *The Sixth Sense* was released in 1999 and their Internet Movie Database (IMBD) scores.



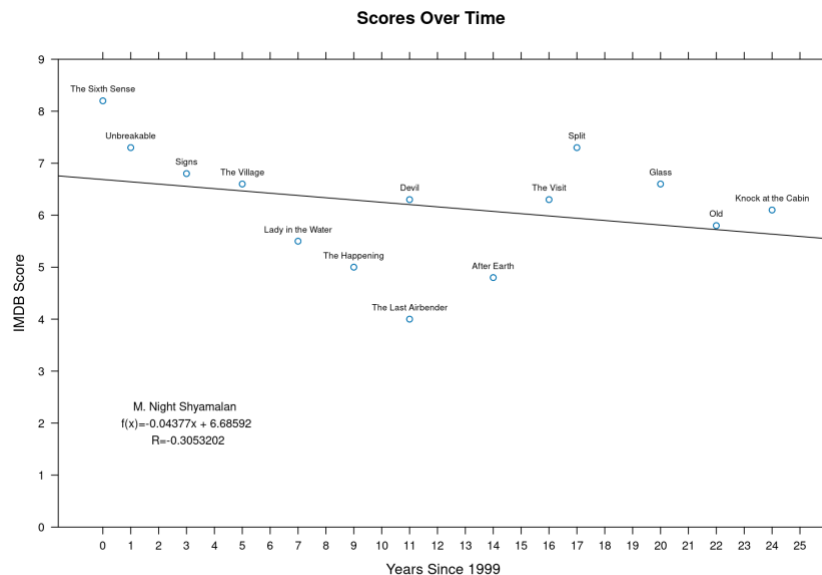
LMR_U4_L11_A



2. In teams, students will discuss and record their responses to the following questions for each plot:



- What kind of plot is this? *Answer: Scatterplot.*
- What do the numbers on the x-axis represent? What do the numbers on the y-axis represent? *Answer: The x-axis shows number of explosions and y-axis shows box office gross in millions of dollars.*
- What is this plot telling us? *Answers will vary. One example could be that if there are more explosions in a movie, then the movie will earn a greater box office gross.*



- What kind of plot is this? *Answer: Scatterplot.*
- What do the numbers on the x-axis represent? What do the numbers on the y-axis represent? *Answer: The x-axis shows the number of years since 1999 and the y-axis shows the movie's IMDB score.*
- What is this plot telling us? *Answers will vary. One example could be that as M. Night Shyamalan has produced more movies, their IMDB ratings have gone down.*

3. Allow students time to discuss and record their answers to the questions.



4. Display both plots, if possible (students may also refer to the plots in their own handout).

Discuss the following questions with the whole class:

- a. What is happening in each plot? What seems to be the trend? *Guide students to understand that the Box Office Gross per Explosion plot shows an increasing trend, while the Scores Over Time plot shows a decreasing trend. An increasing trend is called a positive association and a decreasing trend is called a negative association.*
- b. What does it mean to have an increasing trend and a positive association? *Answer: In Box Office Gross per Explosion, it means that as the number of explosions increase, the movie box office gross also increase.*
- c. What does it mean to have a decreasing trend and a negative association? *Answer: In Scores Over Time, it means that as the years after 1999 pass, the movie IMBD ratings decrease.*
5. Quickwrite: What if we had a plot with **no association**? Ask students to sketch what they think a scatterplot that shows no association looks like. *Answer: A correct sketch will show a scatterplot with data points that show no positive or negative association; no trend or pattern. There would be no association or a very weak one. The data would be scattered.*
6. Select a couple of sketches to share with the whole class. Discuss why the sketches show no association.
7. Ask students to discuss their thoughts about why a line was drawn through the points of the two plots and why there are equations for each plot.
8. Conduct a share out of their observations. Guide students to the understanding that both plots follow a **linear** trend. The line then represents a **model** for the relationship between the two variables. The equations shown in the plots above represent the lines through the points. They provide a description of the data and the relationship between the variables.
9. Ask student teams to refer back to the *What's the Trend?* handout (LMR_U4_L11_A). They should discuss the following questions and record their responses on the *Predicting Values* handout (LMR_U4_L11_B):
 - a. What do you notice about where the points are and where the line is? *Answer: Some points are near the line, others are further away, and one point is just slightly touching the line (but not exactly on it).*
 - b. Recall from Algebra that every line can be represented by an equation in the form $y = mx + b$. In this case, the equation of the regression line is $y = 0.685x + 94.093$. What do the x - and y -values represent in this equation? *Answer: The x -values represent the number of explosions and the y -values represent the predicted box office gross.*
 - c. According to the equation, what is the slope of the line? What does the slope mean in relation to the number of explosions? *Answer: The slope is 0.685. It is the rate of change between the number of explosions and the box office gross. It means that for every explosion increase of 1 the box office gross increases by 0.685 dollars.*
 - d. When the number of explosions (x -value) is zero, what is the box office gross (y -value)? How do you know? What does this mean? *Answer: The box office gross is 94.093 million dollars. Students may use the equation to show that they substituted zero for x , so the y -intercept is the box office gross. It means that if Michael Bay were to make a movie with NO explosions, this would be his projected box office gross.*
 - e. If you wanted to know the box office gross for the point that lies the closest to the line, what would the equation be? Write the equation and solve it. *Answer: Box Office Gross = $0.685(380) + 94.093 \rightarrow$ Box Office Gross = 354.393 or 354,393,000 million dollars.*
 - f. What was the actual box office gross for the point that lies closest to the line? *Answer: The actual profit was 352,400,000 million dollars.*
 - g. What if Michael Bay made a movie that had 275 explosions? What would his predicted box office gross be? Show how you arrived at the solution. *Answer: By substituting 275 in the value of x in the equation, predicted box office gross will be \$282,468,000 or 282.468 in millions of dollars, or by finding the point on the line or both.*

12. If time permits, have students answer the following questions about the *Scores Over Time* scatterplot in LMR_U4_L11_A:

- What do you notice about where the points are and where the line is? *Answer: Some points are near the line, others are further away, and three points are very close to the line (with one slightly touching it but not exactly on the line).*
- What do the x- and y-values represent in this equation? *Answer: The x-values represent the number of years since 1999 and the y-values represent the IMDB score.*
- According to the equation, what is the slope of this line? What does the slope mean? *Answer: The slope is -0.04377 . It is the rate of change between the number of years since 1999 and the IMDB score. It means that for every year since 1999 increase of 1 the IMDB score decreases by 0.04377.*
- When the x-value is zero, what is the y-value? How do you know? What does this mean? *Answer: The IMDB score, the y-value, is 6.68592. Students may use the equation to show that they substituted zero for x, so the y-intercept is the IMDB score. It means that if M. Night Shyamalan were to make a movie in 1999, this would be his projected IMDB score, regardless of the movie type or other factors. While knowing the y-intercept allows us to make predictions using this model, the intercept does not have any meaningful interpretation for this model (it's not possible to go back in time to have M. Night Shyamalan produce another movie in 1999).*
- What would the predicted value of the score be if M. Night Shyamalan released a movie in 2030? How do you know? *Answer: By substituting 31, because $2030 - 1999 = 31$, in the value of x in the equation, the predicted IMDB score will be 5.32905, or by finding the point on the line or both.*

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students will finish answering the questions above about the *Scores Over Time* scatterplot in LMR_U4_L11_A referenced above.

Lesson 12: How Strong Is It?

Objective:

Students will learn that the correlation coefficient is a value that measures the strength in linear associations only.

Materials:

1. *Strength of Association* handout (LMR_U4_L12_A)
2. *Correlation Coefficient* handout (LMR_U4_L12_B)

Advance preparation required: This handout is the resource for the plot cutouts. DO NOT distribute as-is to students (see steps 6 and 10 in the lesson)

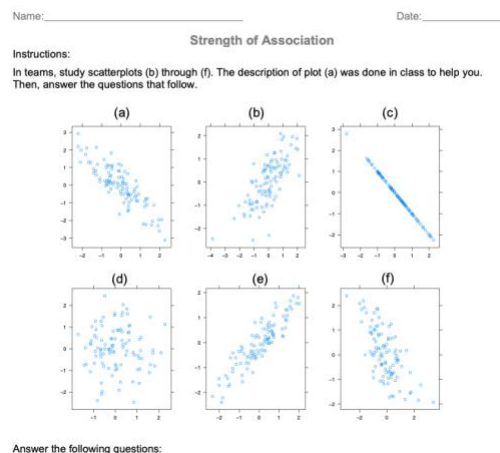
Vocabulary:

correlation coefficient, strength of association

Essential Concept: A high absolute value for correlation means a strong linear trend. A value close to 0 means a weak linear trend.

Lesson:

1. Distribute the *Strength of Association* handout (LMR_U4_L12_A). In teams, students will examine the scatterplots (b) through (e). Their task is to discuss the **strength of the association** for each plot. They will determine which plots they think show strong associations and which ones they believe show weak associations. They must explain how they made their decision. Reasons must reference the plots.



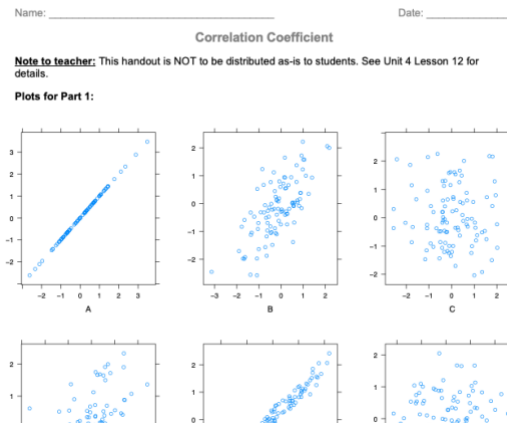
LMR_U4_L12_A

2. As an example, demonstrate how to describe plot (a) in the *Strength of Association* handout.
Possible description: Plot (a) shows a negative association, or decreasing trend. The association appears to be fairly strong because the points are relatively close together, forming a moderate linear pattern.
3. Once all teams have completed the handout, assign one plot to each team for a share out. If two teams have the same plot, one team will share its explanation first and the second team can agree, disagree, or add to the first team's description.
4. Guide students to understand that a strong association has points closer to each other and a weak association has points more scattered.
5. Inform students that, so far, they have been labeling associations as strong, very strong, or weak. A number called the **correlation coefficient** measures strength of association. The correlation coefficient only applies to linear relationships, which must be checked visually with a scatterplot. Later we will learn how to calculate this number using RStudio.



- Distribute the envelopes, with plots from LMR_U4_L12_B (Part 1), to the teams. Students will examine the strength of association in each plot. Their task is to assign a correlation coefficient that corresponds to each plot and to explain why they assigned that correlation coefficient to that particular plot. The only piece of information they will receive is that a correlation coefficient equal to 1 has the strongest linear association and a correlation coefficient equal to 0 has the weakest association.

Advance preparation required: Each team needs one envelope with cutouts of plots A-F in LMR_U4_L12_B (Part 1). Make envelopes according to the number of teams in the class.



LMR_U4_L12_B (Part 1)

- Assign each team one plot. If there are more teams than plots, these teams will be assigned a plot in the next round. Each team will share the correlation coefficient they assigned to their plot and the explanation that goes with it.
- Using the *Voting Cards* strategy (see Instructional Strategies), the rest of the teams will show whether they A – approve, B – disapprove, or C – are uncertain, about the teams' assignment and/or explanation. Repeat for each plot. The correlation coefficients for each plot are:

Plot A: $r = 1.00$

Plot B: $r = 0.72$

Plot C: $r = 0.19$

Plot D: $r = 0.48$

Plot E: $r = 0.98$

Plot F: $r = 0.00$

- The last set of plots showed positive associations. Now students will assign the correlation coefficients for plots G-L for LMR_U4_L12_B (Part 2).



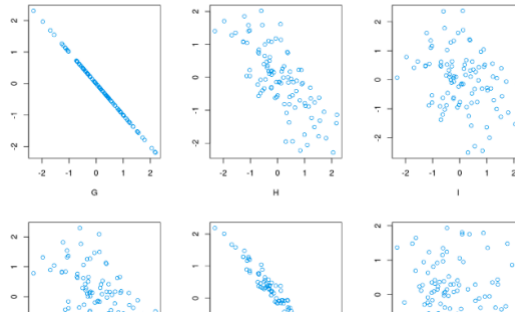
- Distribute the envelopes, with the plots from LMR_U4_L12_B (Part 2), to the teams. Students will examine the strength of association in each plot. Their task is to assign the correlation coefficient that corresponds to each plot and to explain why they assigned that correlation coefficient to that particular plot. The only piece of information they will receive is that a correlation coefficient equal to -1 has the strongest linear association and a correlation coefficient equal to 0 has the weakest association.

Advance preparation required: Each team needs one envelope with cutouts of plots G-L in LMR_U4_L12_B (Part 2). Make envelopes according to the number of teams in the class.

Name: _____ Date: _____

Correlation Coefficient

Plots for Part 2:



LMR_U4_L12_B (Part 2)

11. Teams previously not assigned a plot are now assigned one. Each team will share the correlation coefficient they assigned to their plot and the explanation that goes with it.



12. Using the *Voting Cards* strategy, the rest of the teams will show whether they A – approve, B – disapprove, or C – are uncertain, about the teams' assignment and/or explanation. Lead a class discussion whenever there is disapproval or uncertainty. Repeat for each plot. The correlation coefficients for each plot are:

Plot G: $r = -1.00$

Plot H: $r = -0.72$

Plot I: $r = -0.19$

Plot J: $r = -0.48$

Plot K: $r = -0.98$

Plot L: $r = 0.00$



13. Journal Entry: What is a correlation coefficient, what does it do, and what does it tell us about a scatterplot?

Homework & Next Day

Students will complete the journal entry for homework if not completed in class.

Lab 4D: Interpreting correlations

Complete Lab 4D prior to Lesson 13.

