

Course Title:	Content Area:	Grade Level:	Credit (if applicable)
AP Statistics	Mathematics	Gr. 11 & 12	1.0
Course Description:			
The AP Statistics course introduces students to the major concepts and tools for formulating questions, collecting and analyzing data, and interpreting results from data. The practices and skills for the course have been aligned to help students understand the statistical problem-solving process based on the American Statistical Association (ASA) recommendations. The content, skills, and assessments in the AP Statistics course focus on exploring data, sampling and experimentation, probability and simulation, and statistical inference. Students use technology, investigations, problem solving, and writing as they build conceptual understanding.			
Aligned Core Resources:		Connection to the BPS Vision of the Graduate	
Starnes, Darren and Tabor, Josh (2026). <i>The Practice of Statistics</i> (8th ed.). BFW Publishers AP Classroom Plan - Teachers may consider the following approaches as they plan their instruction before teaching each unit. - Review the overview at the start of each unit guide to identify essential questions, conceptual understandings, and skills for each unit. - Use the Unit at a Glance table to identify related topics that build toward a common understanding, and then plan appropriate pacing for students. Teach - Use the topic pages in the unit guides to identify the required content. - Integrate the content with a skill, considering any appropriate scaffolding. - Use the available resources, including AP Videos, on the topic pages to bring a variety of assets into the classroom. Assess - While teaching each topic, use AP Classroom to assign students topic questions as a way to continuously check student understanding and provide just-in-time feedback. - At the end of each unit, use AP Classroom to assign students progress checks, as homework or an in-class task. - Provide question-level feedback to students through answer rationales; provide formative feedback using Reports. - Create additional practice opportunities using the Question Bank and assign them through AP Classroom AP Statistics Course Exam & Description (CED)		Critical Thinking and Problem Solving & Content Mastery - Formulating Questions & Problem Solving - Quantitative Analysis - Evidence-Based Justification Information, Media, and Technology Literacy - Ethical Data & Media Evaluation - Technology Integration - Critical Evaluation of Claims Effective Communication and Collaboration - Multi-Modal Communication: Students present statistical claims and distributions through written explanations, graphical representations, and oral presentation - Collaborative Problem Solving Civic Awareness & Interdisciplinary Connections - Interdisciplinary Applications: The curriculum connects statistical reasoning to key societal fields	
Additional Course Information:		Link to Completed Equity Audit	
Knowledge/Skill Dependent courses/prerequisites Prerequisites: Grade 11 students may take AP Statistics if they are concurrently enrolled in AP Pre-Calculus, and had a grade of 83 or better in Algebra 2 ACC and were recommended by their teacher. Grade 12 students may take AP Statistics if they had a		AP STATISTICS Equity Curriculum Review Audit (2026)	

grade of 83 or better in Accelerated Algebra 2 and were recommended by their teacher.					
Standard Matrix					
Statistical Practices and Skills	Unit 1	Unit 2	Unit 3	Unit 4	Unit 5
Statistical Practice 1: Formulate Questions					
Determine an investigative question for a statistical study					
1.A Determine a valid investigative question that requires a statistical investigation.	X				
Statistical Practice 2: Collect Data					
Identify and justify methods for collecting data and conducting statistical inference					
2.A Identify information to answer a question or solve a problem.	X	X	X	X	
2.B Justify an appropriate method for ethically gathering and representing data.	X				
2.C Identify appropriate statistical inference methods.			X	X	
2.D Identify types of errors and relationships among components in statistical inference methods.			X	X	
2.E Identify the null and alternative hypotheses.			X	X	
Statistical Practice 3: Analyze Data					
Construct representations of data and calculate numerical statistical outputs.					
3.A Construct tabular and graphical representations of data and distributions.	X	X			X
3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	X	X			X
3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.		X	X	X	
3.D Calculate means, standard deviations, and parameters for probability distributions.		X	X	X	
3.E Calculate appropriate statistical inference method results.			X	X	
Statistical Practice 4: Interpret Results					
Interpret results and justify conclusions and methods					
4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	X	X			X
4.B Justify a claim based on statistical calculations and results.	X	X			X
4.C Describe distributions and compare relative positions of points within a distribution.	X	X			
4.D Interpret statistical calculations and results to assess meaning or a claim.	X	X	X	X	X
4.E Justify the use of a chosen statistical inference method by verifying conditions.			X	X	
4.F Interpret results of statistical inference methods.			X	X	
4.G Justify a claim based on statistical inference method results.			X	X	

Unit Links

[AP Statistics](#)

[Unit 1: Exploring One-Variable Data and Collecting Data](#)

[Unit 2: Probability, Random Variables, and Probability](#)

[Unit 3: Inference for Categorical Data: Proportions](#)

[Unit 4: Inference for Quantitative: Means](#)

[Unit 5: Regression Analysis](#)

Unit Title		
Unit 1: Exploring One-Variable Data and Collecting Data		
Relevant Standards: Bold indicates priority		
1.A, 2.A, 2.B, 3.A, 3.B, 4.A, 4.B, 4.C		
Essential Question(s)		Enduring Understanding(s)
- What kinds of questions can be answered using data from websites? - What kind of data is collected about me on social media? - What methods should I use to collect the data to conduct future inference procedures?		- Meaningful statistical investigations begin with well-defined questions and appropriate methods of data collection. - The quality of statistical conclusions depends on how data are collected and whether sources of bias are minimized. - Data displays and numerical summaries reveal important characteristics of distributions, including variability. - Statistical results must be interpreted within the context of the study and communicated ethically.
Demonstration of Learning		Pacing for Unit
Progress Check 1 - Multiple-choice: ~44 questions - Free-response: 3 questions - Question 1: Multi-Focus on Practices 1 and 2 - Question 2: Multi-Focus on Practices 3 and 4 - Question 3: Multi-Focus on Practices 3 and 4 The additional online question bank will be used to create a mid-unit quiz and other checks for understanding throughout the unit. College Board AP Exam Section I MCQ Weighting: 20–30% <i>(Largest unit weighting on the exam)</i> - Data Collection & Study Design: Distinguishing between observational studies, experiments, and surveys; identifying random sampling methods, sources of bias (nonresponse, undercoverage), and experimental design principles (control groups, randomization, blinding, confounding variables). - Data Representation & Summaries: Constructing and reading graphical displays (histograms, boxplots, dotplots, bar charts) and calculating summary statistics (mean, median, standard deviation, IQR, percentiles). - Distribution Analysis: Describing and comparing single quantitative distributions in context using shape, center, spread, and outliers. See AP Exam Overview for more information		15 Block Periods <i>NOTE:</i> <i>Midterm & AP Exam Review / Mock Exams - 11 Blocks</i> <i>Flexibility/Buffer/Post-AP Exam Project - 6 Blocks</i>
Family Overview (link below)		Integration of Technology:
AP Statistics FAMILY MATERIALS		- Use graphing calculators, Desmos, spreadsheets, or statistical software to create histograms, boxplots, dotplots, and bar charts. - Calculate descriptive statistics using technology. - Analyze authentic datasets using technology
Unit-specific Vocabulary		Aligned Unit Materials, Resources, and Technology (beyond core resources)
Investigative question	Spread	www.mathmedic.com www.desmos.com www.stapplet.com
Population	Variability	
Sample	Outlier	

Census Observational study Experiment Survey Random sampling Bias Sampling frame Variable Categorical variable Quantitative variable Distribution Center	Histogram Boxplot Dotplot Bar chart Mean Median Standard deviation Interquartile range (IQR) Percentile Skewness Association Confounding variable	https://skewthescript.org/	
Opportunities for Interdisciplinary Connections:		Anticipated misconceptions	
- Science: experimental design, measurement error, data collection. - Social Studies: polling, surveys, demographic studies. - Health: public health data and behavioral studies. - Business/Economics: market research and consumer surveys. - English: evaluating claims supported by data.		- Confusing random sampling with random assignment when determining the scope of inference (generalizability vs. causation). - Believing that increasing the sample size can eliminate or correct sources of sampling bias. - Treating bar charts (categorical data) and histograms (quantitative data) as interchangeable. - Assuming the mean and standard deviation are resistant measures against extreme outliers or skewness. - Omitting context or using non-comparative language when describing and comparing distributions.	
Connections to Prior Units		Connections to Future Units	
Students apply algebraic reasoning, graphical analysis, and data interpretation skills developed in previous mathematics courses.		- Provides the foundation for probability models, inference procedures, and regression analysis. - Establishes the language and habits of statistical thinking used throughout the course.	
Differentiation through <i>Universal Design for Learning</i>			
	Engagement	Representation	Action & Expression
1.A	Conduct a live class poll on a topic with inherent variability (e.g., student reaction times or commute lengths) vs. deterministic facts.	Flowchart contrasting deterministic questions ("How many desks are in the room?") with statistical questions expecting variability.	Draft three valid statistical questions based on a real-world dataset (e.g., consumer reports) and pitch one to a peer team.
2.A	"Data Detective" activity sorting survey questions into categorical vs. quantitative variables (discrete vs. continuous).	Color-coded data table highlighting observational units (rows) and variable types (columns) with visual icons.	Annotate an AP-style prompt to identify the cases, variables, and measurement units using digital annotation tools or highlighters.
2.B	Analyze misleading graphs from media headlines (e.g., truncated Y-axes, distorted pie chart slices) to debate ethical design choices.	Side-by-side visual matrix showing "Misleading vs. Ethical" graphical representations of the same univariate dataset.	Redesign a misleading news graphic into an accurate graph and record an audio reflection explaining the correction.
3.A	"Human Histogram" activity where students physically line up by height or shoe size before translating data to paper/software.	Step-by-step visual creation guides detailing axis scaling, bin width selection, and labeling for dotplots, stemplots, and boxplots.	Construct univariate displays using choice of medium: hand-drawn posters, graphing calculators, or online applets.
3.B	"Sports Analytics Challenge" calculating summary statistics and z-scores to compare athletes across different historical eras.	Formula scaffolding sheets contrasting formulas for mean/standard deviation with 5-number summary and z-score equations.	Create an interactive digital calculator guide or record a screencast walkthrough computing summary statistics and z-scores.
4.A	"Distribution Speed Dating" matching mystery graphs to contextual descriptions using explicit comparative language.	Visual anchor chart for the SOCS framework (Shape, Outliers, Center, Spread) emphasizing mandatory comparative conjunctions (<i>whereas</i> , <i>higher than</i>).	Write a comparative analysis essay or record an audio podcast evaluating two competing univariate distributions (e.g., test scores of two classes).

4.B	"Mythbusters: Stat Edition" testing claims about data asymmetry or potential outliers using strict mathematical boundaries.	Claim-Evidence-Reasoning (CER) graphic organizer tailored to the outlier boundaries.	Submit a written CER response or video defense verifying whether a specific data value qualifies as a statistical outlier.
4.C	"Where Do You Stand?" interactive line-up placing students along a scaled Normal curve based on assigned z-score/percentile cards.	Multi-layered Normal distribution diagram displaying z-scores, raw values, percentiles, and Empirical Rule percentages simultaneously.	Present an oral or written explanation of relative standing comparing two individuals in different distributions (e.g., SAT vs. ACT percentiles).
4.D	"Lost in Translation" game where students translate calculated summary metrics into plain-English client advisories.	Sentence-stem bank providing standard AP formulaic phrasing for interpreting standard deviation, z-scores, and percentiles in context.	Write an executive summary or record a video briefing interpreting standard deviation as a measure of typical variation for a business client.