Lab 4D: Interpreting correlations

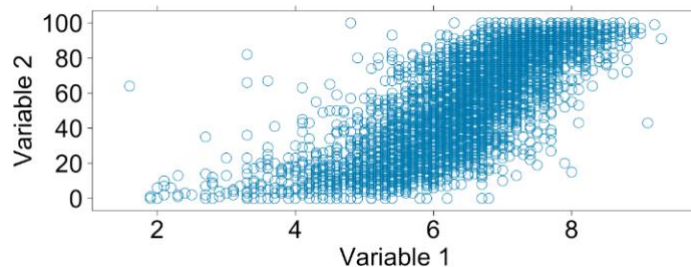
Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Some background...

- So far, we've learned about measuring the success of a model based on how close its predictions come to the actual observations.
- The *correlation coefficient* is a tool that gives us a fairly good idea of how these predictions will turn out without having to make predictions on future observations.
- For this lab, we will be using the **movie** dataset to investigate the following question:
*Which variables are better predictors of a movie's **critics_rating** when the predictions are made using a line of best fit?*

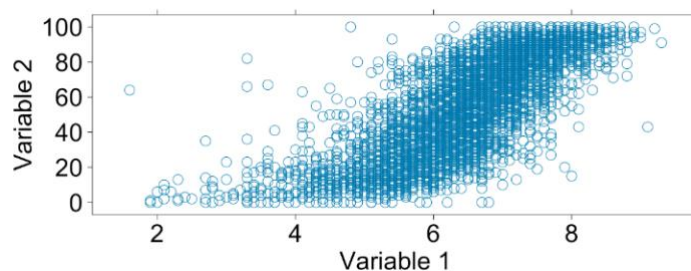
Correlation coefficients

- The *correlation coefficient* describes the *strength* and *direction* of the linear trend.
- It's only useful when the trend is linear and both variables are numeric.



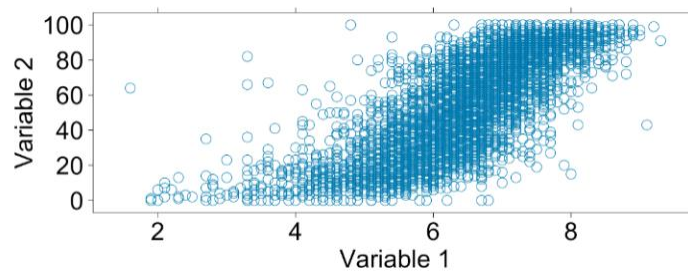
- **(1) Are these variables linearly related? Why or why not?**

Correlation review I



- Correlation coefficients with values close to 1 are very strong with a positive slope. Values close to -1 means the correlation is very strong with a negative slope.
- **(2) Does this plot have a positive or negative correlation?**

Correlation review II



- Recall that if there is no linear relation between two numerical variables, the correlation coefficient is close to 0.
- (3) What do you guess the correlation coefficient will be for these two variables?**

The movie data

- (4) Write and run code loading the movie data using the data command.**
- The data comes from a variety of sources like *IMDB* and *Rotten Tomatoes*.
 - The `critics_rating` contains values between 0 and 100, 100 being the best.
 - The `audience_rating` contains values that range between 0 and 10, 10 being the best.
 - `n_critics` and `n_audience` describe the number of reviews used for the ratings.
 - `gross` and `budget` describes the amount of money the film made and took to make.

Calculating Correlation Coefficients!

- We can use the `cor()` function to find the particular correlation coefficient of the variables from the previous plot, which happen to be `audience_rating` and `critics_rating`.
- But note, the `cor()` function removes any observations which contain an `NA` value in either variable.
- (5) Write and run code calculating the correlation coefficient for these variables using the `cor()` function. The inputs to the functions work just like the inputs of the `xyplot` function.**

Now answer the following.

- (6) What was the value of the correlation coefficient you calculated?**
- (7) How does this actual value compare with the one you estimated previously?**
- (8) Does this indicate a strong, weak, or moderate association? Why?**
- (9) How would the scatterplot need to change in order for the correlation to be stronger?**
- (10) How would it need to change in order for the correlation to be weaker?**

Correlation and Predictions

- (11) Find the two variables that look to have the strongest correlation with `critics_rating`.**
 - (12) Compute the correlation coefficients for `critics_rating` and each of the two variables.**
 - (13) Use the correlation coefficient to determine which variable has a stronger linear relationship with `critics_rating`.**

- (14) Write and run code fitting two `lm()` models to predict `critics_rating` with each variable and compute the MSE for each.
 - (15) Use the MSE to determine which variable is a better predictor of `critics_rating`.
- (16) How are the correlation coefficient and the MSE related?

On your own

- (17) Select two different numerical variables from the `movie` data. Plot the variables using the `xyplot()` function.
 - (18) Would calculating a correlation coefficient for the two variables be appropriate? Justify your answer.
 - (19) Predict what value you think the correlation coefficient will be. Compare this value to the actual value. Finally, interpret what the actual correlation coefficient means.
- (20) Work with your classmates to determine which two variables have the strongest correlation coefficient. Write them down.
 - (21) Why do you think these variables are so strongly related?
 - (22) Is using the correlation coefficient to describe the relationship appropriate and why/why not?

Lesson 13: Improving Your Model

Objective:

Students will learn to describe associations that are not linear.

Materials:

1. *Describe the Association* handout (LMR_U4_L13)

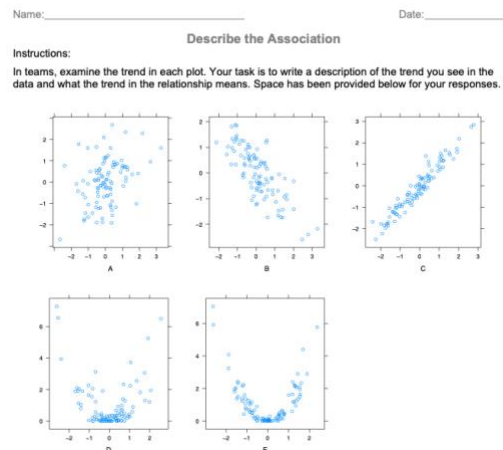
Vocabulary:

non-linear, polynomial trends

Essential Concepts: If a linear model is fit to a non-linear trend, it will not do a good job of predicting. For this reason, we need to identify non-linear trends by looking at a scatterplot or the model needs to match the trend.

Lesson:

1. Remind students that they have been learning a great deal about linear associations. However, there are other types of associations, and today they will learn to describe them.
2. Distribute the *Describe the Association* handout (LMR_U4_L13). In teams, students will examine the trend of each plot. Their task is to write a description of the trend that they see in the data and what the trend means.



LMR_U4_L13



3. Allow students time to discuss and record their descriptions for each plot in their DS journals. Walk around the room monitoring student teamwork. Look for descriptions that are interesting to share with the whole class.
4. Select a team to present a description of one plot to the class. Teams will listen to each presentation, compare it to their description of the plot, and as a team they will agree or disagree. If there is disagreement, lead a discussion that guides students to reason toward the correct description.
5. Summarize the discussion for each plot and ask students to take notes or revise their descriptions in their DS journals.
6. Repeat steps 4 and 5 from above for the rest of the plots.

Plot Descriptions for *Describe the Association* (LMR_U4_L13):

- *Plot A: There is no trend (perhaps some may see a very, very weak linear trend), so there is no/hardly any association. There is a great deal of scatter in the data. It means that y does not depend on x.*
 - *Plot B: There appears to be a linear trend. The association is negative and appears somewhat strong. It means that as x increases, y decreases.*
 - *Plot C: There is a linear trend. The association is positive and it is very strong. It means that the y-value increases at approximately the same rate for every increase in x value. This is a line.*
 - *Plot D: The trend is non-linear. There seems to be a weak association because there is scatter in the data. We cannot tell if the association is positive or negative. It has the shape of a parabola; therefore, it is quadratic. For smaller x-values, the y-value is decreasing and for larger x values, the y value is increasing.*
 - *Plot E: The trend is non-linear. There seems to be a strong association because there is little scatter in the data. It is also in the shape of a parabola, so it is quadratic.*
7. Using the *Cheat Notes* strategy, ask teams to write notes about how to describe associations.
 8. Plots A, B, and C should be familiar to the students by now. However, plots D and E show a different type of trend. Although the trends are **non-linear**, they can still tell us important information about the y-values based on values of x. Ask:
 - a. What happens if we were to fit a linear model to these non-linear trends? Would it still make good predictions? *Answer: Fitting a linear model to a non-linear trend would not properly describe the trend of the data. Therefore no, it would not make good predictions.*
 9. To examine why they would not make good predictors, draw an approximate linear best-fit line and get students to understand that in some regions, the model would almost always over-predict, and in others would almost always under-predict. We want a model that goes, more or less, through the 'middle' of the points. Ask:
 - a. How can we get a model that goes, more or less, through the middle of all the data points? *Answer: We need to change the model.*
 10. Trends like the quadratic ones shown in plots D and E can be described as **polynomial trends**. Other polynomial trends include cubic trends, quartic trends, etc. You may show students several choices of equations (linear, quadratic, cubic) along with their graphs and ask them which might be a good candidate to represent them.
 11. When investigating the data for trends, the model needs to fit the data.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next Day

Students may finish their *Cheat Notes* for homework, if not completed in class.

Lab 4E: Some models have curves

Complete Lab 4E prior to Practicum.

Lab 4E: Some models have curves

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Making models do yoga

- So far, we have only worked with prediction models that fit the *line of best fit* to the data.
- What happens if the true relationship between the data is nonlinear?
- In this lab, we will learn about prediction models that fit *best fitting curves* to data.
- **(1) Before moving on, load the movie data and write and run code splitting it into two sets:**
 - A set named **training** that includes 75% of the data.
 - And a set named **test** that includes the remaining 25%.
 - Remember to use `set.seed`.

Problems with lines

- Before learning how to fit curves, let's first fit a linear model for reference.
- **(2) Write and run code training a linear model predicting audience_rating based on critics_rating for the training data. Assign this model to movie_linear.**
- **(3) Fill in the blanks below to create a scatterplot with audience_rating on the y-axis and critics_rating on the x-axis using your test data.**

```
xyplot(____ ~ ____, data = ____)
```

- Previously, you used `add_line` to plot the *line of best fit*. An alternative function for plotting the *line of best fit* is `add_curve`, which takes the name of the model as an argument.
- Run the code below to add the *line of best fit* for the training data to the plot.

```
add_curve(movie_linear)
```

- **(4) Describe, in words, how the line fits the data. Are there any values for critics_rating that would make obviously poor predictions?**
 - Hint: how does the linear model perform on very low and very high values of `critics_rating`?
- **(5) Compute the MSE of the model for the test data and write it down for later.**
 - Hint: refer to Lab 4C.

Adding flexibility

- You don't need to be a full-fledged Data Scientist to realize that trying to fit a line to curved data is a poor modeling choice.
- If our data is curved, we should try to model it with a curve.
- Instead of fitting a line, with equation of the form

$$y = a + bx$$

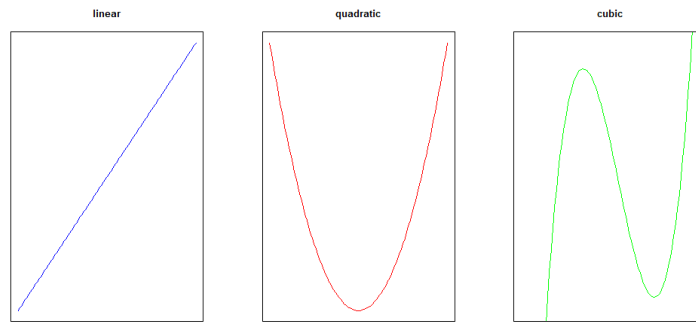
- we might consider fitting a *quadratic curve*, with equation of the form

$$y = a + bx + cx^2$$

- or even a *cubic curve*, with equation of the form

$$y = a + bx + cx^2 + dx^3$$

- In general, the more coefficients in the model, the more flexible its predictions can be.



Making bend-y models

- To fit a quadratic model in R, we can use the `poly()` function.
 - (6) Fill in the blanks below to train a quadratic model predicting `audience_rating` from `critics_rating`, and assign that model to `movie_quad`.

```
movie_quad <- lm(____ ~ poly(____, 2), data = training)
```

- (7) What is the role of the number 2 in the `poly()` function?

Comparing lines and curves

- (8) Fill in the blanks to
 - create a scatterplot with `audience_rating` on the y-axis and `critics_rating` on the x-axis using your test data, and
 - add the *line of best fit* and *best fitting quadratic curve*.
 - Hint: the `col` argument is added to the `add_curve` functions to help distinguish the two curves.

```
xyplot(____ ~ ____, data = ____)  
add_curve(____, col = "blue")  
add_curve(____, col = "red")
```

- (9) Compare how the *line of best fit* and the *quadratic model* fit the data. Which do you think has a lower `test MSE`?
- (10) Compute the `MSE` of the quadratic model for the `test` data and write it down for later.
- (11) Use the `test MSE` to explain why one model fits better than the other.

On your own

- (12) Write and run code creating a model that predicts `audience_rating` using a cubic curve (polynomial with degree 3), and assign this model to `movie_cubic`.
- (13) Write and run code creating a scatterplot with `audience_rating` on the y-axis and `critics_rating` on the x-axis using your test data.
- Using the names of the three models you have trained, add the *line of best fit*, *best fitting quadratic curve*, and *best fitting cubic curve* for the training data to the plot.
- (14) Based on the plot, which model do you think is the best at predicting the `test` data?
- (15) Use the `test MSE` to verify which model is the best at predicting the `test` data.

Practicum: Predictions

Objective:

Students will create a linear model to predict the nutritional component that is most closely associated with the amount of sugar contained in a cereal.

Materials:

1. *Predictions Practicum* (LMR_U4_Practicum_Predictions)

Practicum Predictions

Data about the nutritional components of popular cereal brands has been collected and made available for your team's use. We are interested in determining which other nutritional component is most closely associated with the amount of sugar contained in a cereal.

Your team will use the data to make predictions using linear models and compare the accuracy of your model to the rest of your classmates. Finally, the class will determine which team had the best prediction. Follow the directions below to explore and analyze the data:

1. You will have two datasets: one training named **cereal** and one test named **cereal_test**. Load both datasets. Write down the code you used.
2. Explore the training data. Which variable do you think is the best predictor of sugar? Choose at least 3 variables, make a plot for each one, and fit a linear regression line through each of them. Select the model that you think best makes the best prediction.
3. For the linear model your team selected:
 - a. Describe what the plot shows.
 - b. Explain why you selected that particular model.
 - c. Compute the mean squared error of your model using your test data.
 - d. Now make a set of predictions with your test data. Calculate the mean squared error for the test data. Is it better or worse than for the training data, or about the same?
4. Present your team's linear model to the class. Explain why you chose your model and the typical amount of error in its predictions.
5. Give an example of a prediction for one value of x . State that value, give the predicted sugar, and describe, based on the test data, how far off your prediction might actually be.

Piecing it Together

Instructional Days: 4

Enduring Understandings

Real-life phenomena are often complex. Data scientists use multiple regression models to create simple equations to help explain and predict these phenomena. Data scientists can also use polynomial transformations to add flexibility to rigid linear models.

Engagement

Students will read the article titled *How Long Can a Spinoff Like Better Call Saul Last?* that will set the context for students to begin thinking about more than one explanatory variable to make better predictions. The article can be found at:

<http://fivethirtyeight.com/features/how-long-can-a-spinoff-like-better-call-saul-last/>

Learning Objectives

Statistical/Mathematical:

A-REI D-10: Understand that the graph of an equation in two variables is the set of all its solutions plotted in the coordinate plane, often forming a curve (which could be a line).

S-ID 6: Represent data on two quantitative variables on a scatter plot, and describe how the variables are related.

- a. Fit a function to the data; use functions fitted to data to solve problems in the context of the data. *Use given functions or choose a function suggested by the context. Emphasize linear models.*

Data Science:

Understand that multiple regression can be a better tool for predicting than simple linear regression and know when it is appropriate to use multiple regression versus simple linear regression. Understand when linear models are not appropriate based on the shape of the scatterplot.

Applied Computational Thinking using RStudio:

- Use multiple linear regression models with other predictor variables.
- Fit regression lines to data and predict outcomes.
- Fit polynomials functions to data.

Real-World Connections:

Economists and marketing firms use multiple regression to predict changes in the market and adjust strategies to fit the demands of changes in the marketplace.

Language Objectives

1. Students will read informative texts to evaluate claims based on data.
2. Students will engage in partner and whole group discussions about how adding variables to a model will help or hinder its predictions.
3. Students will construct their own linear model using multiple variables to compare and contrast which model makes the best predictions.

Data File or Data Collection Method

Data File:

1. Movies: `data(movie)`
2. Cereal brands: `data(cereal)`

Data Collection Method:

Team-generated Participatory Sensing Campaign: Students will collect data for a team-selected topic.

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Lesson 14: More Variables to Make Better Predictions

Objective:

Students will see that information from different variables can be used together to create linear models that make more accurate predictions.

Materials:

1. *Advertising Plots Part 1* handout (LMR_U4_L14_A)
2. *Advertising Plots Part 2* handout (LMR_U4_L14_B)
3. Article: *How Long Can a Spinoff Like 'Better Call Saul' Last?*
<http://fivethirtyeight.com/features/how-long-can-a-spinoff-like-better-call-saul-last/>

Vocabulary:

market

Essential Concepts: We can use scatterplots to assess which variables might lead to strong predictive models. Sometimes using several predictors in one model can produce stronger models.

Lesson:

1. Remind students that models are used to make predictions. Ask a volunteer to think of a TV show that had a “spinoff” and to name both of the shows. Ask if he/she knows whether or not the original was more or less successful than the spinoff. Then, ask the class: Is there a way to predict spinoff success?



2. Next, using the *Talking to the Text* instructional strategy, ask students to read the article titled: *How Long Can a Spinoff Like Better Call Saul Last?*

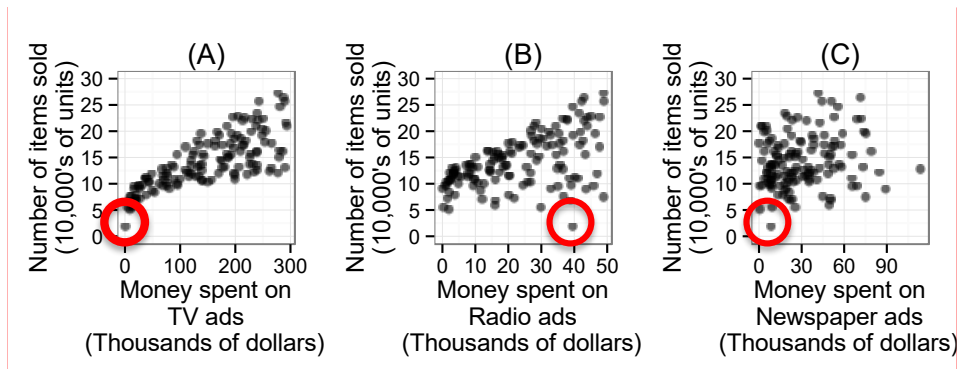
Note: If this is the first time using this strategy with your students, make sure you model/explain it before they begin reading it. See *Instructional Strategies in Teacher Resources* for a description.

3. After reading the article, ask students to discuss three *Talking to the Text* responses with a partner. You may set a time limit for each student to share with his/her partner.



4. Then, in teams, students will answer the following questions pertaining to the article:
 - a. What is the article trying to predict? *Answer: The success of a spinoff show.*
 - b. How many variables are used? *Answer: Four – the original show name, the number of episodes of the original show, the name of the spinoff show, and number of episodes of the spinoff show.*
 - c. What other variables might affect a spinoff? *Possible answers are budget or actors.*
 - d. The dotted line in the plot is not a regression line. How would you draw a regression line to make predictions? *Answers will vary but we would want to try to “fit” a line to the plotted data.*
 - e. What other information would you like to know to predict a spinoff’s success? *Answers will vary but may be similar to (c) above.*
5. Allow students time to discuss and record their answers. Then conduct a share out of their responses to the discussion questions.
6. Discuss the following questions with the class:
 - a. What effect does advertising have on retail sales?
 - b. Where do stores advertise (What mediums do they use)? Does each method of advertisement reach the same people?
 - c. Does each method of advertisement have a similar effect? Or are some methods more effective than others?

7. Distribute the 3 plots from the *Advertising Plots Part 1* handout (LMR_U4_L14_A) and inform the students about the data using the details below:



LMR_U4_L14_A

(Plots are presented separately in the LMR)

- These 3 plots show the number of items sold by a retailer (in 200 different markets) and the amount of money the company spent on TV, Radio, and Newspaper advertisements.
 - The data has 200 observations, one for each different market. A **market** is simply a location where an item is sold. For example, Los Angeles and San Francisco are two different markets.
 - Each observation has 4 variables: (1) The number of items sold (in 10's of thousands of units), (2) the money spent on TV ads (in thousands of dollars), (3) the money spent on radio ads (in thousands of dollars), and (4) the money spent on newspaper ads (in thousands of dollars).
 - The data were collected using an observational study.
8. To illustrate a-d above, ask students to refer to plot A (TV ads) and circle the market in which this retailer sold the least number of items (see circles in plots above). Ask:
- Approximately, how many items did this market sell? *Answer: About 20,000 items. The actual number of items sold was 1.6 (in 10,000's of units) which is 16,000 items.*
 - Approximately, how much money did this retailer spend on TV ads in this market? *Answer: This retailer spent zero dollars on TV ads. The actual amount the retailer spent on TV ads was 0.7 thousands of dollars, which is \$700.*
9. Students should then refer to plot B (Radio ads), find the same market (the one in which the retailer sold about 20,000 items) and circle it. Ask:
- Approximately, how much money did the retailer spend on Radio ads in the same market? *Answer: About 40 thousand dollars. The actual amount spent on Radio ads was 39.6 thousands of dollars, which is \$39,600.*
10. Finally, ask students to refer to plot C (Newspaper ads), find the same market (the one in which the retailer sold about 20,000 items), and circle it. Ask:
- Approximately, how much money did the retailer spend on Newspaper ads in the same market? *Answer: About 10 thousand dollars. The actual amount spent on Newspaper ads is 8.7 thousands of dollars, which is \$8,700.*

| TV | Radio | Newspaper | Sales |
|-----|-------|-----------|-------|
| 0.7 | 39.6 | 8.7 | 1.6 |

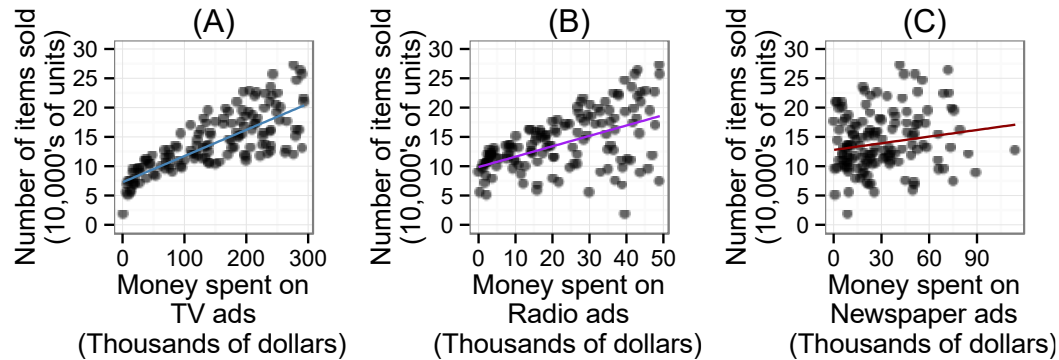


11. Based on the above plots, use a *Pair-Share* to discuss the following:

- Describe the relationship between advertisements and the number of items sold. *Answers will vary but we would expect the number of items sold to increase with increased advertisements.*

- b. Which type of advertisement is the most strongly correlated with the number of units sold? How can you tell? *Answer: Plot A, TV advertisements, appears to be the most strongly correlated with the number of units sold. We can tell because the points are more closely packed/together than in plots B or C.*

12. Distribute the *Advertising Plots Part 2* handout (LMR_U4_L14_B) which contains plots A-C, but now include the line of best fit.



LMR_U4_L14_B

(Plots are presented separately in the LMR)

13. Ask students to recall from Lesson 7 that a method statisticians use to figure out which predicted values is closest to the actual data is the mean squared error (MSE).



14. In teams, ask students to discuss the following:

- a. How would you determine which plot would make the most accurate predictions? *Answers will vary, but you would expect to hear something like: "the prediction line that has the least amount of distance to all the points on the plot would make the most accurate prediction because the predicted values will be closer to the actual data."*



15. Next, have students select a statement they think is best (a or b), then write a justification for their selection based on what they learned in this lesson. This may be completed as homework.

- a. Combining multiple variables (e.g., TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to worse predictions because the variables that make poor predictions will contaminate those that make good predictions.
- b. Combining multiple variables (e.g., TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to better predictions because the model can use more information to make predictions.

16. Inform students that RStudio has the capability of creating models that combine multiple variables to make predictions about another variable. For example, it can make a model to predict number of items sold using both money spent on TV and money spent on Newspaper ads. Students will learn more about it during the next lesson.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework

Students may continue writing their justifications for the selected statement in item 15 if they were unable to finish during class.

Lesson 15: Combination of Variables

Objective:

Students will learn that we can make better predictions by including more variables. Then they will wrestle with how the information should be combined.

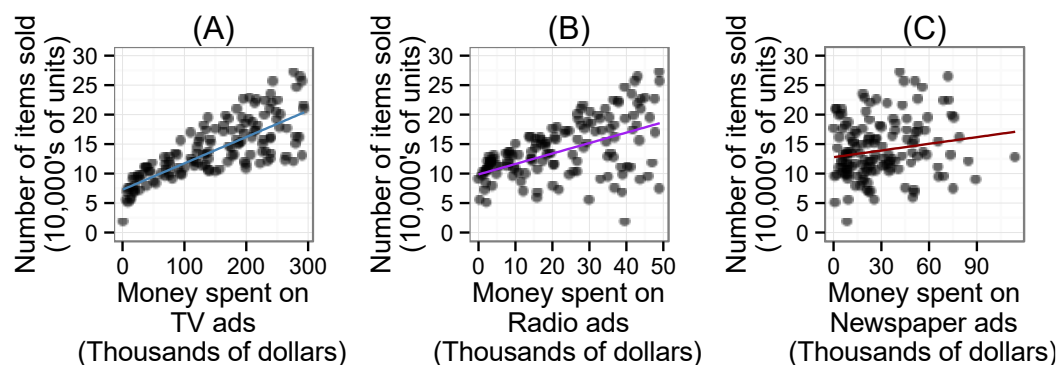
Materials:

1. *Advertising Plots Part 2* handout (LMR_U4_L14_B) from Lesson 14

Essential Concepts: If multiple predictors are associated with the response variable, a better predictive model will be produced, as measured by the mean squared error.

Lesson:

1. Display the plots and statements from the previous day:



LMR_U4_L14_B

(Plots are presented separately in the LMR)

- a. Combining multiple variables (e.g., money spent on TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to worse predictions because the variables that make poor predictions will contaminate those that make good predictions.
- b. Combining multiple variables (e.g., TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to better predictions because the model can use more information to make predictions.