Supporting Multilingual/English Learners (*CELP standards*)

	Emerging	Expanding	Bridging
1.A	Select valid statistical questions (those expecting variability) from visual cards; use sentence frames: "Is '[Question]' statistical? Yes, because [variable] changes."	Formulate statistical questions for one-variable data sets using stems: "What is the typical [measure] of [population]?"	Draft, evaluate, and refine complex investigative questions that clearly distinguish between population parameters and sample statistics.
2.A	Match data collection terms (<i>sample, population, census, variable</i>) to visual icons and simple scenario cards.	Identify population, sample, and variable type (categorical vs. quantitative) from scenario prompts using a graphic organizer.	Analyze study descriptions to identify sampling frames, potential sources of bias (undercoverage, nonresponse), and variable definitions.
2.B	Match sampling method cards (<i>SRS, Stratified, Cluster, Systematic</i>) to visual diagrams showing population selection.	Explain why a sampling method is unbiased or biased using paragraph frames ("This sampling method creates bias because [group] is left out.").	Evaluate sampling and experimental designs (blocking, blinding) and justify their ethical and statistical validity in formal written arguments.
3.A	Label components (axes, scales, titles) on pre-drawn graphs (dotplots, stemplots, histograms, boxplots) using a visual word bank.	Construct histograms, stemplots, and boxplots from data sets using structured step-by-step checklists and template grids.	Independently construct and display comparative graphical representations (side-by-side boxplots, back-to-back stemplots) using precise scaling.
3.B	Calculate mean, median, IQR, range, and z-scores using a formula reference card with annotated visual variables.	Calculate summary statistics and percentiles, recording work in structured computational templates with step-by-step starters.	Compute and interpret summary statistics, linear transformations, and relative positions (z-scores and percentiles) in contextual tasks.
4.A	Use the S.O.C.S. acronym chart (<i>Shape, Outliers, Center, Spread</i>) with picture prompts to list basic distribution features.	Describe distributions using S.O.C.S. and explicit comparative words (<i>higher than, skewed right, similar variability</i>) with sentence frames.	Write comprehensive, contextual comparative analyses of multi-group distributions incorporating precise statistical terminology.
4.B	Complete fill-in-the-blank claim statements comparing data points to summary statistics (e.g., "Point X is an outlier because it is greater than...").	Write structured justifications for claims about data distributions using provided "Claim-Evidence-Reasoning" (CER) writing frames.	Construct rigorous, contextualized arguments justifying claims about data distributions
4.C	Identify shape (symmetric, skewed) and point position (z-score sign/magnitude) using visual distribution overlays.	Describe an individual data point's location relative to the mean/median using comparative sentence frames ("A z-score of +2.1 means...").	Interpret percentiles and z-scores in non-standard distributions, comparing relative standings across different quantitative distributions.
4.D	Match computed summary statistics to their visual meanings on a distribution curve.	Interpret standard deviation and IQR in context using sentence templates ("The typical distance of [variable] from the mean is...").	Evaluate claims about real-world data by interpreting the contextual meaning of center, variability, and position measures.

Unit Outline

Lesson Sequence	Skills	Learning Objective	Essential Knowledge	AP Activities (if available)
Day 1 1.1 <i>Introducing</i>	1.A Determine a valid investigative	1.1.B Determine an investigative	1.1.A.1 A statistical study is a study in which data are collected from a sample to answer an investigative question about a larger population.	

<i>Statistics: What Can We Learn from Data?</i>	question that requires a statistical investigation.	question	1.1.A.2 Statistical studies are necessary when the population is too large or it is too difficult to collect data from every item or individual in the population. 1.1.A.3 A datum (singular form of data) is a piece of information about an item or individual. A collection of data is called a data set. 1.1.A.4 A population consists of all items or individuals of interest. The population size is represented by the symbol N. 1.1.A.5 A sample selected for study is a subset of the population from which data are obtained. The number of items in the sample, called the sample size, is represented by the symbol n. 1.1.A.6 Each component of a statistical study and the resulting calculations can be related to an aspect of the corresponding real-world context from which the components were derived. This identification of a statistical result with the corresponding contextual component is what is meant by “in context.”
	2.A Identify information to answer a question or solve a problem.	1.1.A Identify components within a statistical study	1.1.B.1 An investigative question for a specific study should have a defined purpose and should not be changed based on the data analysis or results. 1.1.B.2 An investigative question should be posed so that the required data can be collected and analyzed.
Day 2 1.2 <i>Variables</i>	2.A Identify information to answer a question or solve a problem.	1.2.A Identify observational units, variables, parameters, and statistics from a statistical study or data set	1.2.A.1 An observational unit is an item or individual from which a datum is collected. 1.2.A.2 A variable is a characteristic that may change from one observational unit to another. 1.2.A.3 Data collected on numerical and categorical variables measured on observational units, including photographs, sounds, videos, and text can convey meaningful information. 1.2.A.4 A parameter is a numerical attribute or summary of the variable of interest for a population. 1.2.A.5 A statistic is a numerical attribute or summary of the variable of interest for a sample. The value of a statistic from a certain sample is often not equal to the unknown value of the population parameter but may provide the basis for making inferences about the population parameter.
		1.2.B Identify types of variables.	1.2.B.1 A categorical variable, also called a qualitative variable, takes on values that are category names or group labels. 1.2.B.2 A quantitative variable, also called a numerical variable, takes on numerical values for a measured or counted q
		1.2.C Identify types of quantitative variables.	1.2.C.1 A discrete quantitative variable can take on a countable number of values. The number of values may be finite or countably infinite, as with the whole numbers. 1.2.C.2 A continuous quantitative variable can take on an infinite number of possible values within a given interval. The number of values the variable can take on is measurable but not countable. This variable can take on all possible values between any pair of values.

Day 3 1.3 <i>Tabular Representation and Summary Statistics for One Categorical Variable</i>	3.A Construct tabular and graphical representations of data and distributions.	1.3.A Construct categorical one-variable tabular	1.3.A.1 A frequency table shows the number of observational units in each category of a categorical variable. 1.3.A.2 A relative frequency table shows the proportion of observational units in each category of a categorical variable.	
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.3.B Describe categorical one-variable tabular representations with summary statistics.	1.3.B.1 Percentages, relative frequencies, and ratios all provide the same information as proportions. 1.3.B.2 Counts and relative frequencies of categorical variables reveal information that can be used to justify claims about the variables in context	
Day 4 1.4 <i>Graphical Representations for One Categorical Variable</i>	3.A Construct tabular and graphical representations of data and distributions.	1.4.A Construct categorical one-variable graphical representations.	1.4.A.1 Bar charts, also called bar graphs, display frequencies (counts) or relative frequencies (proportions) for the categories of a single categorical variable. Each bar on a bar chart represents a category of the categorical variable of interest. The height or length of each bar corresponds to the frequency or relative frequency of the observational units in each category. 1.4.A.2 Pie charts are used to display frequencies (counts) or relative frequencies (proportions) for categorical data. Each slice on a pie chart represents a category of the categorical variable of interest. The area of each slice, as a fraction of the total area, corresponds to the relative frequency of observational units falling within each category. The sum of the slices' areas together will equal 1, or 100% of the total area.	
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.4.C Compare multiple categorical one-variable tabular and graphical representations.	1.4.C.1 Frequency and relative frequency tables, bar charts, and pie charts can be used to compare two or more data sets in t	
	4.B Justify a claim based on statistical calculations and results	1.4.B Justify a claim using categorical one-variable graphical representations.	1.4.B.1 Graphical representations of a categorical variable reveal information that can be used to justify claims about the variable in context.	
Day 5 1.5 <i>Graphical Representations for One Quantitative Variable</i>	3.A Construct tabular and graphical representation	1.5.A Construct quantitative one-variable graph	1.5.A.1 Histograms, stem-and-leaf plots, and dotplots provide a visual representation of the distribution of the values of a quantitative variable. These graphs show the frequency or relative frequency of the quantitative variable values or intervals of values and maintain the natural ordering, smallest to largest, of the quantitative variable. 1.5.A.2 A histogram places the observed values of the quantitative variable into ordered intervals, or bins, along the horizontal axis. Each bar represents an interval or bin, and the height of each bar shows the frequency or relative frequency of the observations within that interval. Altering the interval widths, or bin widths, can change the appearance of the histogram. Alternatively, a histogram can be constructed with bins	Gallery Walk Have students work in groups of four to construct a dotplot, stem-and-leaf plot, histogram, or boxplot for a set of student-generated data (e.g., time in minutes to get to school). After the gallery walk, discuss what information can be seen more easily in each graph (e.g., box plots can easily show the IQR).

			on the vertical axis with bars appearing horizontally. 1.5.A.3 A stem-and-leaf plot splits each value of the quantitative variable into two parts: a “stem” (the first digit or digits) and a “leaf” (usually the single digit after the stem digit or digits). Both stems and leaves are ordered from smallest to largest. 1.5.A.4 A dotplot represents each value of the quantitative variable by a dot. Each dot is placed above the horizontal or beside the vertical axis corresponding to the value of that observation, with nearly identical values stacked on top of each other.	
Day 6 1.6 <i>Descriptions for One Quantitative Variable Distributions</i>	4.A continued on next page Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.6.A Describe distributions of quantitative one-variable graphical representations.	1.6.A.1 Descriptions of the distribution of one quantitative variable include shape, center, and variability (spread) as well as any unusual features such as outliers, gaps, or clusters in context. 1.6.A.2 The shape of the distribution of one quantitative variable is skewed to the right (positively skewed) if the right tail (toward larger values) is longer than the left. The shape of the distribution is skewed to the left (negatively skewed) if the left tail (toward smaller values) is longer than the right. The shape of the distribution is approximately symmetric if the left half is approximately the mirror image of the right half. 1.6.A.3 Distributions of one quantitative variable with one main peak are called unimodal. Distributions with two prominent peaks are called bimodal. A distribution in which each frequency or each relative frequency is approximately the same with no prominent peaks is approximately uniform. 1.6.A.4 Outliers for one quantitative variable are data points that are unusually small or large relative to the rest of the data. 1.6.A.5 A gap is a region in a distribution between two values in which there are no observed data. 1.6.A.6 Clusters are concentrations of values usually separated by gaps.	
	4.B Justify a claim based on statistical calculations and results.	1.6.B Justify a claim using distributions of quantitative one-variable graph	1.6.B.1 Graphical representations of a quantitative variable may reveal information that can be used to justify claims about the variable in context.	
Day 7 1.7 <i>Summary Statistics for One Quantitative Variable</i>	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	1.7.A Calculate measures of center and position for quantitative data	1.7.A.1 Two commonly used measures of center in the distribution of a quantitative variable are the mean and median. 1.7.A.2 The mean is the sum of all the values divided by the number of values and can be found with and without using technology. For a sample, the mean is denoted by $\bar{x}$: $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$, where x_i represents the i^{th} data point in the sample and n represents the number of data values in the sample. 1.7.A.3 The median is the middle value when the data set is ordered from smallest to largest and can be found with and without using technology. One common method for determining the median of a data set with an even number of values is to use the mean of the two middle values. A common method for determining the median of a	

			data set with an odd number of values is to use the value in the middle of all the values. 1.7.A.4 In an ordered data set, the smallest value is the minimum value, and the largest value is the maximum value. 1.7.A.5 The first quartile, denoted by Q1, is the median value of the lower half of the ordered data set from the minimum value to the position of the median. Approximately 25% of the values in the data set are less than or equal to Q1. The third quartile, denoted by Q3, is the median value of the upper half of the ordered data set from the position of the median to the maximum value. Approximately 75% of the values in the data set are less than or equal to Q3. The second quartile, Q2, is also the median of the data set. Q1 and Q3 form the boundaries for the middle 50% of values in an ordered data set. 1.7.A.6 The pth percentile is the value that has p% of the data less than or equal to it when the data set is ordered from smallest to largest. The first and third quartiles are the 25th and 75th percentiles, respectively
		1.7.B Calculate measures of variability for quantitative data.	1.7.B.1 Three commonly used measures of variability (or spread) in the distribution of a quantitative variable are the range, interquartile range, and standard deviation. 1.7.B.2 The range is the difference between the maximum data value and the minimum data value. 1.7.B.3 The interquartile range (IQR) is the difference between the third and first quartiles: Q3 – Q1. 1.7.B.4 The standard deviation is a typical deviation of the data values from their mean and can be found with and without using technology. The sample standard deviation is $s = \sqrt{\frac{1}{n-1} \sum (x_i - \bar{x})^2}$ denoted by s and calculated by $s = \sqrt{\frac{1}{n-1} \sum (x_i - \bar{x})^2}$, where x_i is the data value, $\bar{x}$ is the mean, and n is the number of data values in the sample. The square of the sample standard deviation, s^2, is called the sample variance.
		1.7.C Calculate different units of measurement for summary statistics.	1.7.C.1 Changing units of measurement affects the values of the calculated statistics.
		1.7.D Calculate outliers for quantitative data.	1.7.D.1 There are many methods for determining potential outliers.
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.7.E Compare multiple quantitative one-variable summary statistics	1.7.E.1 Summary statistics can be used to compare features of two or more independent samples, including center, variability, shape, and outliers.
	4.B Justify a claim based on statistical calculations and results	1.7.F Justify the selection of a summary statistic for describing quantitative data.	1.7.F.1 The median and IQR are considered a resistant (or robust) measure of center and measure of variability, respectively, because outliers do not greatly (if at all) affect their values. Because outliers can affect their values greatly, the mean is considered a nonresistant (or non-robust) measure of center, and the range and standard deviation are considered nonresistant (or nonrobust) measures of variability. 1.7.F.2 Summary statistics of a quantitative variable may reveal information that can be used to justify claims about the variable in context

Day 8 1.8 Graphical Representations of Summary Statistics for One Quantitative Variable	3.A Construct tabular and graphical representations of data and distributions.	1.8.A Construct quantitative one-variable graphical representations of summary statistics.	1.8.A.1 A five-number summary is made up of the minimum data value, the first quartile (Q1), the median, the third quartile (Q3), and the maximum data value. 1.8.A.2 A boxplot is a graphical representation of the five-number summary (minimum, first quartile, median, third quartile, maximum). The box represents the middle 50% of data, with a line at the median and the ends of the box corresponding to the quartiles. Lines (“whiskers”) that represent 25% of the data extend from the first quartile to the minimum and from the third quartile to the maximum. If there are outliers in the data, the whiskers extend to the most extreme data values that are not outliers, and outliers are usually	
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.8.B Describe quantitative one-variable graphical representations of summary statistics based on the relationship of the mean and the median.	1.8.B.1 If a distribution is relatively symmetric, then the values of the mean and median are relatively close to each other. If a distribution is skewed right, then the value of the mean is usually larger than the median. If the distribution is skewed left, then the value of the mean is usually smaller than the median.	
Day 9 1.9 Comparisons of the Distributions for One Quantitative Variable	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	1.9.D Calculate z-scores with population parameters.	1.9.D.1 A standardized score measures the number of standard deviations a data value falls above or below the mean. 1.9.D.2 A z-score is calculated as $\frac{x_i - \mu}{\sigma}$, where x_i is the data value, μ is the population mean, and σ is the population standard deviation. A z-score measures how many standard deviations a data value is above (positive z-score) or below (negative z-score) the mean. When the population mean and standard deviation are unknown, the sample mean and standard deviation may be used to determine a z-score	Match Mine Create a blank 3×3 grid as a gameboard and nine index cards depicting distributions of differing shapes, center, skewness, clustering, etc. Pair up students and place a folder between the partners so they cannot see each other’s boards. Distribute a blank grid and a set of cards to each partner. Student A arranges the cards however they choose on their grid and then describes that arrangement so that Student B can attempt to match it on their own board. When the pair thinks they have correctly made all matches, have them compare their arrangements to see how well they did.
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	1.9.A Compare multiple quantitative one-variable graphical representations.	1.9.A.1 Graphical representations of a quantitative variable can be used to compare important features between two or more distributions of the same quantitative variable. Histograms, back-to-back stem-and-leaf plots, and dotplots may be used to compare center, variability, shape, outliers, clusters, or gaps in two or more distributions. Boxplots may be used to compare center, variability, outliers, and skewness (or symmetry).	
		1.9.B Compare multiple quantitative one-variable graphical representations of summary statistics.	1.9.B.1 A comparison of graphical representations for two or more distributions can include any of the numerical summaries (e.g., mean, standard deviation, etc.	
	4.B Justify a claim based on	1.9.C Justify a claim using multiple		

	statistical calculations and results.	quantitative one-variable graphical representations.		
	4.C Describe distributions and compare relative positions of points within a distribution.	1.9.E Compare z-scores as measures of relative position for distributions.	1.9.E.1 z-scores may be used to compare relative positions of individual values within a distribution or between distributions.	
Day 10	Mid-Unit Assessment			
Day 11 1.10 <i>The Investigative Question Revisited and Data Collection</i>	1.A Determine a valid investigative question that requires a statistical investigation.	1.10. A Determine the components of an investigative question within a statistical study.	1.10.A.1 The first component of an investigative question should guide the data collection process and should be phrased in terms of the variable(s) of interest in the study. 1.10.A.2 The second component of an investigative question should guide the data analysis choice. 1.10.A.3 The third component of an investigative question should indicate the type(s) of conclusion(s) applicable from the study. The investigative question should provide the population to which the conclusions will be applicable and, in the case of an experiment that uses random assignment, a cause-and-effect conclusion.	Odd One Out After modeling an odd one out example, have students form groups of four and give each student a description of a statistical study. Explain that three of the studies are of the same type (observational or experimental) and one is different. Have students work together in their groups to determine which study is the odd one out and explain why.
	2.A Identify information to answer a question or solve a problem.	1.10.B Identify a census.	1.10.B.1 A census consists of recording information from all items or individuals in a population.	
		1.10.C Identify an experiment	1.10.C.1 An experiment is a study in which a researcher assigns conditions, or treatments, to experimental units to explore an investigative question of interest about the population. 1.10.C.2 The experimental unit is the observational unit to which the treatment is assigned. When experimental units consist of people, they are sometimes referred to as subjects or participants. 1.10.C.3 An explanatory variable, or factor, is a variable whose different categories, or levels, are imposed on the experimental units. The different categories, or levels, of the explanatory variable are called treatments. When there is more than one explanatory variable, the combinations of the categories, or levels, of the explanatory variables are called treatments. 1.10.C.4 A response variable is an outcome measured on each experimental unit after the treatment has been administered.	
		1.10.D Identify an observational study.	1.10.D.1 An observational study is a study where treatments are not imposed. The	

			researcher records the values of the variables of interest in order to explore an investigative question of interest. 1.10.D.2 A prospective study is one in which the observational units of study are selected at a point in time, and data are gathered both at that time and into the future. 1.10.D.3 A retrospective study is one in which the observational units of study are selected at a point in time and data from the past are gathered. 1.10.D.4 A survey is an observational study in which the data are collected from humans using a standard set of questions. 1.10.D.5 A confounding variable in an observational study provides an alternative explanation for the observed relationship between the explanatory and response variables determined in the study, thereby reducing the possibility of concluding a causal relationship between the explanatory and response variables of interest. To be a confounding variable, a variable must be associated with both the explanatory variable and the response variable.	
	2.B Justify an appropriate method for ethically gathering and representing data.	1.10.E Justify the appropriateness of generalizations for a statistical study.	1.10.E.1 A sample is considered random when all observational units in the sample are selected from the population using some type of random mechanism, such as a random number generator. 1.10.E.2 When observational units, or experimental units, in a sample are randomly selected from a population, it is appropriate to make generalizations about the entire population of individuals from which the sample was selected. 1.10.E.3 A sample is not randomly selected when observational units are deliberately chosen or volunteer themselves to be in the sample. 1.10.E.4 When observational units, or experimental units, in a sample are not randomly selected from a population, it is appropriate to make generalizations only about a population of individuals that are similar to those used in the study	
Day 12 1.11 Random Sampling	2.A Identify information to answer a question or solve a problem.	1.11.A Identify a sampling method given a description of a study.	1.11.A.1 Sampling without replacement is a sampling strategy in which an observational unit from a population can be selected only once. The observational unit is not returned to the population before subsequent selections of observational units are made, so there is no chance that the observational unit can be selected again. 1.11.A.2	Password-Style Games After completing the lessons on sampling and surveying, use the following nine terms in a password-style game: simple random sample, stratified random sample, cluster random sample, systematic

			Sampling with replacement is a sampling strategy in which an observational unit from the population can be selected more than once. The observational unit is returned to the population before subsequent selections of observational units are made, so it is possible that the observational unit could be selected again. 1.11.A.3 In a simple random sample (SRS) of size n, every sample of the size n has the same chance of being selected. This method is the basis for many types of sampling mechanisms. There are several procedures to obtain a simple random sample; for example, using a random number generator or randomly selecting numbered slips of paper. 1.11.A.4 A stratified random sample involves the division of all individuals in a population into non-overlapping groups, called strata, based on one or more shared attributes or characteristics (homogeneous grouping). Within each stratum a simple random sample is selected, and the selected individuals are combined to form one sample. 1.11.A.5 A cluster random sample involves the division of a population into smaller groups, called clusters. Ideally, each cluster mirrors the heterogeneity of the population, with clusters similar to one another. A simple random sample of clusters is selected from the population to form the sample of clusters. Data are collected from all observational units in each of the selected clusters. 1.11.A.6 A systematic random sample is a method in which sample members from a population are selected according to a random starting point and a fixed, periodic interval between successive sampling units.	random sample, bias, voluntary response bias, undercoverage bias, nonresponse bias, and response bias. Pair students and ask one student per pair to sit facing the back of the classroom. Project or write the terms on the board one at a time while the front-facing students give clues to their partners, who try to guess the terms. The winner is the pair that guesses the most terms correctly.
	2.B Justify an appropriate method for ethically gathering and representing data.	1.11.B Justify the appropriateness of a sampling method.	1.11.B.1 Each random sampling method has different characteristics that make it more appropriate for sampling populations depending on the question being investigated.	
Day 13 1.12 <i>Potential Problems with Sampling</i>	2.A Identify information to answer a question or solve a problem.	1.12. A Identify potential sources of bias in sampling methods.	1.12.A.1 Bias in a sampling method is a systematic error in the sampling procedure that results in a statistic being consistently larger or consistently smaller than the parameter the statistic is used to estimate. 1.12.A.2 Voluntary response bias is a bias that may occur when a sample consists entirely of volunteers. 1.12.A.3 Undercoverage bias may occur when the sampling method fails to include part of the	

			population or a part of the population is less likely to be selected based on the sampling method. 1.12.A.4 Nonresponse bias may occur because of a failure to obtain responses from some individuals chosen to be sampled. The respondents and nonrespondents could differ significantly in ways that are important for the study. 1.12.A.5 Response bias may occur when responses to a survey or measurements of observational units tend to differ from the “true” value in one direction. Examples include questions that are confusing or leading (question wording bias) or self-reported responses. 1.12.A.6 Nonrandom sampling methods (e.g., samples chosen by convenience or voluntary response) introduce potential bias because they do not use random chance to select the individuals.	
Day 14 1.13 <i>Experimental Design</i>	2.A continued on next page Identify information to answer a question or solve a problem.	1.13.A Identify elements of a well-designed experiment.	1.13.A.1 A well-designed experiment should include the following: • Comparisons of at least two treatment groups, one of which could be a control group • Random assignment of treatments to experimental units • Replication • Direct control of potential extraneous sources of variation in the response 1.13.A.2 A control group is a collection of experimental units that are created for comparison. A control group may be given a treatment different from the treatment of interest to determine if the treatment of interest has an effect (e.g., a treatment with an inactive substance, a placebo, may be given). 1.13.A.3 The placebo effect is the difference between the average response to a placebo and the average response to no treatment. 1.13.A.4 In a single-blind, also called single-masked, experiment, participants do not know which treatment they are receiving, but members of the research team who interact with them know which treatment each participant is receiving, or vice versa. 1.13.A.5 In a double-blind, also called double-masked, experiment, neither the participants nor the members of the research team who interact with them know which treatment each participant is receiving. 1.13.A.6 An extraneous source of variation, also referred to as an extraneous variable, is a variable that is known (or believed) to affect the response but is not an explanatory variable being studied. 1.13.A.7 The purpose of random assignment is to create treatment groups that are as similar as possible with respect to extraneous sources of variation. If random assignment is successful, the respective distributions of each extraneous variable will be approximately the same for all the treatment groups. 1.13.A.8 A confounding variable in an experiment is a variable that is related to the explanatory variable in such a way that it is difficult to determine which variable, explanatory or confounding, is influencing the change in the response variable. However, in a well-designed experiment, the	

			potential for confounding variables is reduced. 1.13.A.9 Replication within an experiment means more than one experimental unit is assigned to each treatment 1.13.A.10 Direct control in an experiment means keeping the settings of certain potential extraneous sources of variation in the response variable the same from experimental unit to experimental unit.
		1.13.B Identify experimental designs.	1.13.B.1 In a completely randomized design, treatments are assigned to experimental units completely at random. Often the number of experimental units assigned to each treatment will be the same, but the sample sizes in each treatment do not have to be the same. 1.13.B.2 A blocking variable is a source of extraneous variation in the response variable. In a randomized block design, the experimental units are first grouped according to similar values of a blocking variable. These groups are called blocks. Units within the same block are homogeneous with respect to the blocking variable. After the blocks are formed, the treatments are randomly assigned to experimental units within each block so that all treatments occur within every block. 1.13.B.3 The purpose of blocking is to separate the variation in the response caused by the blocking variable from the rest of the extraneous variation in the response. Blocking allows for more precise comparisons of the response across the treatments. Within a block, the treatments can be compared without having to worry about variation in the response caused by changes in the blocking variable. 1.13.B.4 A matched pairs design is a randomized block design with only two treatments. Experimental units are arranged in pairs by matching on one or more extraneous sources of variation in the response variable. Each pair receives both treatments by randomly assigning one treatment to one member of the pair and the other treatment to the second member of the pair. Alternatively, each experimental unit may get both treatments while the order of the treatments is randomized.
	2.B Justify an appropriate method for ethically gathering and representing data.	1.13.C Justify the appropriateness of a particular experimental design.	1.13.C.1 One experimental design may be more appropriate than another experimental design based on the goals of the investigative study, the characteristics of the population, and the sample and variables involved.
		1.13.D Justify the appropriateness of the conclusions based on a well-designed experiment.	1.13.D.1 Using random assignment of treatments to experimental units allows for cause-and-effect conclusions between the explanatory and the response variables because the potential for confounding variables is reduced. 1.13.D.2 Depending on the experimental unit, it may be unethical or difficult to randomly select experimental units to participate in an experiment. In that case, the study's experimental units are obtained from volunteers and will represent the population of experimental units similar to those who participated in the study
Day 15	Unit Test: AP Progress Check 1		