2. Ask the students to share out their opinions using the *Agree/Disagree* strategy.
3. Next, inform teams that they will have 2 minutes to come up with as many combinations of ads (variables) as they can think of (e.g., TV + Newspaper ads, TV+ Radio ads, TV + Radio + Newspaper ads, etc.)
4. After 2 minutes, list all the different combinations by conducting a *Whip Around* and eliciting a combination from each team.
5. By a show of hands, ask students to select which combination or single model will be the best predictor for the number of items sold by the retailer.
6. Then inform students that we will determine which of the statements is true by comparing the mean squared error (MSE) of single models (like the ones we showed in the previous lesson) vs. combined models. But first, use the line of best fit for the combined variables:

$$\widehat{sales} = 0.044889(tv) + 0.194199(radio) - 0.007392(newspaper) + 3.154843$$

Note: The function that produced the line of best fit using RStudio was:

```
lm(Sales ~ TV + Radio + Newspaper, data = retail)
```

- a. Use this equation to predict the amount of sales for the same market they circled in the previous lesson. *Student calculations should yield the predicted value in (b), below.*

Note: Remind students that they need to substitute the values as they appear in the x-axis of the plots without converting to thousands of dollars. For example, the circled market spent about 10 thousand dollars on newspaper ads, so students should substitute 10 instead of the expanded value in the equation.

| TV | Radio | Newspaper | Sales |
|-----|-------|-----------|-------|
| 0.7 | 39.6 | 8.7 | ??? |

- b. Does the predicted value (10.81224) seem like a plausible number of sales? Why? *Answer: It is not a plausible number of sales because the prediction is too high. The prediction says the retailer will sell about 108,122 units, when the actual sales were about 16,000 units. Although the model did not make a very good prediction for this market, it is not surprising because as LMR_U4_L14_B displays, that market did not fit the overall pattern in any of the scatterplots.*

| TV | Radio | Newspaper | Sales |
|-----|-------|-----------|-------|
| 0.7 | 39.6 | 8.7 | 1.6 |

7. Reveal that RStudio calculated the mean squared error for different combinations plus the single models, and the results are displayed on the table below. This means that, for example, when using the TV model to predict number of items sold, our predictions will typically be off by about 3.373579 (in 10,000s) of units or 33,736 units. Then ask students the questions below.

Note: Remember that the MSE will always be in square units. In order to convert back to the original units, simply take the square root of the MSE.

| Model | Mean Squared Error (MSE) | Square Root MSE |
|--------------------|--------------------------|-----------------|
| TV | 11.38103 | 3.373579 |
| Radio | 25.35521 | 5.035396 |
| Newspaper | 31.44164 | 5.607285 |
| TV-Radio | 4.456745 | 2.1111 |
| TV-Newspaper | 9.246567 | 3.040817 |
| Radio-Newspaper | 27.26889 | 5.221963 |
| TV-Radio-Newspaper | 4.70338 | 2.168728 |

- a. Which model is the best predictor of number of items sold? *Answer: The TV-Radio model is the best predictor of number of items sold because it had the least amount of error, on average. When using the TV-Radio model to predict number of items sold, our predictions will typically be off by 21,111 units.*
- b. Which model was the least reliable in predicting the number of items sold? *Answer: The Newspaper model is the least reliable predictor of number of items sold because it had the most amount of error, on average. When using the Newspaper model to predict number of items sold, our predictions will typically be off by 56,073 units.*
- c. What else do you notice about the models? *Answer: It appears that combining the variables into one model is much better than any of the single-variable models.*



8. Think back to our statements from the last lesson and the beginning of this lesson:
- a. Combining multiple variables (e.g., money spent on TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to worse predictions because the variables that make poor predictions will contaminate those that make good predictions.
- b. Combining multiple variables (e.g., TV and Newspaper ads, TV and Radio ads, TV, Radio, and Newspaper ads, etc.) into one model will lead to better predictions because the model can use more information to make predictions.



Engage students in a discussion to check understanding of using multiple variables in a model.

- a. Has their statement selection changed or stayed the same? Why?
- b. What evidence do they have to support their statement selection? *Answers will vary because students can make an argument for either statement, both statements, or neither statements being true. There is evidence in step 7 above that supports that more variables will lead to worse predictions (like TV-Radio-Newspaper) but there is also evidence that more variables lead to better predictions (like TV-Radio).*

9. Inform the students that, in the next lab, they will find out how to create the line of best fit for models that include many variables.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next Day

Ask students to think of a reason or reasons about why it would not be a good idea to make a scatterplot for models that include more than 3 predictor variables? *The answer is mainly because humans are limited to seeing things in 3 dimensions. For example, the model that combines all of the variables together is a 4-dimensional model. What does that look like?*

Lab 4F: This model is big enough for all of us

Complete Lab 4F prior to Lesson 16.

Lab 4F: This model is big enough for all of us!

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Building better models

- So far, in the labs, we've learned how to make predictions using the *line of best fit*, also known as *linear models* or *regression models*.
- We've also learned how to measure our model's prediction accuracy by cross-validation.
- In this lab, we'll investigate the following question:
Will including more variables in our model improve its predictions?

Divide & Conquer

- (1) Start by loading the **movie** data and write and run code splitting it into two sets (see Lab 4C for help).
 - A set named **training** that includes 75% of the data.
 - A set named **test** that includes the remaining 25%.
 - Remember to use **set.seed**.
- (2) Write and run code creating a linear model, using the **training** data, that predicts **gross** using **runtime**.
 - (3) Write and run code creating the MSE of the model by making predictions for the **test** data.
- (4) Do you think that a movie's **runtime** is the only factor that goes into how much a movie will make? What else might affect a movie's **gross**?

Including more info

- Data scientists often find that including more relevant information in their models leads to better predictions.
- (5) Fill in the blanks below to predict **gross** using **runtime** and **reviews_num**.

```
lm(____ ~ ____ + ____, data = training)
```
- (6) Does this new model make more or less accurate predictions? Describe the process you used to arrive at your conclusion.
- (7) Write down the code you would use to include a 3rd variable, of your choosing, in your **lm()**.

On your own

- (8) Write down which other variables in the **movie** data you think would help you make better predictions.
 - (9) Are there any variables that you think would not improve our predictions?
- (10) Write and run code creating a model for all of the variables you think are relevant.
 - (11) Assess whether your model makes more accurate predictions for the **test** data than the model that included only **runtime** and **reviews_num**.
- (12) With your neighbors, determine which combination of variables leads to the best predictions for the **test** data.

Decisions, Decisions

Instructional Days: 3

Enduring Understandings

Decision trees are used to classify observations into similar groupings based on known characteristics. Questions are asked, then the observations are sorted based on the responses to the questions. After a specified number of iterations, a final group membership is decided. One particular modeling tool we use for decision trees is known as CART (Classification and Regression Trees).

Engagement

Students will be presented with the question about whether they would rather trust a doctor or a data scientist to diagnose them if they were having a chest pains. This will set the context for decision trees and how they are used to make predictions.

Learning Objectives

Statistical/Mathematical:

S-IC 2: Decide if a specified model is consistent with results from a given data-generating process, e.g., using simulation.

Data Science:

Understand that classification and regression trees can be used to predict membership in groups.

Applied Computational Thinking using RStudio:

- Create classification and regression trees.

Real-World Connections:

Cardiologists may use a decision tree to diagnose whether people are or are not having a heart attack. Since the late 1870s, this method has been found to correctly diagnose a heart attack in over 95% of cases compared to correct diagnoses based on individual doctors' expertise, which ranged between 75 and 90%.

Language Objectives

1. Students will engage in partner and whole group discussions to express their understanding of classification trees.
2. Students will explain orally and in writing how to determine the accuracy of their non-linear model.
3. Students will make connections between decision trees and linear models in writing.

Data File or Data Collection Method

Data File:

1. USMNT/NFL dataset
Optional: `data(futbol)`
2. Titanic: `data(titanic)`

Data Collection Method:

Team-generated Participatory Sensing Campaign: Students will collect data for a team-selected topic.

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Lesson 16: Football or Futbol?

Objective:

Students will learn what decision trees look like and how they can be used to classify people or objects into groups. They will engage in an activity to see how making slight changes to the tree can lead to drastic rises or reductions in misclassifications.

Materials:

1. *Decision Tree for Heart Attack Risk* graphic (LMR_U4_L16_A)
2. *CART Activity Player Stats* (LMR_U4_L16_B)
3. *CART Activity Round 1 Questions* (LMR_U4_L16_C)
4. *CART Activity Round 2 Questions* (LMR_U4_L16_D)

Advance preparation required: LMR_U4_L16_B, C, & D need to be cut (see step 8 in lesson).

Vocabulary:

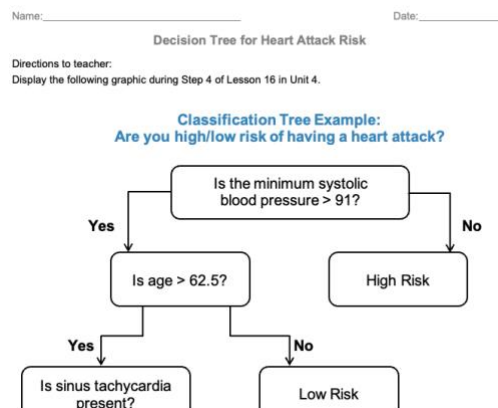
classify, decision tree, Classification and Regression Trees (CART), nodes

Essential Concepts: Some trends are not linear, so the approaches we've done so far won't be helpful. We need to model such trends differently. Decision trees are a non-linear tool for classifying observations into groups when the trend is non-linear.

Lesson:

1. Ask students the following question:
If you were having chest pains, who would you trust more to diagnose you – a data scientist or a doctor?
2. Give the students some time to think about the question and have a few of them share out their responses with the class.
Note: It's likely that most students will choose to go to a doctor.
3. As it turns out, back in the late 1970s, a cardiologist (and early data scientist) named Lee Goldman developed a decision tree based on millions of patient observations. It was made to diagnose whether people were or were not having a heart attack. The accuracy of the decision tree compared to the accuracy of actual doctor diagnoses are shown below:
 - a. Correct diagnoses using the decision tree were above 95%.
 - b. Correct diagnoses based on individual doctors' expertise was anywhere between 75-90%.
4. Display the graphic from the *Decision Tree for Heart Attack Risk* handout (LMR_U4_L16_A) and explain that this is one example of what the decision tree that Goldman developed might have looked like.

Note: This is NOT the actual tree Goldman developed.



LMR_U4_L16_A



5. Using a *Pair-Share*, ask students to discuss the following questions using the graphic above.
Note: Answers will vary. These questions are meant to gauge student thinking before defining decision trees in the following steps of the lesson.
 - a. What are decision trees?
 - b. How do they work at classifying data into groups?
6. Remind students that this unit has focused on linear models and making predictions. In the real world, data can be modeled in a variety of ways, many of which are non-linear, and because of this, we can't easily write down a mathematical equation to help us make predictions. However, we can use what we have learned so far to determine whether or not other models can provide a good fit to the data.
7. Let students know that one method of modeling data in a non-linear way is with **decision trees**, like the one we saw with the heart attack classification. Explain that decision trees are "grown" by using algorithms, or rules, to test many, many different decision trees to find the one that makes the best predictions.
8. A decision tree is basically a series of questions that are asked sequentially. Observations start by answering the first question (at the root of the tree) and then proceed along the different branches based on the answers they give to the questions that follow. At the end, based on all of the questions asked, observations are then classified as one of k classifications.
9. Remind students that algorithms are a series of steps that are repeated a large number of times. For decision trees, this enables us to (1) explore many possible paths, beginning from the same initial point, or (2) find different starting points based on where we ended during the previous iteration.
10. Inform students that, during today's lesson, they will be participating in an activity to try to **classify** professional athletes into one of two groups: (1) soccer players on the US Men's National Team, OR (2) football players in the National Football League (NFL).
11. Ask students to recall that they created and worked with *linear models* earlier in the unit. We are continuing our work with models and will learn another method of modeling called **CART**, which stands for **Classification and Regression Trees**. This is an umbrella term to refer to the following types of decision trees:
 - a. Classification Trees: the leaves predict the values of a categorical variable
 - b. Regression Trees: the leaves predict a numerical value
12. CART Activity: to get a sense of how classification trees work, the students will see one in action. We are going to try to classify 15 professional athletes into either soccer or football players based on some of their characteristics.
Advance preparation required: The cards in LMRs B, C, and D listed above (and previewed below) need to be cut out prior to class time.

Name: _____ Date: _____
CART Activity Player Stats

Directions for teacher:
Create "player" cards by cutting out each player's statistics from the table.

| | | |
|--|---|--|
| Player 1
Name: Cade Cowell
Team: San Jose
Height (inches): 72
Weight (pounds): 172
Age: 20
League: USMNT | Player 2
Name: Jared Goff
Team: Detroit
Height (inches): 76
Weight (pounds): 222
Age: 28
League: NFL | Player 3
Name: Kirk Cousins
Team: Minnesota
Height (inches): 75
Weight (pounds): 202
Age: 34
League: NFL |
| Player 4
Name: Bryan Reynolds
Team: Westerlo
Height (inches): 73
Weight (pounds): 170
Age: 22
League: USMNT | Player 5
Name: Jalen Neal
Team: Los Angeles
Height (inches): 73
Weight (pounds): 163
Age: 20
League: USMNT | Player 6
Name: Brock Purdy
Team: San Francisco
Height (inches): 73
Weight (pounds): 212
Age: 23
League: NFL |
| Player 7
Name: Desmond Ridder
Team: Atlanta
Height (inches): 75
Weight (pounds): 211 | Player 8
Name: Baker Mayfield
Team: Los Angeles
Height (inches): 73
Weight (pounds): 215 | Player 9
Name: Geno Smith
Team: Seattle
Height (inches): 75
Weight (pounds): 221 |

LMR_U4_L16_B

Name: _____ Date: _____

CART Activity Round 1

Directions for teacher:
Create nodes by cutting out each question below.

| |
|--|
| 1 – Root Node
Is your team located in the United States?
YES: Go to your right.
NO: Go to your left. |
| 2 – Internal Node
Are you 33 years old or older?
YES: Go to your right.
NO: Go to your left. |
| 3 – Leaf Node
You play for the US Men's National Soccer Team (USMNT). |
| 4 – Leaf Node
You play for the National Football League (NFL). |

LMR_U4_L16_C

Name: _____ Date: _____

CART Activity Round 2

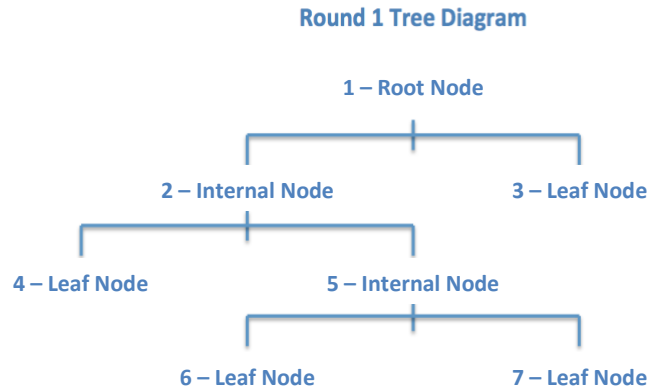
Directions for teacher:
Create nodes by cutting out each question below.

| |
|--|
| 1 – Root Node
Are you 74 inches tall or taller?
YES: Go to your right.
NO: Go to our left. |
| 2 – Leaf Node
You play for the National Football League (NFL). |
| 3 – Internal Node
Do you weigh more than 200 pounds?
YES: Go to your right.
NO: Go to your left. |
| 4 – Leaf Node
You play for the National Football League (NFL). |

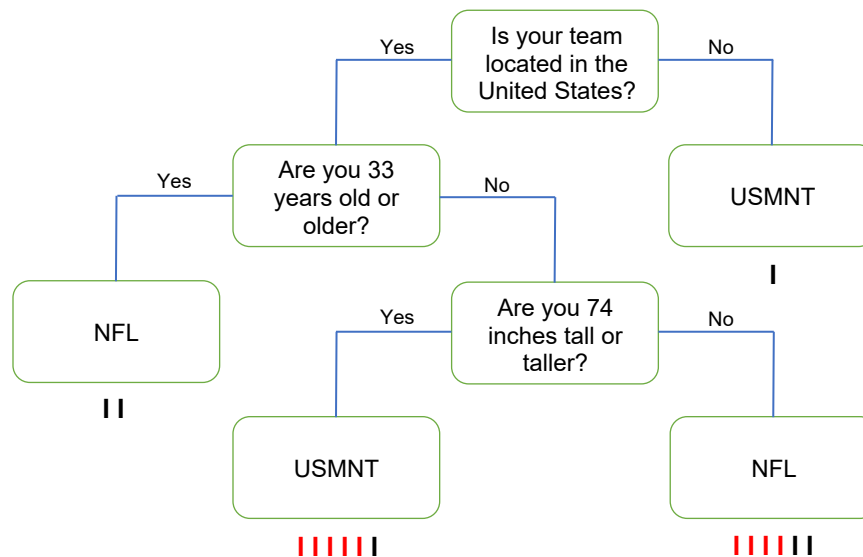
LMR_U4_L16_D

13. Ask for 15 volunteers and hand each of them a data card from the *CART Activity Player Stats* handout (LMR_U4_L16_B). These students will be known as the “players.” Each card lists the following variables for 15 different professional athletes.
 - a. team location
 - b. name
 - c. age
 - d. height (inches)
 - e. weight (pounds)
 - f. league
14. The “players” will only be allowed to say “yes” or “no” in this activity. No other talking is permitted.
15. Now, ask for 7 additional volunteers to be the **nodes** on the classification tree. There are three types of nodes in a decision tree:
 - a. Root node – the initial question/characteristic that splits the population/dataset
 - b. Internal nodes – a question/characteristic that splits the data
 - c. Leaf nodes – an “end” where no more splitting is possible
16. Distribute one question/classification from the *CART Activity Round 1 Questions* handout (LMR_U4_L16_C) to each node.

17. Arrange the 7 nodes in the room as depicted by the graphic below:



18. Now, each “player,” one at a time, will approach 1 – Root Node, who will ask the “player” the question listed on his/her card. Depending on the player’s answer, 1 – Root Node will direct the “player” to the next node.
19. The “player” continues through the nodes until a leaf declares the “player” to be either (1) a soccer player on the US Men’s National Team, OR (2) a football player in the National Football League (NFL).
20. Allow all the “players” to go through the nodes until each one is classified as either a soccer or football player.
21. After each player has been classified, record the classifications in the table below. The tally marks below correspond with classified incorrectly (red) and classified correctly (black), if all the player stats cards are used.

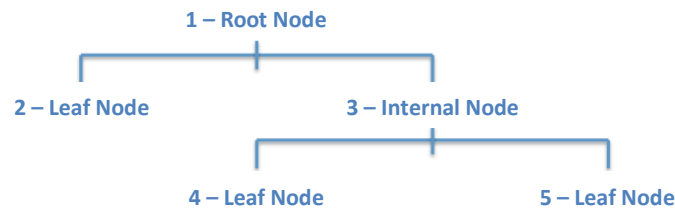


Note: This table will be referenced again later. A filled-out version is available in the next lesson.

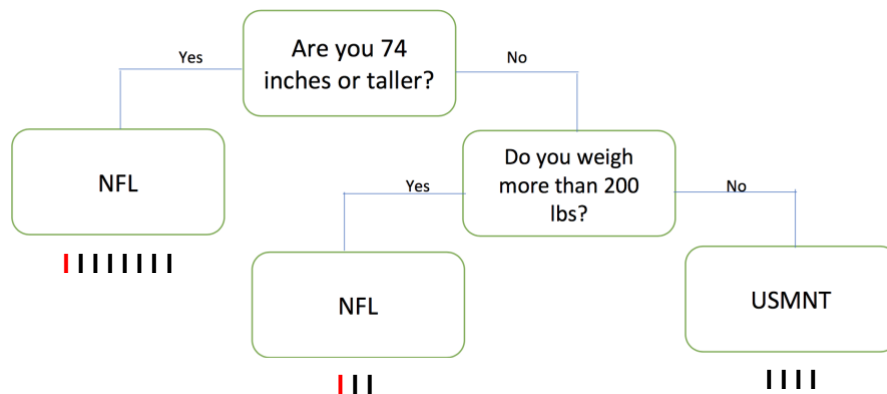
| Round 1 | Classified Correctly | Classified Incorrectly | Total Classified |
|---------------|----------------------|------------------------|------------------|
| 3 – Leaf Node | | | |
| 4 – Leaf Node | | | |
| 6 – Leaf Node | | | |
| 7 – Leaf Node | | | |

22. Ask students, “How successful were we in classifying players correctly?” *Answers will vary but if all the activity player stats cards were used then 6/15, or 40%, were classified correctly. Students might measure success by saying that we misclassified players 9/15, or 60%, of the time – this is called the misclassification rate (MCR) and will be defined in the next lesson.*
23. After proceeding through “Round 1,” ask an additional 5 students to come up as more nodes, distribute the cards from the CART Activity Round 2 Questions strips (LMR_U4_L16_D), and arrange the students like the diagram below:

Round 2 Tree Diagram



24. Have each “player” go through this new set of nodes until they are re-classified by these new rules.
25. After each player has been classified, record the classifications in the table below. The tally marks below correspond with classified incorrectly (red) and classified correctly (black), if all the player stats cards are used.



Note: This table will be referenced again later. A filled-out version is available in the next lesson.

| Round 2 | Classified Correctly | Classified Incorrectly | Total Classified |
|---------------|----------------------|------------------------|------------------|
| 2 - Leaf Node | | | |
| 4 - Leaf Node | | | |
| 5 - Leaf Node | | | |

26. Ask students, “How successful were we in classifying players correctly in this second round?” *Answers will vary but if all the activity player stats cards were used then 13/15, or 87%, were classified correctly. Students might measure success by saying that we only misclassified players 2/15, or 13%, of the time – this is called the misclassification rate (MCR) and will be defined in the next lesson.*



27. Once the activity has been completed, ask students the following questions:

- a. How do decision trees classify objects/people as being a member of a group? *Answer: By asking a series of questions, one at a time, and sending the participant down a particular path until he/she is classified.*
- b. Did we do as well, worse, or better in Round 2 compared to Round 1 at correctly guessing which sport the “players” participate in? Explain. *Answers will vary according to results of the activity. If all activity player stats cards were used then we did better in Round 2, with 87% success, than Round 1, with 40% success, in correctly guessing which sport the “players” participate in.*
- c. How can we figure out what questions to ask and in what order to minimize the number of incorrect classifications (also known as *misclassifications*)? *Answers will vary. This one might not be obvious but the point is for the students to wrestle with how they might think it can be done.*

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework



Students will find a decision tree online that aligns with their interests and:

- a. identify the nodes – root, internal, and leaf
- b. describe the population that it applies to
- c. describe what the decision tree is predicting, i.e., what are the classifications at the leaves?

Lesson 17: Grow Your Own Decision Tree

Objective:

Students will create their own decision trees based on training data (i.e., the data from the previous day's lessons), and then see how well their decision tree works on new test data.

Materials:

1. *Make Your Own Decision Tree* handout (LMR_U4_L17)
Optional: Provide code to students to create the training data in RStudio (see step 8 in lesson)

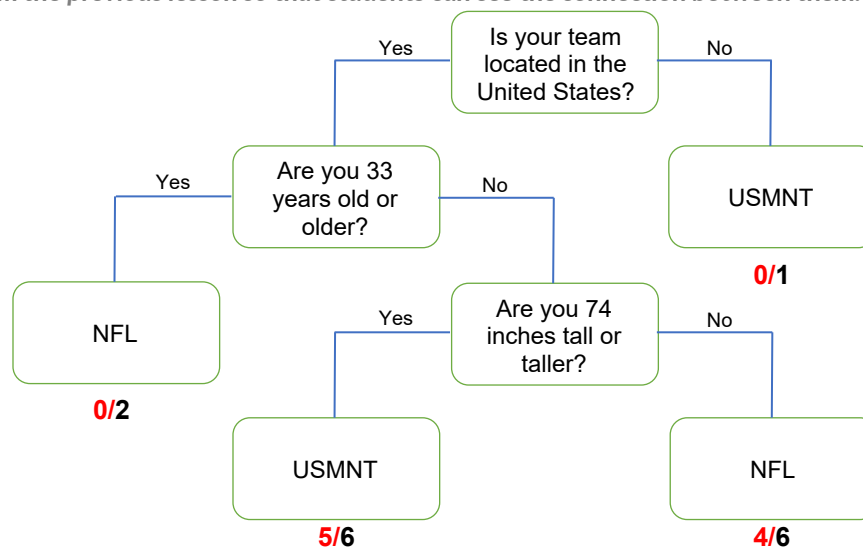
Vocabulary:

misclassification rate (MCR), training data, test data

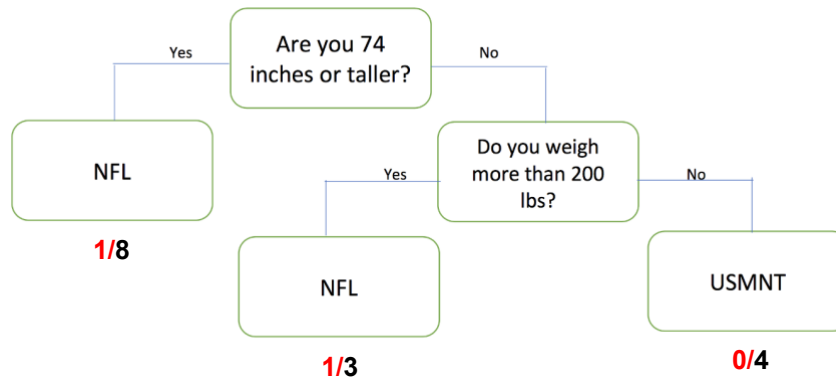
Essential Concepts: We can determine the usefulness of decision trees by comparing the number of misclassifications in each.

Lesson:

1. Begin the lesson by asking the following question: *How did we assess whether a linear model made good predictions for a set of data? Answers will vary but so far in this unit, we have used Mean Squared Error (MSE) and Mean Absolute Error (MAE).*
2. Tell students that much like linear models, classification trees also have a method for determining how well they make predictions for a set of data.
3. Classification trees use a **misclassification rate (MCR)**, which is the proportion of observations who were predicted to be in one category but were actually in another.
4. Refer back to your tallied decision trees and correct/incorrect classification tables from the previous lesson. What were the overall proportion of incorrect classifications for Round 1? What about for Round 2? *Answers will vary. If all activity player stats cards were used then the MCR for Round 1 is 9/15, or 0.60, and for Round 2 is 2/15, or 0.13. You can reference the decision trees and tables from the previous lesson so that students can see the connection between them.*



| Round 1 | Classified Correctly | Classified Incorrectly | Total Classified | |
|---------------|----------------------|------------------------|------------------|------------------|
| 3 - Leaf Node | 1 | 0 | 1 | |
| 4 - Leaf Node | 2 | 0 | 2 | |
| 6 - Leaf Node | 1 | 5 | 6 | |
| 7 - Leaf Node | 2 | 4 | 6 | |
| MCR | | 0 + 0 + 5 + 4 | 1 + 2 + 6 + 6 | = 9/15 (or 0.60) |



| Round 2 | Classified Correctly | Classified Incorrectly | Total Classified | |
|---------------|----------------------|------------------------|------------------|-------------------------|
| 2 - Leaf Node | 7 | 1 | 8 | |
| 4 - Leaf Node | 2 | 1 | 3 | |
| 5 - Leaf Node | 4 | 0 | 4 | |
| MCR | | 1 + 1 + 0 | 8 + 3 + 4 | = 2/15 (or 0.13) |

- These proportions of incorrect classifications are our misclassification rates.
- Today we'll be creating our own classification tree and assess its prediction accuracy.
- Display the following data (the same data from the player cards used in the previous lesson):

| Team | Player | Height (inches) | Weight (pounds) | Age | League |
|---------------|----------------|-----------------|-----------------|-----|--------|
| Atlanta | Desmond Ridder | 75 | 211 | 23 | NFL |
| Carolina | Sam Darnold | 75 | 225 | 25 | NFL |
| Cincinnati | Matt Miazga | 76 | 183 | 28 | USMNT |
| Cleveland | Deshaun Watson | 74 | 220 | 27 | NFL |
| Columbus | Aidan Morris | 69 | 159 | 22 | USMNT |
| Detroit | Jared Goff | 76 | 222 | 28 | NFL |
| Los Angeles | Jalen Neal | 73 | 163 | 20 | USMNT |
| Los Angeles | Baker Mayfield | 73 | 215 | 27 | NFL |
| Miami | Julian Gressel | 73 | 207 | 30 | USMNT |
| Minnesota | Kirk Cousins | 75 | 202 | 34 | NFL |
| New Orleans | Andy Dalton | 74 | 220 | 35 | NFL |
| San Francisco | Brock Purdy | 73 | 212 | 23 | NFL |
| San Jose | Cade Cowell | 72 | 172 | 20 | USMNT |
| Seattle | Geno Smith | 75 | 221 | 32 | NFL |
| Westerlo | Bryan Reynolds | 73 | 170 | 22 | USMNT |

- Distribute the *Make Your Own Decision Tree* handout (LMR_U4_L17) and give students time to come up with their own decision trees based on the **training data** they are given. Students may work in pairs or teams. They should follow the directions on page 1 of the handout and come up with a series of possible yes/no questions that they could ask to classify each player into his correct league (the NFL or the USMNT).