Unit Title		
Unit 2: Probability, Random Variables, and Probability		
Relevant Standards: Bold indicates priority		
3.A, 3.B, 3.C, 3.D, 4.A, 4.B, 4.C, 4.D		
Essential Question(s)		Enduring Understanding(s)
- How might we represent the reading habits of people with and without an electronic reading device to help describe similarities and/or differences between the two groups? - How can the weather be predicted? - Why might the data we collect influence decisions being made?		- Probability provides a framework for quantifying uncertainty and modeling chance behavior. - Random variables and probability distributions describe the long-run behavior of random processes. - Expected value and variability help inform decisions involving risk and uncertainty. - Simulations can be used to investigate and approximate complex probability situations.
Demonstration of Learning		Pacing for Unit
Progress Check 2 - Multiple-choice: ~42 questions - Free-response: 2 questions - Question 1: Multi-Focus on Practices 3 and 4 - Question 2: Multi-Focus on Practices 3 and 4 The additional online question bank will be used to create a mid-unit quiz and other checks for understanding throughout the unit. College Board AP Exam Section I MCQ Weighting: 15–25% - Probability Rules: Applying rules for complements, unions, intersections, conditional probability, and testing for independence. - Random Variables & Distributions: Calculating expected values (mean) and standard deviations for discrete random variables. - Special Distributions: Solving problems involving normal distributions (z-scores, percentiles) and binomial probability models. - Sampling Distributions: Understanding sampling variability and applying the Central Limit Theorem (CLT) to understand sample behavior. See AP Exam Overview for more information		14 Block Periods <i>NOTE:</i> Midterm & AP Exam Review / Mock Exams - 11 Blocks Flexibility/Buffer/Post-AP Exam Project - 6 Blocks
Family Overview (link below)		Integration of Technology:
AP Statistics FAMILY MATERIALS		- Use graphing calculators, spreadsheets, or simulation software to model random events. - Generate sampling distributions through simulation. - Explore probability distributions graphically and numerically. - Use graphing calculators and Desmos to calculate probabilities and expected values.
Unit-specific Vocabulary		Aligned Unit Materials, Resources, and Technology (beyond core resources)
Chance process Probability Simulation Random outcome Sample space Event	Probability distribution Expected value Mean of a random variable Variance Standard deviation	www.mathmedic.com www.desmos.com www.stapplet.com https://skewthescrpt.org/

Complement Independent events Conditional probability Random variable Discrete random variable	Binomial setting Binomial random variable Binomial probability Normal distribution Parameter Statistic		
Opportunities for Interdisciplinary Connections:		Anticipated misconceptions	
• Science: genetics and inheritance probabilities. • Business: risk analysis and expected profit. • Economics: forecasting and decision-making under uncertainty. • Computer Science: simulations and randomization algorithms. • Games and Sports Analytics: modeling outcomes and probabilities.		1. Believing in the "Law of Averages" (Gambler's Fallacy) for independent outcomes. 2. Confusing independent events with mutually exclusive (disjoint) events. 3. Subtracting variances or directly adding standard deviations when combining random variables. 4. Confusing a single sample distribution with a population distribution or a sampling distribution. 5. Misinterpreting the Central Limit Theorem as stating that the population or sample data distribution becomes Normal as sample size increases.	
Connections to Prior Units		Connections to Future Units	
• Builds on concepts of variability and data collection introduced in Unit 1. • Uses distributions and graphical interpretation skills developed in Unit 1.		• Provides the theoretical foundation for statistical inference. • Supports understanding of sampling distributions, confidence intervals, and hypothesis testing. • Introduces concepts used in regression inference.	
Differentiation through <i>Universal Design for Learning</i>			
	Engagement	Representation	Action & Expression
2.A	Investigate bivariate relationships in popular culture (e.g., screen time vs. sleep duration) to formulate explanatory/response questions.	Matrix mapping bivariate questions to variable types (Quantitative vs. Quantitative; Categorical vs. Quantitative).	Write two bivariate investigative questions for a self-selected real-world dataset and justify why they require two-variable analysis.
3.A	"Variable Match" challenge identifying explanatory (X) and response (Y) variables in real-world studies (e.g., study time vs. exam score).	Color-coded scatterplot axes diagram identifying independent variable (X-axis) and dependent variable (Y-axis).	Construct a bivariate variable map identifying lurking variables that could influence bivariate relationships.
3.B	Evaluate real-world scatterplots where improper extrapolation led to faulty public policy or business decisions.	Visual guide highlighting the dangers of extrapolation beyond the range of observed X-data.	Write an ethical evaluation report assessing whether a linear model can be legitimately applied to a specific prediction scenario.
3.C	"Bivariate Graph Off" creating scatterplots, residual plots, two-way tables, and segmented bar charts from live class data.	Step-by-step visual checklist for scatterplot construction (scaling, plotting points, drawing Least-Squares Regression Lines (LSRL)).	Generate scatterplots and residual plots using choice of technology (TI-84, Desmos, R, or Python) or physical graphing tools.
3.D	"Predict the Future" lab calculating LSRL equations to predict outcomes (e.g., predicting movie revenue from budget).	Dual-column reference chart comparing manual slope/intercept formulas with computer software regression outputs.	Solve prediction problems by submitting worked calculations, annotated calculator screenshots, or recorded video demonstrations.
4.B	"Scatterplot Sorting Game" categorizing bivariate graphs by Direction, Unusual features, Form, and Strength (DUFS).	Anchor chart breaking down the DUFS acronym with visual examples of linear/nonlinear forms and varying correlation strengths.	Deliver an audio recording or written essay describing the relationship in a scatterplot using contextual DUFS terminology.
4.D	"Residual Detective" evaluating residual plots to determine if a linear model is appropriate for a given dataset.	Side-by-side visual comparison showing a random residual scatter (linear appropriate) vs. a curved residual pattern (linear inappropriate).	Write a formal justification argument verifying or refuting model linearity based on residual plot patterns and r^2 values.

Supporting Multilingual/English Learners (CELP standards)

	Emerging	Expanding	Bridging
2.A	Identify given probabilities from a word problem using a color-coded text highlighter guide.	Extract conditional probability statements and random variable parameters from word problems using a structured probability organizer.	Parse complex conditional and independent probability scenarios, identifying underlying probability distributions (Binomial vs. Geometric).
3.A	Complete blank 2-way tables, Venn diagrams, or probability distribution tables using simple addition/subtraction prompts.	Construct tree diagrams, Venn diagrams, and probability distribution histograms from scenario prompts using guided templates.	Construct complex two-way tables, tree diagrams, and discrete/continuous probability density functions for multi-step random variables.
3.B	Calculate simple probabilities using visual counting cards.	Calculate expected value and standard deviation of discrete random variables using guided formula sheets.	Compute linear combinations of random variables and normal distribution proportions.
3.C	Calculate basic probabilities using addition and multiplication formulas with color-coded key charts.	Calculate conditional probabilities and binomial/geometric probabilities using formula templates or calculator guides.	Calculate complex joint, conditional, binomial, geometric, and normal distribution probabilities across multi-step scenarios.
3.D	Identify formulas for Binomial mean and standard deviation on an annotated formula card.	Calculate parameters for binomial and geometric random variables using structured computational guides.	Calculate and transform parameters for combinations and transformations of random variables, justifying variance additions.
4.B	Match probability calculation results to visual "Likely / Unlikely" decision cards.	Justify whether two events are independent or mutually exclusive using structured mathematical proof templates.	Formulate logical arguments assessing whether an observed event is statistically unusual based on calculated probability distributions.
4.D	Identify whether calculated expected values represent long-run averages using true/false sentence strips.	Interpret expected value and standard deviation of a random variable in context using sentence frames ("Over many trials, the average...").	Critically interpret expected value, variance, and probability outcomes in real-world risk, insurance, and decision-making contexts.

Unit Outline

Lesson Sequence	Skills	Learning Objective	Essential Knowledge	AP Activities (if available)
Day 1 2.1 <i>Tabular and Graphical Representations for the Distributions of Two Categorical Variables</i>	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	2.1.A Compare tabular and graphical representations for the relationship between two categorical variables.	2.1.A.1 A two-way table, also called a contingency table, can be used to summarize and compare data for two categorical variables. The entries in the cells of the table can be frequencies (i.e., counts) or relative frequencies (i.e., proportions). 2.1.A.2 Side-by-side bar charts, segmented bar charts, and mosaic plots are examples of graphs used to display the relationship between two categorical variables. In these graphs, the frequency or relative frequency of each category, or level, of one of the categorical variables is displayed for each category of the other categorical variable. 2.1.A.3 Graphical representations of two categorical variables can be used to compare the relationship of one categorical variable across the levels of the other categorical variable and determine whether the two variables are associated.	Quiz-Quiz-Trade Give students a card containing a question with a two-way table and have them write the answer on the back. Then have them stand up and find a partner. One student quizzes the other, and then they reverse roles. Have them switch cards, find a new partner, and repeat the process.
	4.B Justify a claim based on statistical	2.1.B Justify a claim using tabular and graphical	2.1.B.1 Tabular and graphical representations for the distributions of two categorical variables may reveal information that can be used to justify	

	calculations and results.	representations for the distributions of two categorical variables	claims about the variable in context.	
Day 2 Part 1 2.2 <i>Summary Statistics for Two Categorical Variables</i>	3.B Calculate summary statistics, relative positions of points within a distribution, , and predicted responses.	2.2.A Calculate summary statistics from two-way tables.	2.2.A.1 A joint relative frequency in a two-way table is a cell frequency divided by the total for the entire table. 2.2.A.2 A marginal relative frequency in a two-way table is a row total divided by the total for the entire table or a column total divided by the total for the entire table. 2.2.A.3 A conditional relative frequency is a relative frequency computed by restricting to a particular level, or category of interest. A conditional relative frequency can be a cell frequency in a row divided by the total for that row or it can be a cell frequency in a column divided by the total for that column.	
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	2.2.B Compare summary statistics for two categorical variables.	2.2.B.1 Summary statistics for two categorical variables can be used to compare distributions for evidence of association between the two variables.	
	4.B Justify a claim based on statistical calculations and results.	2.2.C Justify a claim using summary statistics for two categorical variables.	2.2.C.1 Summary statistics for two categorical variables may reveal information that can be used to justify claims about the variables in context.	
Day 2 Part 2 2.3 <i>Estimating Probabilities Using Simulation</i>	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	2.3.A Estimate probabilities using simulations.	2.3.A.1 A random process generates results that are determined by chance. 2.3.A.2 An outcome is the result of one trial of a random process. 2.3.A.3 An event is a collection of outcomes. 2.3.A.4 Simulation is a way to model random events such that simulated outcomes closely match real-world outcomes. All possible outcomes are associated with a value to be determined by chance. Record the counts of simulated outcomes and the count total. 2.3.A.5 The probability of an outcome or event is its long-run relative frequency—that is, its relative frequency over a large number of trials. 2.3.A.6 The relative frequency of an outcome or event determined from empirical data can be used to estimate the actual, or true, probability of that outcome or event. 2.3.A.7 The law of large numbers states that for independent trials, as the number of trials increases, the long-run relative frequency of the outcome or event gets closer and closer to a single value.	
Day 3 2.4 <i>Introduction to Probability</i>	3.C Calculate and estimate expected	2.4.A Calculate probabilities for events and	2.4.A.1 The sample space of a random process is the set of all possible nonoverlapping outcomes. The probability of the sample space is 1.	Think-Pair-Share Provide students with a set of five probability questions: one for each of the

	counts, percentages, probabilities, and intervals.	their complements.	2.4.A.2 If all outcomes in the sample space are equally likely, then the theoretical probability an event E will occur is $\frac{\text{number of outcomes in event } E}{\text{total number of outcomes in the sample space}}$ The probability of event E occurring is written as P(E) . 2.4.A.3 The probability of an event is a number between 0 and 1, inclusive. 2.4.A.4 The probability of the complement of an event E, which can be written as E', $\bar{E}$, or E^C (i.e., the probability of “not E”) is equal to $1- P(E)$.	complement rule, a mutually exclusive event, a conditional probability, the general multiplication rule, and an independent event. Ask students to individually identify the formula needed to solve each problem without doing the final calculations. Then have them share their thoughts with a partner.
Day 4 2.5 <i>Mutually Exclusive Events</i>	4.B Justify a claim based on statistical calculations and results.	2.5.A Justify why two events are mutually exclusive (or disjoint) using joint probability.	2.5.A.1 The probability that events A and B both will occur, sometimes called the joint probability, is the probability of the intersection of A and B. Joint probability is defined as P (A intersect B) or P (A ∩ B) . 2.5.A.2 Two events are mutually exclusive, or disjoint, if they cannot occur at the same time. This means that if two events are mutually exclusive, then P (A ∩ B) = 0.	
Day 5 2.6 <i>Conditional Probability</i>	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	2.6.A Calculate conditional probabilities.	2.6.A.1 The probability that event A will occur given that event B has occurred is called a conditional probability and is written as P(A B). Conditional probability is defined as $P(A B)=\frac{P(A\cap B)}{P(B)}$ 2.6.A.2 The general multiplication rule states that the probability that events A and B both will occur is equal to the probability that event A will occur multiplied by the conditional probability that event B will occur given that event A has occurred. The multiplication rule is defined as $P(A\cap B)=P(A)\cdot P(B A)$	
Day 6 2.7 <i>Independent Events and Unions of Events</i>	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	2.7.A Calculate probabilities for independent events and for the union of two events	2.7.A.1 Events A and B are independent if and only if knowing whether event A has occurred (or will occur) does not change the probability that event B will occur. When events A and B are independent, then P(A B) = P(A), P(B A) = P(B), and P (A ∩ B) = P(A)• P(B). 2.7.A.2 The probability that event A or event B (or both) will occur is the probability of A union B. The probability of the union is defined as P (A ∪ B). 2.7.A.3 P (A union B) = P (A) + P(B) - P(A intersect B) , or P (A ∪ B) = P(A) +P(B) - P(A ∩ B).	
Day 7	Mid-Unit Assessment			

Day 8 2.8 <i>Introduction to Random Variables and Probability Distributions</i>	3.A Construct tabular and graphical representations of data and distributions.	2.8.A Construct a probability distribution for a discrete random variable.	2.8.A.1 A random variable is a variable whose values have numerical outcomes that result from a random phenomenon. 2.8.A.2 A probability distribution for a discrete random variable shows the probability associated with every possible value of the random variable. The sum of the probabilities over all possible values of a discrete random variable is 1. 2.8.A.3 A discrete probability distribution can be determined using the rules of probability or estimated with a simulation. 2.8.A.4 A discrete probability distribution can be represented as a graph, table, or function showing the probabilities associated with values of a random variable. 2.8.A.5 A cumulative probability distribution can be represented as a table or function and shows the probability of being less than or equal to each value of the discrete random variable.	
Day 9 2.9 <i>Parameters of Random Variables</i>	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	2.9.A Calculate the mean and standard deviation for a discrete random variable.	2.9.A.1 A numerical value measuring a characteristic of a probability distribution of a random variable, or a population, is a parameter. The value of a parameter is a single, fixed value. 2.9.A.2 The expected value (or mean) of a probability distribution is a parameter and is denoted by $E(X)$ or μ_x . For a discrete random variable X, the expected value is calculated as $\mu_x = \sum x_i \cdot P(x_i)$ where x_i is the possible value of the random variable and $P(x_i)$ is the probability of the possible value of the random variable. The expected value can be interpreted as the long-run average outcome of the random variable. The discrete random variable can only take on values that are countable or finite. 2.9.A.3 The standard deviation of a probability distribution is a parameter represented by $SD(X)$ or σ_x . For a discrete random variable X, the standard deviation is calculated as $\sigma_x = \sqrt{\sum (x_i - \mu_x)^2 \cdot P(x_i)}$ where x_i is the possible value of the random variable, μ_x is the mean, and $P(x_i)$ is the probability of the possible value of the random variable. The standard deviation can be interpreted as the typical deviation of the values of the random variable from the mean value (or expected value) of the random variable over the long run. The square of the standard deviation of a random variable is called the variance of the random variable and is denoted as $V(X)$ or σ_x^2 .	
	4.D Interpret statistical calculations and results to assess meaning or a claim.	2.9.B Interpret the mean and standard deviation for a discrete random variable.	2.9.B.1 The mean and standard deviation for the probability distribution of a discrete random variable should be interpreted in the context of a specific population.	
Day 10 2.10 <i>The Binomial Distribution</i>	4.B Justify a claim based on statistical calculations and results.	2.10.A Justify why a random variable is or is not a binomial random variable.	2.10.A.1 A binomial random variable, X, is a discrete random variable that counts the number of successes in repeated independent trials, n, that have only two possible outcomes (success or failure), with the probability of success p and the probability of failure 1– p.	Predict and Confirm Before performing a class simulation, have students discuss with their neighbor what values the distribution may take and how often they expect/predict those values to occur. For instance, what might the
	3.D Calculate	2.10.B Calculate the	2.10.B.1	

	means, standard deviations, and parameters for probability distributions.	mean and standard deviation for a binomial distribution.	If a random variable is binomial, its mean, μ_x , is np and its standard deviation, σ_x , is $\sqrt{np(1-p)}$.	distribution look like for a binomial setting with $p = 0.5$ and $n = 16$ when we perform 20 trials? Have students perform the simulation and compare the results to their predictions.
	4.D Interpret statistical calculations and results to assess meaning or a claim.	2.10.C Interpret the mean, standard deviation, and probabilities for a binomial distribution.	2.10.C.1 The mean, standard deviation, and probabilities for a binomial distribution should be interpreted in context.	
	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	2.10.D Estimate probabilities of binomial random variables using data from a simulation.	2.10.D.1 A probability distribution can be constructed using the rules of probability or estimated with a simulation.	
		2.10.E Calculate probabilities for a binomial distribution.	2.10.E.1 The probability that a binomial random variable, X , has exactly x successes for n independent trials, when the probability of success is p , is calculated as $P(X = x) = \binom{n}{x} p^x (1-p)^{n-x}, \quad x = 0, 1, 2, \dots, n$ This is called the binomial probability function.	
Day 11-12 2.11 <i>The Normal Distribution</i>	4.C Describe distributions and compare relative positions of points within a distribution.	2.11.A Describe a normal distribution.	2.11.A.1 A continuous random variable is a variable that can take on any value within a specified domain. Every interval within the domain has a probability associated with it. 2.11.A.2 Many continuous random variables are well modeled by a normal distribution. 2.11.A.3 A normal distribution can be described as a continuous, unimodal, bell-shaped, and symmetric curve. 2.11.A.4 A normal curve can be used to model a distribution of data and a continuous random variable. 2.11.A.5 The normal distribution, or the normal curve, is identified by two parameters, the mean, μ , and the standard deviation, σ . The smaller the standard deviation, the taller and more concentrated the normal curve is around its mean. The larger the standard deviation, the shorter and less concentrated the normal curve is around its mean.	Reversing Interpretations Give pairs of students four pictures of normal distributions with various parts shaded. Have students create the question that could have resulted in the picture shown. For example, if a value of 15 is labeled and the distribution is shaded to the right of 15, students could write, "What is the probability that a value is more than 15?"
	3.D Calculate means, standard deviations, and parameters for probability	2.11.B Calculate the mean and standard deviation for a normal distribution	2.11.B.1 A standard normal distribution is a normal distribution with mean $\mu = 0$ and standard deviation $\sigma = 1$.	

	distributions.		
	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	2.11.C Calculate percentages from a normal distribution using the empirical rule.	2.11.C.1 The empirical rule can be used to estimate the area of a region under the graph of the normal distribution curve. For a normal distribution, approximately 68% of observations are within 1 standard deviation of the mean, approximately 95% of observations are within 2 standard deviations of the mean, and approximately 99.7% of observations are within 3 standard deviations of the mean. This is called the empirical rule, or the 68–95–99.7 rule.
		2.11.D Calculate the probability that a particular value lies within a given interval of a normal distribution.	2.11.D.1 If the distribution of a random variable is approximately normal, the probability that the random variable takes on values within a particular interval of the random variable is determined by the area under the normal curve within that interval. The total probability or area under the normal curve is 1.
		2.11.E Calculate the associated intervals and areas of a normal distribution.	2.11.E.1 The boundaries of an interval associated with a given area in a normal distribution can be determined using technology or using z-scores and a standard normal table. 2.11.E.2 Intervals associated with a given area in a normal distribution can be determined by assigning appropriate inequalities to the boundaries of the intervals. To determine the intervals, p is defined as a number between 0 and 100, x_a is the lower bound, and x_b is the upper bound on a normal distribution. 2.11.E.2.i $P(X < x_a) = \frac{p}{100}$ means that the lowest p% of the values lie to the left of x_a. 2.11.E.2.ii $P(x_a < X < x_b) = \frac{p}{100}$ means that p% of the values lie between x_a and x_b. 2.11.E.2.iii $P(X > x_b) = \frac{p}{100}$ means that the highest p% of the values lie to the right of x_b. 2.11.E.2.iv To determine the most extreme p% of values on both sides requires dividing the area associated with p% into two equal areas on either extreme of the distribution: $P(X < x_a) = \frac{1}{2} \left(\frac{p}{100} \right) \text{ and } P(X > x_b) = \frac{1}{2} \left(\frac{p}{100} \right)$ mean that half of the p% most extreme values lie to the left of x_a and half of the p% most extreme values lie to the right of x_b.
	4.C Describe	2.11.F Compare	2.11.F.1 Percentiles and proportions may be used to

	distributions and compare relative positions of points within a distribution.	measures of relative position for distributions.	compare relative positions of individual values within a normal distribution or between normal distributions.	
Day 13 2.12 <i>Sampling Distributions and the Central Limit Theorem</i>	4.C Describe distributions and compare relative positions of points within a distribution	2.12.A Describe sampling distributions with simulations	2.12.A.1 A sampling distribution of a statistic is the distribution of values of the statistic for all possible samples of a given size from a given population. 2.12.A.2 The sampling distribution of a statistic can be simulated by repeatedly generating a large number of random samples from the population assuming known value(s) for the parameter(s). The value of the statistic is determined and recorded for each sample. The resulting distribution of the sample statistic values approximates the sampling distribution of the statistic. 2.12.A.3 A randomization distribution is the distribution of a statistic generated by simulation from repeatedly randomly reallocating, or reassigning, the response values to treatment groups. The value of the statistic is determined and recorded for each reallocation, or reassignment. The resulting distribution of the statistic values approximates the sampling distribution of the statistic. 2.12.A.4 The central limit theorem (CLT) states that the sampling distribution of a mean of a random sample has a shape that can be approximated by a normal distribution. The larger the sample is, the better the approximation will be.	
Day 14	Unit Test: AP Progress Check 2			