Optional (for LMR_U4_L17): Provide code below if students would like to recreate the training data in RStudio for manipulation.

Name: _____ Date: _____

Make Your Own Decision Tree

Directions:

1. Use the **training data** and blank decision tree provided below to create your own classification tree to help separate the NFL players from the USMNT players.
2. Start at the top of the decision tree and write **yes/no** questions in each of the **nodes/boxes**.
3. Continue to draw additional nodes/boxes to either write more questions or to classify a player.
4. Once you have completed your tree, sort the 6 players from the **test data** (on page 2) using your classifications and record them in the data table. Afterwards, your teacher will reveal who each player actually is, and you can determine how many you classified correctly.

Training Data

| Team | Player | Height (inches) | Weight (pounds) | Age | League |
|---------------|-----------------|-----------------|-----------------|-----|--------|
| Atlanta | Diamond Riddler | 75 | 211 | 23 | NFL |
| Carolina | Sam Darnold | 75 | 225 | 25 | NFL |
| Cincinnati | Matt Misoga | 76 | 183 | 28 | USMNT |
| Cleveland | Deshawn Watson | 74 | 220 | 27 | NFL |
| Columbus | Aidan Morris | 69 | 199 | 22 | USMNT |
| Detroit | Jared Goff | 76 | 222 | 28 | NFL |
| Los Angeles | Jalen Neal | 73 | 163 | 20 | USMNT |
| Los Angeles | Blair Mayfield | 73 | 215 | 27 | NFL |
| Miami | Julian Gressel | 73 | 207 | 30 | USMNT |
| Minnesota | Kirk Cousins | 75 | 202 | 34 | NFL |
| New Orleans | Andy Dalton | 74 | 220 | 35 | NFL |
| San Francisco | Brook Purney | 73 | 212 | 23 | NFL |
| San Jose | Cade Cowell | 72 | 172 | 20 | USMNT |
| Seattle | Geno Smith | 75 | 221 | 32 | NFL |
| Westerlo | Bryan Reynolds | 73 | 170 | 22 | USMNT |

Your Decision Tree

LMR_U4_L17

Optional code:

```
set.seed(42236789)
futbol_rows <- sample(1:nrow(futbol), 15)
futbol_training <- slice(futbol, futbol_training)
```

9. Once the students have finished creating their classification trees, ask the following questions:
 - a. Will you be able to classify other players from a new dataset correctly using this particular classification tree?
 - b. Do you think this classification tree is too specific to the training data?
10. Inform the students that they should now use the **test data** on page 2 of the handout to try to classify 5 *mystery players* into one of the two leagues. They should record the classification that their tree outputs in the data table on page 2 of the *Make Your Own Decision Tree* handout (LMR_U4_L17).
11. Let the students compare their classification trees and league assignments with one another. Hopefully, there will be a bit of variety in terms of the trees and the classifications.
12. Next, show students the correct league classifications for the 5 mystery players. The mystery player names are also included in this table.

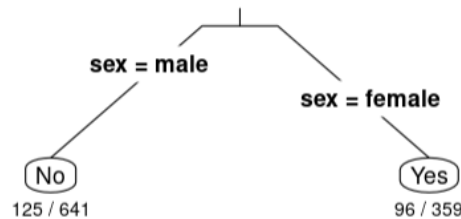
| Team | Player | Height (inches) | Weight (pounds) | Age | League |
|------------|-----------------|-----------------|-----------------|-----|--------|
| Baltimore | Lamar Jackson | 74 | 212 | 26 | NFL |
| Buffalo | Josh Allen | 77 | 237 | 26 | NFL |
| Cincinnati | Brandon Vázquez | 74 | 196 | 25 | USMNT |
| Eupen | Gabriel Slonina | 76 | 192 | 19 | USMNT |
| New York | John Tolkin | 67 | 132 | 21 | USMNT |

13. By a show of hands, ask:
 - a. How many students misclassified all of the players in the test data?
 - b. How many misclassified 4 of the 5 players?
 - c. How many misclassified 3 of the 5 players?
 - d. How many misclassified 2 of the 5 players?
 - e. Did anyone correctly classify ALL 5 mystery players? If so, ask those students to share their classification trees with the rest of the class.

14. Inform students that, when faced with much more data, creating classification trees becomes much harder to make by hand. It is so difficult, in fact, that data scientists rely on software to grow their trees for them. Students will learn how to create decision trees in RStudio in the next lab.

Note: Included below is a front-load of the calculation/interpretation of MCR in RStudio.

15. In the next lab, you will use RStudio to create tree models that will make good predictions without needing a lot of branches. RStudio can also calculate the misclassification rate. However, you might find the visual a little confusing to interpret, so we will preview and break down one of the first outputs you'll see in the lab.



16. Project the image above and explain to the students that the lab uses the titanic data, so we are looking at a total of 1000 observations (in this case, people). In Unit 2, we investigated whether those who survived paid a higher fare than those who died. This lab begins with predicting whether a person survived or not based on their sex – testing the idea that women were given preference on lifeboats. Much like the ratios we saw for our Round 1 and Round 2 classification trees, RStudio gives us ratios for each of the leaf nodes where a classification was made. The denominator tells us how many observations ended up in that node and the numerator tells us the number of misclassifications. Also notice that the questions are not shown in the classification tree, but the characteristics are visible instead. Ask students:

- What does the output 125/641 represent? *Answer: The 641 tells us that six-hundred-forty-one people were classified as not surviving based solely on the fact that they were male. The 125 represents people who were misclassified (they actually survived), which means that 516 people were classified correctly (they indeed did not survive).*
- How would you calculate the misclassification rate (MCR)? *Answer: You would add all of the numerators which represent the misclassifications and divide by the total number of observations which you could obtain by adding all the denominators. $(125+96)/(641+359)=221/1000$ or 0.221.*

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next Day



Students will record their responses to the following discussion questions:

- How is a decision tree/CART similar to or different than a linear model? *Answers will vary but a similarity is that they both make predictions, and a difference is that decision trees can make predictions for both numerical and categorical data while linear models only make predictions for numerical data.*
- Why is a decision tree considered a model? *Possible answer: We consider a decision tree a model because it still represents relationships between variables.*
- Describe the role that training data and test data play in creating a classification tree. *Possible answer: We use training data to build a classification tree and then use test data to see how well our tree makes predictions – we shouldn't use the entirety of our data because then the model will be too specific/overfitted and will not make good predictions when presented with new data.*

Lab 4G: Growing trees

Complete Lab 4G prior to Lesson 18.

Lab 4G: Growing trees

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Trees vs. Lines

- So far in the labs, we've learned how we can fit linear models to our data and use them to make predictions.
- In this lab, we'll learn how to make predictions by growing trees.
 - Instead of creating a line, we split our data into branches based on a series of *yes* or *no* questions.
 - The branches help sort our data into *leaves* which can then be used to make predictions.
- **Start by loading the titanic data.**

Our first tree

- **(1) Write and run code using the `tree()` function to create a *classification* tree that predicts whether a person *survived* the Titanic based on their *sex*.**
 - A *classification* tree tries to predict which category a categorical variable would belong to based on other variables.
 - The syntax for `tree` is similar to that of the `lm()` function.
 - **Assign this model the name `tree1`.**
- **(2) Why can't we just use a *linear model* to predict whether a passenger on the Titanic *survived* or not based on their *sex*?**

Viewing trees

- **(3) To actually look at and interpret our `tree1`, place the model into the `treepplot` function.**
 - **(4) Write down the labels of the two *branches*.**
 - **(5) Write down the labels of the two *leaves*.**
- Answer the following, based on the `treepplot`:
 - **(6) Which *sex* does the model predict will survive?**
 - **(7) Where does the plot tell you the number of people that get sorted into each leaf? How do you know?**
 - **(8) Where does the plot tell you the number of people that have been sorted *incorrectly* in each leaf?**

Leafier trees

- **(9) Similar to how you included multiple variables for a linear model, write and run code creating a *tree* that predicts whether a person *survived* based on their *sex*, *age*, *class*, and where they embarked.**
 - **Assign this model the name `tree2`.**
- **(10) Write and run code creating a `treepplot` for this model and answer the following questions:**
 - **(11) Mrs. Baxter was a 50-year-old female with a 1st class ticket from Cherbourg. Does the model predict that she survived?**
 - **(12) Which variable ended up not being used by `tree2`?**

Tree complexity

- By default, the `tree()` function will fit a *tree model* that will make good predictions without needing lots of branches.
- We can increase the complexity of our trees by changing the complexity parameter, `cp`, which equals `0.01` by default.
- We can also change the minimum number of observations needed in a leaf before we split it into a new branch using `minsplit`, which equals `20` by default.
- **(13) Using the same variables that you used in `tree2`, write and run code creating a model named `tree3` but include `cp = 0.005` and `minsplit = 10` as arguments.**
 - **(14) How is `tree3` different from `tree2`?**

Predictions and Cross-validation

- Just like with *linear models*, we can use cross-validation to measure how well our *classification trees* perform on unseen data.
- First, we need to compute the predictions that our model makes on test data.
 - Use the data function to load the `titanic_test` data.
 - **(15) Fill in the blanks below to predict whether people in the `titanic_test` data survived or not using `tree1`.**
 - Note: the argument `type = "class"` tells the `predict` function that we are predicting a categorical variable and not a numerical variable.

```
titanic_test <- mutate(____, prediction = predict(____, newdata = ____, type = "class"))
```

Measuring model performance

- Similar to how we use the *mean squared error* to describe how well our model predicts numerical variables, we use the *misclassification rate* to describe how well our model predicts categorical variables.
 - The *misclassification rate* (MCR) is the number of people who were predicted to be in one category but were actually in another.
- Run the following command to see a side-by-side comparison of the actual outcome and the predicted outcome:

```
View(select(titanic_test, survived, prediction))
```

- **(16) Where does the first misclassification occur?**

Misclassification rate

- In order to tally up the total number of misclassifications, we need to create a function that compares the actual outcome with the predicted outcome. The **not equal to** operator (`!=`) will be useful here.
- **(17) Fill in the blanks to create a function to calculate the MCR.**
 - Hint: `sum(____ != ____)` will count the number of times that the left-hand side does not equal the right-hand side.
 - We want to count the number of times that actual does not equal predicted and then divide by the total number of observations.

```
calc_mcr <- function(actual, predicted) {  
  sum(____ != ____) / length(actual)  
}
```

- Then run the following to calculate the MCR.

```
summarize(titanic_test, mcr = calc_mcr(survived, prediction))
```

On your own

- (18) In your own words, explain what the *misclassification rate* is.
- (19) Which model (`tree1`, `tree2`, or `tree3`) had the lowest misclassification rate for the `titanic_test` data?
- (20) Write and run code creating a 4th model using the same variables used in `tree2`. This time though, change the *complexity parameter* to `0.0001`. Then answer the following.
 - (21) Does creating a more complex *classification tree* always lead to better predictions? Why not?
- A *regression tree* is a tree model that predicts a numerical variable.
- (22) Write and run code creating a *regression tree* model to predict the Titanic's passengers' ages and calculate the MSE.
 - Plots of regression trees are often too complex to plot.

Ties that Bind

Instructional Days: 3

Enduring Understandings

Clustering is another way to classify data into groups. We classify observations based on numerical characteristics and their similarities. We use k-means to determine the mean value for each group of k clusters by randomly assigning an initial value for the mean and then moving the mean based on its proximity to the points.

Networks classify people into groupings based on who knows whom. Nodes are formed when a relationship between two people is present.

Engagement

Students will determine which points in a plot should be grouped as football players and which points should be grouped as swimmers based on clustering of characteristics.

Learning Objectives

Statistical/Mathematical:

S-IC 2: Decide if a specified model is consistent with results from a given data-generating process, e.g., using simulation.

Data Science:

Understand what RStudio is doing when using the k-means function to find clusters in a group of data and when creating networks in order to learn how to classify data into groups.

Applied Computational Thinking using RStudio:

- Use the k-means function to find clusters in a group of data.
- Plot the data with the cluster assignments based on the k-means function.

Real-World Connections:

Network analysis is used by many private and public entities such as the National Security Agency when they want to find terrorist networks to have maximum impact on communications. The k-means algorithm is a technique for grouping entities according to the similarity of their attributes. For example, dividing countries into similar groups using k-means to make fair comparisons is applicable.

Language Objectives

1. Students will write, in their own words, an explanation of k-means clustering.
2. Students will describe the differences between time spent on videogames and time spent on homework, from their own class data.
3. Students will create visualizations and numerical summaries to explain and justify, orally and in writing, a recommendation to better their community.

Data File or Data Collection Method

Data File:

1. USMNT and NFL: data(futbol)
2. Students' TimeUse campaign data

Data Collection Method:

Team-generated Participatory Sensing Campaign: Students will collect data for a team-selected topic.