Unit Title		
Unit 3: Inference for Categorical Data: Proportions		
Relevant Standards: Bold indicates priority		
2.C, 2.D, 2.E, 3.C, 3.D, 3.E, 4.B, 4.C, 4.D, 4.E, 4.F, 4.G		
Essential Question(s)		Enduring Understanding(s)
- How can we measure changing opinions in populations over time? - When can we perform inference calculations using data from websites to answer everyday questions? - How can we determine whether the difference between two groups is statistically significant? - Are the distributions of various types of wild birds different in urban areas than in rural areas?		- Sample results vary naturally, and statistical inference accounts for this variability. - Confidence intervals provide plausible estimates of population proportions. - Hypothesis tests evaluate claims about population proportions using evidence from sample data. - Statistical conclusions should consider both statistical significance and practical importance.
Demonstration of Learning		Pacing for Unit
Progress Check 3 - Multiple-choice: ~65 questions - Free-response: 3 questions - Question 1: Inference - Question 2: Inference - Question 3: Inference The additional online question bank will be used to create a mid-unit quiz and other checks for understanding throughout the unit. College Board AP Exam Section I MCQ Weighting: 15–25% - Point Estimators & Sampling Distributions: Estimating population parameters using sample proportions. - Proportion Inference Procedures: Constructing and interpreting 1-proportion and 2-proportion z-confidence intervals and hypothesis tests. - Chi-Square Procedures: Conducting Chi-Square tests for homogeneity and independence to analyze categorical relationships across two-way tables. - Reasoning & Conditions: Verifying random sampling and Large Counts conditions, calculating test statistics/p-values, and drawing contextual conclusions. See AP Exam Overview for more information		18 Block Periods <i>NOTE:</i> <i>Midterm & AP Exam Review / Mock Exams - 11 Blocks</i> <i>Flexibility/Buffer/Post-AP Exam Project - 6 Blocks</i>
Family Overview (link below)		Integration of Technology:
AP Statistics FAMILY MATERIALS		- Construct confidence intervals for proportions using a graphing calculator or Desmos.. - Perform one- and two-sample z-tests for proportions using a graphing calculator or Desmos. - Interpret calculator and software output..
Unit-specific Vocabulary		Aligned Unit Materials, Resources, and Technology (beyond core resources)
Sampling distribution	Alternative hypothesis	www.mathmedic.com www.desmos.com www.stapplet.com https://skewthscript.org/
Statistic	P-value	
Parameter	Statistical significance	
Point estimate	Practical significance	

Margin of error	Type I error		
Confidence interval	Type II error		
Confidence level	Power		
Standard error	One-proportion z-interval		
Inference	One-proportion z-test		
Hypothesis test	Two-proportion z-interval		
Null hypothesis	Two-proportion z-test		
Opportunities for Interdisciplinary Connections:		Anticipated misconceptions	
- Political Science: election polling. - Public Health: vaccination and treatment studies. - Sociology: opinion surveys and demographic comparisons. - Psychology: behavioral research. - Journalism: interpretation of polling data.		- Misinterpreting a confidence interval as the probability that individual data values or a parameter fall into a specific calculated interval. - Stating that a high p-value proves or "accepts" the null hypothesis. - Mixing up when to use pooled versus unpooled standard errors in two-proportion inference procedures. - Believing that Type I error probability (α) and Type II error probability add up to 1. - Confusing Chi-Square tests of Homogeneity with Chi-Square tests of Independence.	
Connections to Prior Units		Connections to Future Units	
- Uses probability models and sampling variability from Unit 2. - Applies principles of data collection and study design from Unit 1.		Establishes the framework for all subsequent inference procedures. Introduces confidence intervals, hypothesis testing, and error analysis used again with means and regression.	
Differentiation through <i>Universal Design for Learning</i>			
	Engagement	Representation	Action & Expression
2.C	"Simulation Showdown" comparing physical card-shuffling re-randomization tests with computer applet simulations to evaluate experimental chance variation.	Flowchart decision tree mapping study design (Random Assignment vs. Random Selection) to the appropriate simulation-based inference technique (permutation test vs. bootstrap resampling).	Build an interactive flowchart or record an audio walkthrough explaining why a re-randomization simulation is chosen over a standard observational survey model for an experiment.
2.D	"False Discovery Lab" simulating drug trials where the treatment has zero effect to observe how often random assignment chance leads to a false positive result (Type I error).	Visual 2x2 matrix (True Effect vs. Study Decision) overlaid with levers showing how increasing sample size n or reducing assignment variation reduces false negative risks (Type II error).	Construct an infographic or record a video reflection explaining how small sample sizes in experimental designs increase the risk of Type II errors.
2.E	"By Chance Alone?" challenge formulating hypotheses for famous experiments (e.g., "Lady Tasting Tea") testing whether results could happen purely by random assignment chance.	Sentence-template guide mapping experimental claims to symbolic and verbal null/alternative forms (H_0 : Treatment has no effect / difference is due to chance; H_a : Treatment causes a genuine difference).	Write verbal and symbolic H_0 and H_a for three experimental scenarios and record an audio explanation defining the parameters/treatment effects in context.
3.C	"Card Shuffling / Applet Lab" performing 100+ simulated re-randomizations of experimental data and counting how many yield a treatment difference as extreme as observed.	Visual dotplot of a simulation distribution with the observed experimental result highlighted and the tail area shaded to display the estimated p-value.	Estimate p-values from simulation dotplots using choice of tool (applet, manual tally, or Python script) and submit annotated calculation logs.
3.E	"Difference in Means / Proportions Lab" calculating the observed test statistic from raw experimental group data prior to running simulations.	Worked example cards showing step-by-step subtraction of treatment group means/proportions alongside computer applet input fields.	Submit a calculation portfolio featuring both hand-calculated group statistic differences and applet-generated simulation test statistics.
4.E	"Scope of Inference Auditor" auditing study scenarios to verify whether Random Assignment was used	2x2 Scope of Inference Grid paired with a condition checklist (Random	Produce an "Experimental Design Audit Report" verifying whether study

	(allowing causal inference) and whether Random Selection was used (allowing population generalization).	Assignment, Independent Allocation, Sufficient Simulation Trials (≥ 1000).	conditions justify making a causal claim via simulation-based inference.
4.F	"Translate the Simulation" converting raw simulation p-values into plain-English statements about the probability of getting observed experimental differences by chance alone.	Fill-in-the-blank template for AP-style simulation interpretations: "Assuming [treatment has no effect], the probability of getting a difference in [statistic] of [value] or greater purely due to random assignment chance is [p-value]."	Record a video or podcast segment explaining what a simulation-based p-value means to an audience reviewing a clinical trial result.
4.G	"Jury Verdict: Causation vs. Chance" evaluating simulation p-values against a threshold ($\alpha = 0.05$) to declare whether an experimental treatment genuinely works or if results are inconclusive.	Decision tree linking p-value threshold to the final claim: "Because p-value is small, we have convincing evidence that [treatment causes effect]."	Write a complete CER justification essay or deliver a mock courtroom closing argument stating whether an experiment proves causation and specifying the target population.

Supporting Multilingual/English Learners (CELP standards)

	Emerging	Expanding	Bridging
2.C	Match categorical inference problem cards to method names using a visual flow chart.	Identify the correct categorical inference procedure from scenario prompts using a decision-tree graphic organizer.	Discriminate between 1-prop, 2-prop, χ^2 goodness-of-fit, χ^2 independence, and χ^2 homogeneity tests across ambiguous real-world research contexts.
2.D	Match Type I Error, Type II Error, and Power to visual scenarios using a decision truth-table card.	Describe Type I/II errors and Power in context using structured scenario frames	Analyze trade-offs between α , β , sample size n, and effect size, explaining relationships in formal written evaluations.
2.E	Select correct null and alternative hypothesis statements using parameter symbol cards.	Write H_0 and H_a hypotheses for proportions in symbols and words using structured sentence frames	Formulate parameter definitions and hypotheses for 1-proportion, 2-proportion, and χ^2 tests with non-directional or directional claims.
3.C	Calculate expected counts for χ^2 tests using the formula with color-coded table templates.	Calculate standard errors, margins of error, and confidence intervals for proportions using annotated formula guides.	Compute confidence intervals, test statistics, p-values, and χ^2 expected counts across single-sample and two-sample tasks.
3.E	Identify the test statistic and p-value from a calculator/software output screenshot using highlighter pens.	Calculate z-test statistics and corresponding p-values for proportions using structured step-by-step formula worksheets.	Independently execute full computational procedures for 1-prop, 2-prop, and χ^2 inference tests using technology and formula mechanics.
4.E	Check off verified conditions (Random, 10% condition, Large Counts) on a visual condition checklist.	Verify and write out conditions for proportion inference methods using structured paragraph stems.	Evaluate and justify inference conditions (including expected cell counts > 5 for χ^2 , addressing minor violations in context.
4.F	Interpret confidence intervals using a fill-in template: "We are [95]% confident that the true proportion of [context] is between [L] and [U]."	Interpret confidence intervals and p-values in context using complete sentence starters	Interpret confidence levels, p-values, and χ^2 components in context, explaining the practical meaning of statistical significance.
4.G	Match p-value decisions (reject or fail to reject) using visual rule cards.	Write formal conclusions for hypothesis tests using decision templates	Construct rigorous, contextualized statistical arguments justifying claims about population proportions based on confidence intervals or hypothesis test results.

Unit Outline

Lesson Sequence	Skills	Learning Objective	Essential Knowledge	AP Activities (if available)
Day 1 3.1 Estimators	4.B Justify a claim based on statistical	3.1.A Justify why an estimator is or is not	3.1.A.1	

	calculations and results.	unbiased.	When estimating a population parameter, an estimator is unbiased if, on average, the value of the estimator does not underestimate or overestimate the population parameter.
	3.D Calculate means, standard deviations, and parameters for probability distributions.	3.1.B Calculate estimates for a population parameter.	3.1.B.1 A sample statistic is a point estimator of the corresponding population parameter and can be thought of as the estimate of the population parameter. For example, the sample proportion p is a point estimator for the population proportion p .
Day 2 3.2 Sampling Distributions for Sample Proportions	3.D Calculate means, standard deviations, and parameters for probability distributions.	3.2.A Calculate the mean and standard deviation of a sampling distribution for a sample proportion.	3.2.A.1 For a population with population proportion p , when the sampled values are independent, the sampling distribution of a sample proportion $\hat{p}$ has a mean $\mu_{\hat{p}} = p$ and a standard deviation $\sigma_{\hat{p}} = \sqrt{\frac{p(1-p)}{n}}$.
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.2.B Justify the appropriateness of conditions for the sampling distribution of a sample proportion.	3.2.B.1 Sampling without replacement requires that two conditions be met: 3.2.B.1.i The randomization condition—the data should be collected using a random sample. 3.2.B.1.ii The 10% condition—the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size. 3.2.B.2 The sampling distribution of the sample proportion p is approximately normal provided the sample size is large enough. To ensure the sample size is large enough, the following condition must be met: $np \geq 10$ and $n(1-p) \geq 10$, where np is the expected number of successes and $n(1-p)$ is the expected number of failures.
	4.D Interpret statistical calculations and results to assess meaning or a claim.	3.2.C Interpret the mean, standard deviation, and probabilities for a sampling distribution of a sample proportion.	3.2.C.1 The mean, standard deviation, and probabilities for a sampling distribution of a sample proportion should be interpreted in the context of a specific population.
Day 3 3.3 Constructing a Confidence Interval for a Population Proportion	2.C Identify appropriate statistical inference methods.	3.3.A Identify an appropriate confidence interval procedure including the parameter for a population proportion.	3.3.A.1 A confidence interval is an interval estimate for a population parameter. Based on the sample proportion, a confidence interval can be calculated to estimate the value of a single population proportion. The appropriate confidence interval procedure is a one-sample z -interval for a population proportion. 3.3.A.2 The parameter for a confidence interval for a population proportion should reference the proportion, the response variable, and the population in context.
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.3.B Justify the appropriateness of constructing a confidence interval for a population proportion by verifying conditions.	3.3.B.1 A one-sample z -interval for a population proportion requires that three conditions be met: 3.3.B.1.i The randomization condition—the data should be collected using a random sample. 3.3.B.1.ii The 10% condition—when sampling without replacement, the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size.