Legend for Activity Icons



Video clip



Discussion



Articles/Reading



Assessments



Class Scribes

Lesson 18: Where Do I Belong?

Objective:

Students will learn what clustering is and how to classify groups of people into clusters based on unknown similarities.

Materials:

1. Find the Clusters handout (LMR_U4_L18)

Vocabulary:

clustering, cluster, k-means

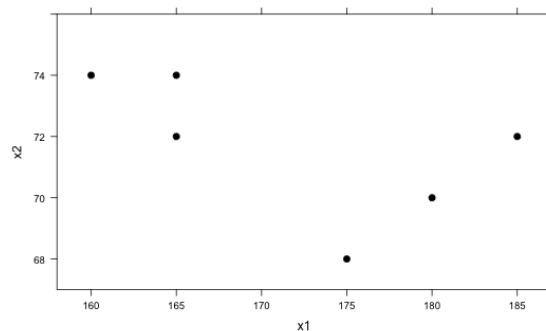
Essential Concepts: We can identify groups, or “clusters,” in data based on a few characteristics. For example, it is easy to classify a group of people into football players and swimmers, but what if you only knew each person's arm span? How well could you classify them into football players and swimmers now?

Lesson:

1. Inform the students that they will continue to explore different types of models, and today they will be focusing on **clustering**. Clustering is the process of grouping a set of objects (or people) together in such a way that people in the same group (called a **cluster**) are more similar to each other than to those in other groups.
2. Have the students recall that, in the previous lessons, they used decision trees/CART to classify people into different groups based on whether or not a person had a specific characteristic (e.g., whether or not a professional athlete's team is based in the US).
3. But sometimes we don't know what these specific characteristics are. We are simply given numerical variables and asked to find similarities. This is where clustering comes in – similar people will congregate towards each other, and we want to see if we can identify their groupings.
4. We will look at a very basic example first. Suppose the following 6 observations are given:

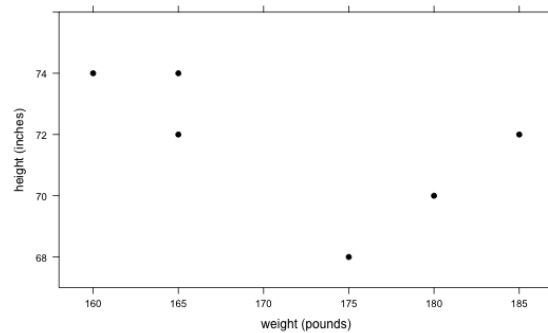
| Obs | X_1 | X_2 |
|-----|-------|-------|
| 1 | 160 | 74 |
| 2 | 165 | 72 |
| 3 | 165 | 74 |
| 4 | 175 | 68 |
| 5 | 180 | 70 |
| 6 | 185 | 72 |

5. Plot the X_1 and X_2 points on a scatterplot either on the board or on poster paper (X_1 can be on the horizontal axis and X_2 can be on the vertical axis). The graph should look like the one below:

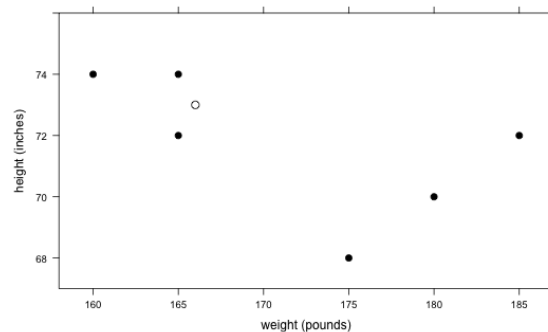


6. Ask students if they think there are any clusters, or groups, that stand out to them. It is likely that they will say there are 2 clusters in the graph: the top left corner 3 points, and the bottom right 3 points.

7. Now pose the following scenario that further describes the data:
 - a. A doctor provides yearly physicals to the football and swimming teams at a local high school.
 - b. The doctor has collected data over the past few years on each player's weight (in pounds) and height (in inches). She informs us that weight was coded as the variable X_1 , and height was coded as the variable X_2 . You can re-label the scatterplot with this new information.



- c. Unfortunately, the doctor never recorded what sport each person played.
8. Using the information about height and weight, ask the students to decide:
 - a. Which group of points most likely represents players from the swimming team? *Answer: The points in the upper left corner are probably swimmers because swimmers are usually tall (and have large arm spans) and thin.*
 - b. Which group of points most likely represents players from the football team? *Answer: The points in the bottom right corner are probably football players because they tend to be heavier and more muscular.*
9. Now suppose a new player comes into the doctor's office for a physical. His weight and height are recorded as 166 pounds and 73 inches, respectively, but the doctor forgets to ask what sport he plays. Plot this point on the graph and ask students to determine which sport they think this student plays. *Answer: This student is most likely a swimmer because he is tall and thin, and his point is near the swimming cluster.*



10. That was an easy one! But what if a player comes in and has the following measurements:
weight = 173 pounds, height = 73 inches?
11. Distribute the *Find the Clusters* handout (LMR_U4_L18) and tell the students that the new point has been added to the "Round 0" graph.

Name: _____ Date: _____

Find the Clusters

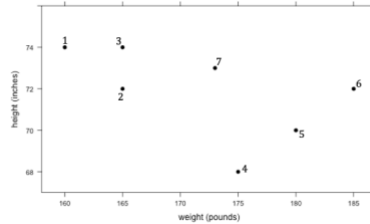
Round 0 – Initialization

Pick random points for your initial cluster centers.

Cluster A center: (_____, _____)

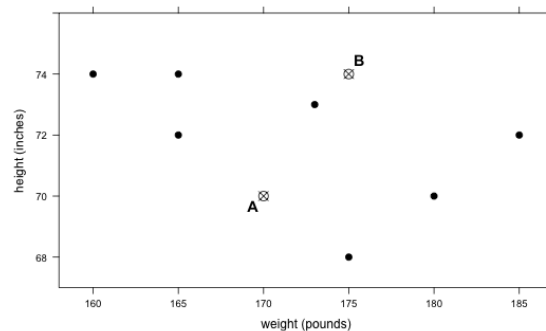
Cluster B center: (_____, _____)

Plot your two points for A and B on the graph below.



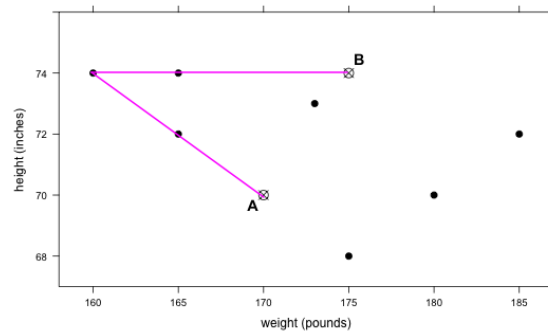
LMR_U4_L18

12. Ask students:
 - a. On which team do you think this person plays? *Answer: It is much more difficult to tell now because it looks like it is right in between the two clusters.*
13. In order to determine group placement, we can use a process called **k-means clustering**. With this method, we select k clusters that we want to identify. Since we know we only have 2 types of athletes, football players and swimmers, we will be finding $k = 2$ clusters.
14. To introduce the students to this idea, circle the 3 points in the upper left corner (the ones that are likely the swimmers) and have students find the “mean point.” This means that they should find the mean x-value and the mean y-value of the 3 points. They can then plot this new point and use it as the mean of this particular group, or cluster.
15. The goal of this algorithm is to keep recalculating means as the clusters change. To begin, we will pick 2 points, A and B, to represent the center of each cluster. We will be referring to this as the initialization step. If you would like to use the point found in step 14 above and label it as “A,” that is completely fine. You can simply pick just one other random point near the other cluster and label it as “B.”
16. **Initialization:** For now, we will start with the following two points as our initial cluster centers of each group: A: (170, 70) and B: (175, 74). In the “Round 0” plot on the *Find the Clusters* handout, each student should plot and label these two points.

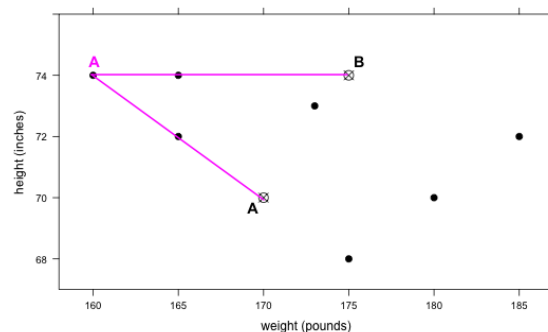


17. **Assignment Step:** Inform the students that they will be determining the distance between the 7 observations and point A and point B. Then, they will decide if the point is closer to cluster center A or cluster center B and label the point with that letter.

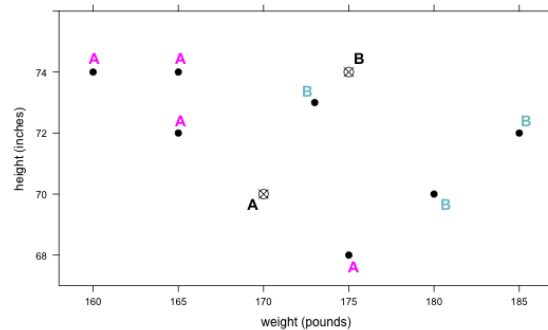
- a. Lines have been drawn from the top left point to the cluster centers in the plot below as a guide. You can draw this on the board as a reference for the students as well.



- b. Since the line to point A is smaller, we would classify that point as being in cluster A (as shown below).



- c. The students may draw similar lines for every point on the graph, or they can simply eyeball it, to make a decision as to which cluster each point belongs in. Even if they guess incorrectly, the algorithm should be able to find the correct groups after some time. The correct classifications for Round 0 are as follows, using our points from step 16 above:



18. **Update Step:** Once the class has agreed on the Round 0's cluster classifications, they should compute new values for points A and B by using the clustered point means. For point A, they simply need to find the mean x-value for the 4 points and the mean y-value for the 4 points. Repeat this process to find the new point B – these will be our new cluster centers as we move onto Round 1. Calculating means to derive cluster centers, like points A and B, for grouping data points is part of the **k-means** algorithm.

- a. The new points for A and B have been calculated below. The students should be calculating these on their own and recording their new cluster centers on the handout.

$$\text{x-value for A} = (160 + 165 + 165 + 175)/4 = 166.25$$

$$\text{y-value for A} = (74 + 72 + 74 + 68)/4 = 72$$

$$\text{x-value for B} = (173 + 180 + 185)/3 = 179.3$$

$$\text{y-value for B} = (73 + 70 + 72)/3 = 71.67$$

new A = (166.25, 72)
new B = (179.3, 71.67)

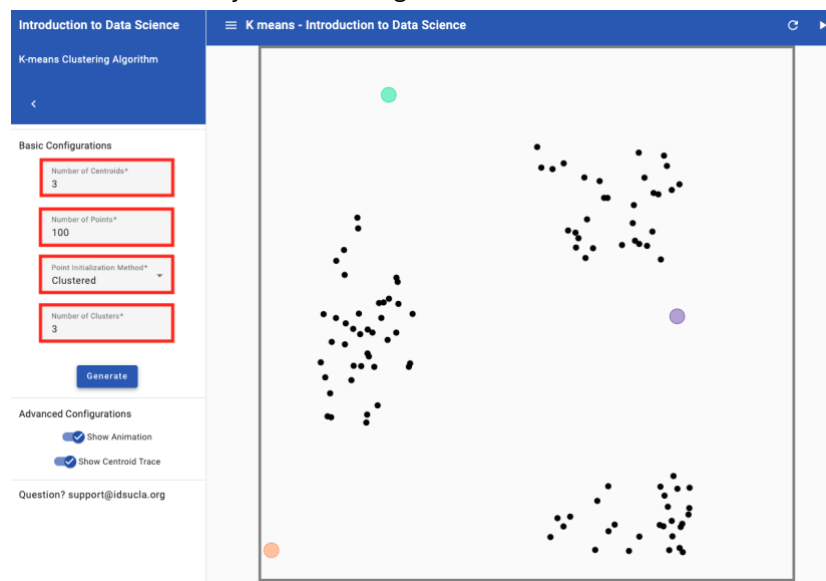
19. Have the students continue working through the handout until the cluster membership remains the same between 2 consecutive rounds. This means that, from one iteration to the next, the points in each cluster do not change.

Note: If students used the initial cluster centers from step 16 above, they would finish clustering in Round 2 with points 1, 2, and 3 belonging to cluster A and points 4, 5, 6, and 7 belonging to cluster B.

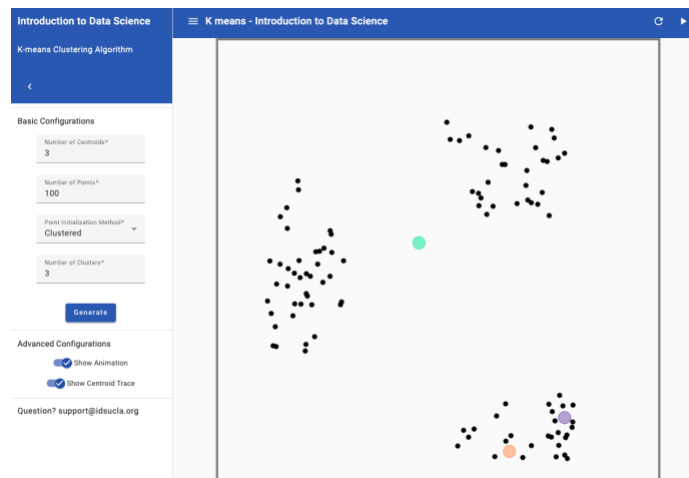
20. Where you choose your initial points matters in determining which points end up in which clusters. Demonstrate this to the class using the K-means Clustering App, located on the Applications page on Portal under Explore

(<https://portal.thinkdataed.org/#curriculum/applications/>).

- a. In the app, “centroids” is the academic term for cluster centers. For this example, we will use 3 centroids and choose a “Clustered Initialization” with 100 points and 3 clusters. See image below for how to adjust the settings.