			3.3.B.1.iii The normality condition—the observed number of successes, $n\hat{p}$, and the observed number of failures, $n(1-\hat{p})$, should be at least 10.	
3.E Calculate appropriate statistical inference method results.	3.3.C Calculate an appropriate confidence interval for a population proportion.	3.3.C.1 z^* denotes a critical value, such that $-z^*$ and $+z^*$ represent the boundaries enclosing the middle C% of the standard normal distribution, in which C% is an approximate confidence level with which the population proportion is estimated. 3.3.C.2 An interval estimate can be constructed as point estimate (+/- margin of error). For a population proportion, the one-sample z-interval estimate is $\hat{p} \pm z^* \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}$.		
	3.3.D Calculate the standard error and margin of error of a sample statistic for a confidence interval for a population proportion, and estimate a given sample size from the margin of error	3.3.D.1 The standard error (SE) of a statistic is an estimate of the standard deviation of the sampling distribution of the statistic. The standard error of the sample proportion $\hat{p}$ is $SE_{\hat{p}} = \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}$. 3.3.D.2 The standard error quantifies the typical amount that a statistic will vary from the value of the corresponding population parameter. 3.3.D.3 The margin of error of $\hat{p}$ is half the width of the confidence interval and is calculated as the critical value (z^*) times the standard error (SE) of $\hat{p}$, which equals $z^* \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}$. 3.3.D.4 The formula for the margin of error (MOE) can be rearranged to $n = \frac{(z^*)^2 (\hat{p})(1-\hat{p})}{(MOE)^2}$, solve for n, the minimum sample needed to achieve a given margin of error. For this purpose, if $\hat{p}$ is not defined or unable to be calculated, use $\hat{p} = 0.5$ in order to find the upper bound for the sample size that will result in a given margin of error.		
Day 4 3.4 <i>Justifying a Claim Based on a Confidence Interval for a Population Proportion</i>	4.F Interpret results of statistical inference methods.	3.4.A Interpret a confidence interval in context for a population proportion.	3.4.A.1 Because the confidence interval for a population proportion is calculated based on a sample from a population, the computed interval may or may not contain the value of the population proportion. 3.4.A.2 The interpretation of the confidence level is that in repeated random sampling with the same sample size, approximately C% of confidence intervals calculated will capture the population proportion, with C representing the numerical value of the confidence level used. 3.4.A.3 When interpreting a C% confidence interval for a population proportion, we say we are C% confident that the interval (a, b) contains the true value of the parameter for the population, where a	The Scribe and the Calculator Have students work with a partner to construct and interpret a confidence interval for a population proportion. Only one partner is allowed to use the calculator, and only the other partner is allowed to write. When a calculation needs to be made, only the scribe can describe to the calculator operator what buttons to push; when writing needs to be done, only the calculator operator can describe to the scribe what needs to

			represents the lower limit and b represents the upper limit. An interpretation of a confidence interval for a population proportion includes a reference to the parameter with details about the population it represents in the context of the study.	be written. Have students switch roles when constructing and interpreting a confidence interval for the difference of two population proportions.
	4.G Justify a claim based on statistical inference method results.	3.4.B Justify a claim based on a confidence interval for a population proportion.	3.4.B.1 A confidence interval for a population proportion provides a range of plausible values that may serve as convincing evidence to support a particular claim about the population proportion.	
	2.D Identify types of errors and relationships among components in statistical inference methods.	3.4.C Identify the relationships among sample size, confidence interval width, confidence level, and margin of error for a population proportion.	3.4.C.1 For a given sample, increasing the confidence level will result in the following: 3.4.C.1.i The critical value will increase. 3.4.C.1.ii The margin of error will increase. 3.4.C.1.iii The width of the confidence interval will increase. 3.4.C.2 Increasing the sample size decreases the standard error. Thus, when all other things remain the same, the width of the confidence interval for a population proportion tends to decrease as the sample size increases. For a confidence interval for a population proportion with a given confidence level, the width of the interval is approximately proportional to $\frac{1}{\sqrt{n}}$.	
Day 5 3.5 <i>Setting Up a Test for a Population Proportion</i>	2.C Identify appropriate statistical inference methods.	3.5.A Identify an appropriate testing method for a population proportion including the parameter for the population proportion.	3.5.A.1 A hypothesis test is a statistical inference procedure that is used to make a decision about the value of a population parameter. The appropriate hypothesis testing procedure is a one-sample z-test for a population proportion. 3.5.A.2 The parameter for a hypothesis test for a population proportion should reference the population parameter, the response variable, and the population in context.	Error Analysis Give student pairs a worksheet with 20 sets of hypotheses, each with a common student mistake. Be sure to include hypotheses for a test for a population proportion, a test for the difference between two proportions, a chi-square test for independence, and a chi-square test for homogeneity. Have students circle the incorrect part, write why the circled component is incorrect, and then write the correct hypotheses. Include errors such as using statistics instead of parameters and interchanging the = and > in the two hypotheses.
	2.E Identify the null and alternative hypotheses.	3.5.B Identify the null and alternative hypotheses for a population proportion.	3.5.B.1 In the hypothesis testing procedure, the null hypothesis, H_0 , is the statement about a parameter that is assumed to be correct unless there is convincing statistical evidence suggesting otherwise. It is the status quo condition. The alternative hypothesis, H_a , is the claim or belief about a parameter for which evidence is being collected. A researcher's claim or belief about the population parameter is represented by the alternative hypothesis. 3.5.B.2	

			The null hypothesis contains an equality reference ($=$, $\geq$, or $\leq$). Although the null hypothesis for a one-sided test may include an inequality symbol, in AP Statistics it is tested at the boundary of equality. The alternative hypothesis with $<$ or $>$ is called one-sided, and the alternative hypothesis with $\neq$ is called two-sided. 3.5.B.3 The null hypothesis for a one-sample z-test for a population proportion is as follows: $H_0: p = p_0$, where p_0 is the null hypothesized value for the population proportion. A one-sided alternative hypothesis for a one-sample z-test for a population proportion is either $H_a: p < p_0$ or $H_a: p > p_0$. A two-sided alternative hypothesis is $H_a: p \neq p_0$.	
	4.E Justify the use of a chosen statistical inference method by verifying conditions	3.5.C Justify the appropriateness of a hypothesis test for a population proportion by verifying conditions.	3.5.C.1 A one-sample z-test for a population proportion requires that three conditions be met: 3.5.C.1.i The randomization condition—the data should be collected using a random sample. 3.5.C.1.ii The 10% condition—when sampling without replacement, the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size. 3.5.C.1.iii The normality condition—the expected number of successes, np_0, and the expected number of failures, $n(1 - p_0)$, should be at least 10.	
Day 6 3.6 <i>p</i> -Values	4.F Interpret results of statistical inference methods.	3.6.A Interpret the p-value of a hypothesis test for a population proportion	3.6.A.1 Given the null hypothesis is true, there is a probability distribution of the test statistic called the null distribution. Using the null distribution, the p-value is the probability of obtaining a test statistic as extreme or more extreme (i.e., in the direction of the alternative hypothesis) than the test statistic that is observed given that the null hypothesis is true. That is, when x is the test statistic, the p-value is determined by finding the following: 3.6.A.1.i The probability at or above the observed value of the test statistic ($P(z \geq x)$), if the alternative is $>$ 3.6.A.1.ii The probability at or below the observed value of the test statistic ($P(z \leq x)$), if the alternative is $<$ 3.6.A.1.iii The probability less than or equal to the negative of the absolute value of the test statistic plus the probability greater than or equal to the absolute value of the test statistic,	

			$\left(P\left(z \leq - x \right)\right)+\left(P\left(z \geq x \right)\right)$, if the alternative is $\neq$ 3.6.A.2 If the distribution of the test statistic has been simulated, the p-value is the proportion of values in the null distribution that are as extreme or more extreme than the observed value of the test statistic. This is as follows: 3.6.A.2.i The proportion at or above the observed value of the test statistic, if the alternative is $>$ 3.6.A.2.ii The proportion at or below the observed value of the test statistic, if the alternative is $<$ 3.6.A.2.iii The proportion less than or equal to the negative of the absolute value of the test statistic plus the proportion greater than or equal to the absolute value of the test statistic, if the alternative is $\neq$ 3.6.A.3 An interpretation of the p-value of a hypothesis test for a population proportion should include a statement that the p-value is computed by assuming the null hypothesis is true (i.e., by assuming the true population proportion is equal to the particular value stated in the null hypothesis in context). 3.6.A.4 Small p-values indicate that the observed value of the test statistic would be unusual if the null hypothesis were true and therefore provide evidence for the alternative hypothesis. The lower the p-value, the more convincing the statistical evidence for the alternative hypothesis. 3.6.A.5 p-values that are not small indicate that the observed value of the test statistic would not be unusual if the null hypothesis were true and therefore do not provide convincing statistical evidence for the alternative hypothesis, nor do they provide evidence that the null hypothesis is true.	
Day 7 3.7 Carrying Out a Test for a Population Proportion	3.E Calculate appropriate statistical inference method results.	3.7.A Calculate an appropriate test statistic and p-value for testing a hypothesis about a population proportion	3.7.A.1 The test statistic for testing a population proportion is $z=\frac{\hat{p}-p_0}{\sqrt{\frac{p_0\left(1-p_0\right)}{n}}}$. The z-statistic has a standard normal distribution when the null hypothesis is true. 3.7.A.2 The distribution of the test statistic assuming the null hypothesis is true (null distribution) can be approximated by a probability model (e.g., a theoretical distribution such as the standard normal distribution). 3.7.A.3 The p-value of a one-sample z-test for a population proportion is found from the standard normal distribution using a table or technology.	Sentence Starters For a given question, provide students with a set of hypotheses, p-value, significance level, and context. Have them compare the p-value to the significance level to determine whether to reject the null hypothesis. Use a sentence starter with blanks to fill in, and have students write a sentence in context explaining whether they have enough evidence to reject H0 or if they will fail to reject H0. Make sure students avoid the common mistake of implying that evidence supports an accept H0 conclusion or a reject H0 conclusion.
	4.G Justify a claim based on statistical inference method results.	3.7.B Justify a claim about the population based on the results of a hypothesis test for a population proportion.	3.7.B.1 The significance level of a hypothesis test, denoted by α , is the predetermined probability of rejecting the null hypothesis given that it is true. The significance level may be given or determined by the researcher. The relationship between a p-value and the significance level of a hypothesis test	

			determines whether a result is statistically significant. 3.7.B.2 A formal decision in a hypothesis test explicitly compares the p-value to the significance level, α. If the p-value $\leq \alpha$, then reject the null hypothesis, $H_0: p = p_0$. If the p-value $> \alpha$, then fail to reject the null hypothesis. 3.7.B.3 Rejecting the null hypothesis means there is convincing statistical evidence to support the alternative hypothesis. Failing to reject the null hypothesis means there is not convincing statistical evidence to support the alternative hypothesis. 3.7.B.4 A hypothesis test can lead to rejecting or not rejecting the null hypothesis but can never lead to concluding or proving that the null hypothesis is true. Lack of statistical evidence for the alternative hypothesis is not the same as evidence for the null hypothesis. 3.7.B.5 The results of a hypothesis test for a population proportion can serve as the statistical reasoning to support the answer to an investigative question about the population that was sampled. 3.7.B.6 A conclusion for the hypothesis test for a population proportion is stated in context consistent with, and in terms of, the alternative hypothesis using non-definitive language. The conclusion should contain a reference to the parameter and the population.	
Day 8	Mid-Unit Assessment			
Day 9 3.8 <i>Potential Errors When Performing Tests</i>	2.D Identify types of errors and relationships among components in statistical inference methods.	3.8.A Identify Type I and Type II errors.	3.8.A.1 A Type I error occurs when there is convincing statistical evidence that the alternative hypothesis is true (due to the small p-value), but it is not. 3.8.A.2 A Type II error occurs when there is not convincing statistical evidence that the alternative hypothesis is true (due to the large p-value), but it is. 3.8.A.3 The power of a hypothesis test is the probability that a hypothesis test will correctly reject the false null hypothesis.	
	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	3.8.B Calculate the probability of Type I and Type II errors.	3.8.B.1 The probability of making a Type I error is defined as the significance level, α . For a given study and hypothesis test, the probability of making a Type I error is typically set to a small value (e.g., 0.01, 0.05, 0.10) prior to collecting the data. 3.8.B.2 The probability of making a Type II error is 1 - power.	
	2.D Identify types of errors and relationships among	3.8.C Identify the factors that affect the probability of errors	3.8.C.1 For a given study and hypothesis test, the probability of a Type II error should ideally be small, and thus, the power will be large (e.g., $P(\text{Type II error}) = 0.20$ and power = 0.80). The probability of a Type	

	components in statistical inference methods.	in hypothesis testing.	If error decreases and the power increases when any one of the following occurs, provided the others do not change: 3.8.C.1.i Sample size(s) increases. 3.8.C.1.ii Standard error decreases. 3.8.C.1.iii True parameter value is farther from the null hypothesis. 3.8.C.1.iv Significance level (α) of a test increases.
	4.D Interpret statistical calculations and results to assess meaning or a claim.	3.8.D Interpret Type I and Type II errors.	3.8.D.1 In some studies, making a Type I error may have more serious consequences than making a Type II error. In other studies, making a Type II error may have more serious consequences than making a Type I error. The consequences of each error should be considered prior to conducting the study. 3.8.D.2 Because the significance level, α, is the probability of making a Type I error, the consequences of a Type I error influence decisions about a significance level. 3.8.D.3 Because sample size influences the probability of making a Type II error, the consequences of a Type II error influence decisions about how large the sample size should be.
Day 10 3.9 <i>Sampling Distributions for the Difference Between Sample Proportions</i>	3.D Calculate means, standard deviations, and parameters for probability distributions.	3.9.A Calculate the mean and standard deviation of the sampling distribution for the difference between two sample proportions.	3.9.A.1 For two independent populations, with population proportions p_1 and p_2, when the sampled values are independent, the sampling distribution for the difference in sample proportions, $\hat{p}_1 - \hat{p}_2$, has a mean, $\mu_{\hat{p}_1 - \hat{p}_2} = p_1 - p_2$ and standard deviation, $\sigma_{\hat{p}_1 - \hat{p}_2} = \sqrt{\frac{p_1(1-p_1)}{n_1} + \frac{p_2(1-p_2)}{n_2}}$
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.9.B Justify the appropriateness of conditions for the sampling distribution for the difference between two sample proportions.	3.9.B.1 When sampling without replacement, two conditions must be met: 3.9.B.1.i The randomization condition—the data should be collected using two independent random samples. 3.9.B.1.ii The 10% condition—the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N$ and $n_2 \leq 10\%N_2$, where N is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2. 3.9.B.2 If the data come from an experiment, the data only need to meet the randomization condition. The treatments must be randomly assigned to the experimental units to meet the randomization condition. 3.9.B.3 The sampling distribution for the difference between sample proportions, $\hat{p}_1 - \hat{p}_2$, will have an approximately normal distribution provided both sample sizes are large enough. To ensure that both samples are large enough, the data must meet the following conditions: $n_1 p_1 \geq 10$, $n_1(1-p_1) \geq 10$, $n_2 p_2 \geq 10$, and $n_2(1-p_2) \geq 10$ are the expected number of successes and $n_1(1-p_1)$ and $n_2(1-p_2)$ are the expected number of failures.