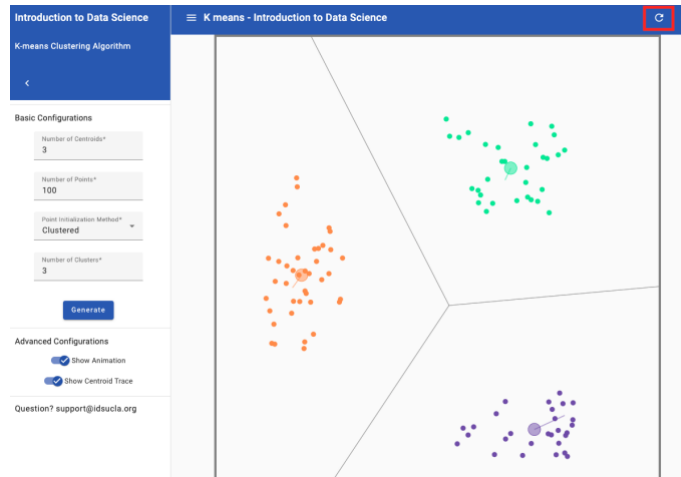
- b. Move two of the centroids so that they are close to/within one cluster. Move the remaining centroid in between the other two clusters. An example is shown below.



- c. Click “Next:” until no points change from the previous update. See image below of end of clustering for the example from (b).



- d. While we still have three clusters, we can clearly see that our initial points have influenced the end result of those clusters. You can click “Restart and play again” to move the centroids to the center of the clusters and click “Next:” until no points change from the previous update to illustrate this point (see image below of correctly clustered groups).



Note: Clicking “Restart and play again” allows you to move your cluster centers to new positions while still using the same 100 points. You can repeat the same process as above to create different additional groupings.

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Homework & Next Day



Write a paragraph that describes k-means clustering in your own words.

Lab 4H: Finding clusters

Complete Lab 4H prior to Lesson 19.

Lab 4H: Finding clusters

Directions: Follow along with the slides, completing the questions in **blue** on your computer, and answering the questions in **red** in your journal.

Clustering data

- We've seen previously that data scientists have methods to predict values of specific variables.
 - We used *regression* to predict numerical values and *classification* to predict categories.
- *Clustering* is similar to classification in that we want to group people into categories. But there's one important difference:
 - In *clustering*, we don't know how many groups to use because we're not predicting the value of a known variable!
- In this lab, we'll learn how to use the k-means clustering algorithm to group our data into clusters.

The k-means algorithm

- The k-means algorithm works by splitting our data into k different clusters.
 - The number of clusters, the value of k , is chosen by the data scientist.
- The algorithm works *only* for numerical variables and *only* when we have no missing data.
- **To start, use the data function to load the futbol dataset.**
 - This data contains 24 players from the US Men's National Soccer team (USMNT) and 33 quarterbacks from the National Football League (NFL).
- **(1) Write and run code creating a scatterplot of the players' `ht_inches` and `wt_lbs` and color each dot based on the `league` they play for.**

Running k-means

- After plotting the player's heights and weights, we can see that there are two clusters, or different types, of players:
 - Players in the NFL tend to be taller and weigh more than the shorter and lighter USMNT players.
- **(2) Fill in the blanks below to use k-means to cluster the same height and weight data into two groups:**

```
kclusters(____~____, data = futbol, k = ____)
```

- **(3) Use this code and the `mutate` function to add the values from `kclusters` to the `futbol` data. Name the variable `clusters`.**

k-means vs. ground-truth

- In comparing our football and soccer players, we *know* for certain which league each player plays in. We call this knowledge **ground-truth**.
- Knowing the *ground-truth* for this example is helpful to illustrate how k-means works, but in reality, data scientists would run k-means not knowing the *ground-truth*.
- **(4) Compare the clusters chosen by k-means to the *ground-truth*. How successful was k-means at recovering the `league` information?**

On your own

- Load your class's `timeuse` data (remember to run `timeuse_format` so each row represents the mean time each student spent participating in the various activities).
- (5) Write and run code creating a scatterplot of `homework` and `videogames` variables.
- (6) Based on this graph, identify and remove any outliers by using the `filter` function.
- (7) Write and run code using `kclusters` with `k=2` for `homework` and `videogames`.
- (8) Describe how the groups differ from each other in terms of how long each group spends playing `videogames` and doing `homework`.

Lesson 19: Our Class Network

Objective:

Students will participate in an activity to map out their own network based on acquaintances between two people.

Materials:

1. *Friend Network Graphic* handout (LMR_U4_L19_A)
2. Index cards
3. *Network Code* file (LMR_U4_L19_B)

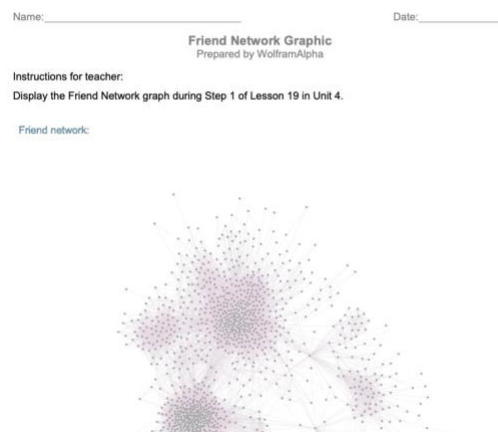
Vocabulary:

network

Essential Concepts: Networks are made when observations are interconnected. In a social setting, we can examine how different people are connected by finding relationships between other people in a network.

Lesson:

1. Display the *Friend Network Graphic* handout (LMR_U4_L19_A), which shows a WolframAlpha visualization of someone's Facebook friends. Inform the students that this type of model is called a **network**, which is simply a group of people or things that are interconnected in some way.

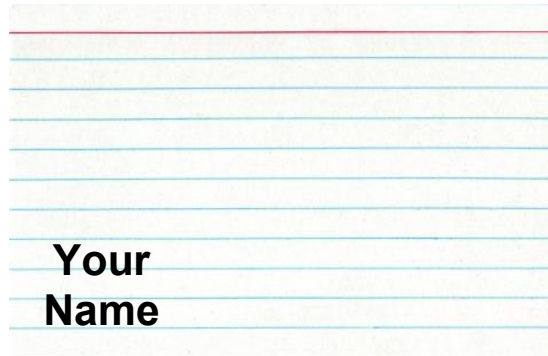


LMR_U4_L19_A



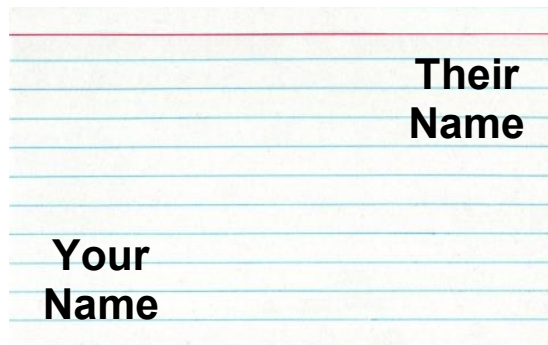
2. Ask the following questions about the graphic:
 - a. What does each dot represent? *Answer: Each dot represents one person.*
 - b. What does each line represent? *Answer: Each line represents a friendship between two people.*
 - c. How are all the people in this graphic connected to each other? *Answer: They are all friends with the person whose Facebook this is.*
 - d. Why are some areas denser than others? *Answer: A lot of people in the darker spots know each other, so there are more connections/friendships.*
 - e. Why are some people not in groups at all (the dots at the edges of the graphic)? *Answer: The main person does not have any friends in common with this person.*
 - f. What might some of the groupings (the denser spots) represent? *Answers will vary. Some examples include high school friends, college friends, graduate school friends, family members, or people who participate in similar hobbies.*
3. Ask the students what other types of social networks, other than Facebook, they belong to? Responses will most likely include TikTok, Twitter, Instagram, Snapchat, LinkedIn, Google+, etc.

4. Next, inform the students that networks can be as big or as small as we want. We can even determine our own class's social network and create visualizations from it!
5. Network Activity:
 - a. Distribute index cards to students. Each student will need enough cards to make a connection with every other person in the class. For example, if there are 20 students in a class, then each student needs 19 cards.
 - b. On EVERY index card, the student should write his/her first AND last name in the lower left-hand corner (see image below).



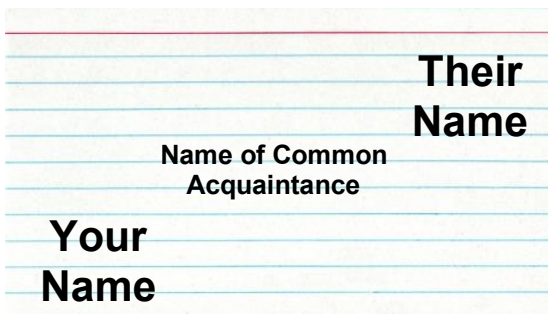
An index card with horizontal blue lines. The text "Your Name" is written in bold black font in the bottom-left corner.

- c. Next, each student will walk around the classroom and put another student's first AND last name in the top right-hand corner of an index card (see image below).



An index card with horizontal blue lines. The text "Your Name" is written in bold black font in the bottom-left corner, and "Their Name" is written in bold black font in the top-right corner.

- d. In the center of the index card, the students should write the name of the *closest 3rd* person that they BOTH know (see image below). The person can be someone in the class, someone outside of the class, or someone who doesn't even attend the same school.



An index card with horizontal blue lines. The text "Your Name" is written in bold black font in the bottom-left corner, "Their Name" is written in bold black font in the top-right corner, and "Name of Common Acquaintance" is written in bold black font in the center.

- e. Once all the students have completed their cards, they will turn them in to the teacher so the teacher can create a visualization of the network.

Note: This will probably take an entire class period to complete, which is fine because the graphics can be created and shown the next day.
6. At this point, the teacher will need to manually input the data from the index cards into a spreadsheet. ***It is recommended that the spreadsheet be saved as a .csv file.*** Two sample index cards are included on the next page, along with how you would input the data.



Note: The first index card corresponds to rows 1 and 2 in the spreadsheet (the purple box). The second index card corresponds to rows 3 and 4 in the spreadsheet (the pink box). So, each card will take up two rows in the spreadsheet.

Note: It is probably best to input the data after class and present the visualization during the next day.

7. Once all data has been input into a spreadsheet, use the code provided in the *Network Code* file (LMR_U4_L19_B) to produce graphs for the class's social network.

Note: The R Script file can be opened and viewed in the "source" pane of RStudio after it has been uploaded as a file into your RStudio project. There are 2 places where the code needs to be edited by the teacher:

- a. Be sure to change the file name when reading in the .csv file in Line 7 of the code.
- b. Read the comments in Lines 91-96 to help find the 5 most popular people in the class's network. This may require some edits to Lines 97 and 108.

```

1 ~ #####
2 # Load and clean the data
3 ~ #####
4
5 # Spreadsheet needs to be a .csv file for this code to work
6 # Be sure to replace "name_of_file" with your actual file name
7 connect <- read.csv("name_of_file.csv", head=FALSE, stringsAsFactors=FALSE)
8
9 # Assign variable names to columns 1 and 2 in the data set
10 names(connect) <- c("person1", "person2")
11
12 # Create the connections between people
13 connect$person1 <- tolower(connect$person1)
14 connect$person2 <- tolower(connect$person2)
15 connect$person1 <- gsub(connect$person1, pattern = "-", replacement = " ")
16 connect$person2 <- gsub(connect$person2, pattern = "-", replacement = " ")
17
18 # Find all unique persons in the data set
19 uni_connect <- c(unique(connect$person1), unique(connect$person2)))
20

```

LMR_U4_L19_B

Class Scribes:



One team of students will give a brief talk to discuss what they think the 3 most important topics of the day were.

Next Day

Students will end their Team Participatory Sensing campaign data collection after today's lesson. Starting the next day, they will analyze their data as part of the End of Unit 4 Project.

End of Unit 4 Project: Modeling a Community Issue

Objective:

Students will apply their learning of the third and fourth units of the curriculum by completing an end of unit modeling activity project.

Materials:

1. Computers
2. *End of Unit 4 Project: Modeling a Community Issue* (LMR_U4_End_of_Unit_Project)

End of Unit 4 Project Modeling a Community Issue

At the beginning of this unit, you explored data from the Trash Participatory Sensing campaign as well as from drought.gov (U.S. drought). You also created a Participatory Sensing campaign to investigate an issue in your community.

For this assignment, you will use the results of your Participatory Sensing campaign, and possibly an official dataset, to apply what you have learned in Unit 4 and to answer the research question you chose with your group at the beginning of the unit.

Your assignment is as follows:

1. Research which government entity is responsible for the community issue that your team chose – it could be local, state, federal, or even an international entity.
2. Propose one or two recommendations to the government entity in step 1 from above to help raise public awareness and/or alleviate the issue. Specifically, you will create a presentation in which you answer the following questions:
 - a. What is/are the specific recommendation(s) you are proposing to increase public awareness and/or alleviate the issue?
 - b. Why do you think this will work? What evidence do you have to support this? Include any necessary plots and analysis.
3. Your 5-minute presentation comprising 4-5 slides should include:
 - a. An introduction – Who are you? Introduce yourself and your team.
 - b. The issue – What is the issue? Why is this issue important? Why should we care about it?
 - c. The Participatory Sensing campaign – Explain your campaign. What was your research question? What statistical question(s) were you hoping to answer?
 - d. Your recommendation – What is your recommendation? How will it raise public awareness and/or alleviate the issue? Why do you think this will work? This is where you include your evidence:
 - i. How does your article connect to your recommendation?
 - ii. How does your Participatory Sensing campaign data support your recommendation? Or how does it show that there is a lack of something needed?
 - iii. If you have an official dataset, how does it support your recommendation?
 - iv. Include visualizations, numerical summaries, and/or statistics.
 - e. A closing – Summarize your point into a few closing sentences.

Each person must participate in the presentation. In addition to the presentation, submit a 2–4-page double-spaced summary of your analysis including plots/graphs.