	4.D Interpret statistical calculations and results to assess meaning or a claim.	3.9.C Interpret the mean, standard deviation, and probabilities for the sampling distribution for the difference between two sample proportions.	3.9.C.1 The mean, standard deviation, and probabilities for the sampling distribution for the difference between two sample proportions should be interpreted within the context of two specific populations.
Day 11-12 3.10 <i>Constructing a Confidence Interval for the Difference Between Two Population Proportions</i>	2.C Identify appropriate statistical inference methods.	3.10.A Identify an appropriate confidence interval procedure including the parameters for the difference between two population proportions.	3.10.A.1 Based on the sample data, a confidence interval can be calculated to estimate the difference between two population proportions. The appropriate confidence interval procedure is a two-sample z-interval for a difference between population proportions. 3.10.A.2 The parameters of a confidence interval for the difference between two population proportions should refer to the difference in the proportions, the response variable, and the populations in context.
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.10.B Justify the appropriateness of constructing a confidence interval for the difference between two population proportions by verifying conditions.	3.10.B.1 A two-sample z-interval for a difference between two population proportions requires that three conditions be met: 3.10.B.1.i The randomization condition—the data should be collected using two independent random samples or a randomized experiment. 3.10.B.1.ii The 10% condition—when sampling without replacement, the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N_1$ and $n_2 \leq 10\%N_2$, where N_1 is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2 . (Note: This condition is unnecessary when the data are from a randomized experiment.) 3.10.B.1.iii The normality condition—the number of observed successes, $n_1\hat{p}_1$, and $n_2\hat{p}_2$, and number of observed failures, $n_1(1-\hat{p}_1)$ and $n_2(1-\hat{p}_2)$, for both samples are all at least 10.
	3.E Calculate appropriate statistical inference method results.	3.10.C Calculate an appropriate confidence interval for the difference between two population proportions.	3.10.C.1 The point estimate for the difference between two population proportions is $(\hat{p}_1 - \hat{p}_2)$. 3.10.C.2 For the difference between two population proportions, the interval estimate can be constructed as point estimate +/- (margin of error). The interval estimate for the difference between two population proportions is $(\hat{p}_1 - \hat{p}_2) \pm z^* \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}}$
		3.10.D Calculate the standard error and margin of error for estimating the difference between two population proportions.	3.10.D.1 The standard error (SE) for the difference between two population proportions is $SE_{\hat{p}_1 - \hat{p}_2} = \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}}$. 3.10.D.2 For the difference between two population proportions, the margin of error is the critical value (z^*) times the standard error (SE) of the difference between the two proportions, which equals

			$z^* \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}}$	
Day 13 3.11 <i>Justifying a Claim Based on a Confidence Interval for the Difference Between Two Population Proportions</i>	4.F Interpret results of statistical inference methods.	3.11.A Interpret a confidence interval in context for the difference between two population proportions.	3.11.A.1 Because the confidence interval for the difference between two population proportions is calculated based on samples from two populations, the computed interval may or may not contain the value for the difference between those two population proportions. 3.11.A.2 The interpretation of the confidence level is as follows: In repeated random sampling with the same sample sizes from the same populations, approximately C% of confidence intervals created will capture the difference between the two population proportions, where C represents the numerical value of the confidence level used. 3.11.A.3 When interpreting a C% confidence interval for the difference between two population proportions, we say we are C% confident that the interval (a, b) , contains the parameter for the difference between the populations, where a represents the lower limit and b represents the upper limit. An interpretation of a confidence interval for the difference between two population proportions includes a reference to the parameter with the details about the populations it represents in the context of the study.	The Scribe and the Calculator (see Topic 3.4)
	4.G Justify a claim based on statistical inference method results.	3.11.B Justify a claim based on a confidence interval for the difference between two population proportions.	3.11.B.1 A confidence interval for the difference between two population proportions provides an interval of values that may provide convincing evidence to support a particular claim about the difference between the two population proportions. For example, if the interval contains 0, then there is insufficient evidence to conclude there is a difference between the two population proportions. If the interval does not contain 0, there is sufficient evidence to conclude there is a difference between the two population proportions	
Day 14 3.12 <i>Setting Up a Test for the Difference Between Two Population Proportions</i>	2.C Identify appropriate statistical inference methods.	3.12.A Identify an appropriate testing method for the difference between population proportions including the parameters.	3.12.A.1 The appropriate testing method for the difference between two population proportions is a two-sample z-test for the difference between two population proportions. 3.12.A.2 The parameters for a hypothesis test for the difference between two population	Error Analysis (see Topic 3.5)

			proportions should reference the population parameters, the response variables, and the populations in context.	
	2.E Identify the null and alternative hypotheses.	3.12.B Identify the null and alternative hypotheses for the difference between population proportions.	3.12.B.1 For a two-sample z-test for the difference between two population proportions, the null hypothesis indicates no difference. The null hypothesis for the difference between two population proportions can be written as either $H_0: p_1 = p_2 \text{ or } H_0: p_1 - p_2 = 0$ A one-sided alternative hypothesis for the difference between two population proportions can be written as either $H_a: p_1 < p_2 \text{ or equivalently } H_a: p_1 - p_2 < 0, \text{ or } H_a: p_1 > p_2$ or equivalently $H_a: p_1 - p_2 > 0$. A two-sided alternative hypothesis for the difference between two population proportions can be written as either $H_a: p_1 \neq p_2 \text{ or equivalently } H_a: p_1 - p_2 \neq 0$	
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.12.C Justify the appropriateness of a hypothesis test for the difference between two population proportions by verifying conditions.	3.12.C.1 A two-sample z-test for a difference between two population proportions requires that three conditions be met: 3.12.C.1.i The randomization condition—the data should be collected using two independent random samples or a randomized experiment. 3.12.C.1.ii The 10% condition—when sampling without replacement, the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N_1 \text{ and } n_2 \leq 10\%N_2$, where N_1 is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2. (Note: This condition is unnecessary when the data are from a randomized experiment.) 3.12.C.1.iii The normality condition— $n_1 \hat{p}_c, n_1 (1 - \hat{p}_c), n_2 \hat{p}_c, \text{ and } n_2 (1 - \hat{p}_c)$, must all be at least 10, $\hat{p}_c = \frac{n_1 \hat{p}_1 + n_2 \hat{p}_2}{n_1 + n_2}$ with $\hat{p}_c$ being the combined (or pooled) proportion	

			assuming that H_0 is true $(H_0: p_1 = p_2 \text{ or } H_0: p_1 - p_2 = 0)$.	
Day 15 3.13 <i>Carrying Out a Test for the Difference Between Two Population Proportions</i>	3.E Calculate appropriate statistical inference method results.	3.13.A Calculate an appropriate test statistic and p-value for testing a hypothesis for the difference between two population proportions.	3.13.A.1 The test statistic for the difference between two population proportions is $z = \frac{(\hat{p}_1 - \hat{p}_2) - 0}{\sqrt{\hat{p}_c(1 - \hat{p}_c)} \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$, where $\hat{p}_c = \frac{n_1 \hat{p}_1 + n_2 \hat{p}_2}{n_1 + n_2}$ is the proportion of successes for the two groups combined. The z-statistic has a standard normal distribution when the null hypothesis is true. 3.13.A.2 The p-value for a two-sample z-test for the difference between two population proportions can be found from the standard normal distribution using a table or technology.	Sentence Starters (see Topic 3.7)
	4.F Interpret results of statistical inference methods.	3.13.B Interpret the p-value of a hypothesis test for the difference between two population proportions.	3.13.B.1 The p-value is the probability of obtaining a test statistic as extreme or more extreme than the test statistic that was observed (i.e., in the direction of the alternative hypothesis) given that the null hypothesis is true. An interpretation of the p-value of a hypothesis test for a difference between two population proportions should include a statement that the p-value is computed assuming the null hypothesis is true (i.e., by assuming that the true population proportions are equal to each other in context).	
	4.G Justify a claim based on statistical inference method results.	3.13.C Justify a claim about the populations based on the results of a hypothesis test for the difference between two population proportions.	3.13.C.1 A formal decision in a hypothesis test for the difference between two population proportions explicitly compares the p-value to the significance level, α . If the p-value $\leq \alpha$ then reject the null hypothesis, $H_0: p_1 = p_2$ or $H_0: p_1 - p_2 = 0$. If the p-value $> \alpha$, then fail to reject the null hypothesis. 3.13.C.2 The results of a hypothesis test for the difference between two population proportions can serve as the statistical reasoning to support the answer to an investigative question about the two populations that were sampled. 3.13.C.3 A conclusion for the hypothesis test for the difference between two population	

			proportions is stated in context consistent with, and in terms of, the alternative hypothesis using nondefinitive language. The conclusion should contain a reference to the parameters and the populations.	
Day 16 3.14 <i>Setting Up a Chi-Square Test for Homogeneity or Independence</i>	4.C Describe distributions and compare relative positions of points within a distribution.	3.14.A Describe chi-square distributions.	3.14.A.1 The chi-square statistic measures the distance between observed and expected counts relative to expected counts. 3.14.A.2 Chi-square distributions have positive values and are skewed right. Within this family of density curves, the skew becomes less pronounced with increasing degrees of freedom.	Error Analysis (see Topic 3.5)
	2.C Identify appropriate statistical inference methods.	3.14.B Identify an appropriate testing method for comparing distributions in two-way tables of categorical data including the populations and variables	3.14.B.1 To determine whether the distributions of a categorical variable for two or more populations are different, the appropriate test is the chi-square test for homogeneity. 3.14.B.2 A chi-square test for homogeneity should reference the categorical variable and the populations in context. 3.14.B.3 To determine whether row and column variables in a two-way table of categorical data might be associated in the single population from which the data were sampled, the appropriate test is the chi-square test for independence. 3.14.B.4 A chi-square test for independence should reference the categorical variables and the population in context.	
	2.E Identify the null and alternative hypotheses.	3.14.C Identify the null and alternative hypotheses for a chi-square test for homogeneity or independence.	3.14.C.1 The appropriate null hypothesis for a chi-square test for homogeneity is H_0 : there is no difference in the distributions of the categorical variable across populations or treatments. The appropriate alternative hypothesis for a chi-square test for homogeneity is H_a : there is a difference in the distributions of the categorical variable across populations or treatments. 3.14.C.2 The appropriate null hypothesis for a chi-square test for independence is H_0 : there is no association between two categorical variables in a given population or the two categorical variables in a given population are independent of each other. The appropriate alternative hypothesis for a chi-square test for independence is H_a : there is an association between two categorical variables in a given population or the two categorical	

			variables in a given population are not independent of each other.	
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	3.14.D Justify the appropriateness of a chi-square test for independence or homogeneity by verifying conditions.	3.14.D.1 A chi-square test for homogeneity or independence requires that three conditions must be met: 3.14.D.1.i The randomization condition—the test of independence states that the data should be collected using a random sample. The test for homogeneity states that the data should be collected using independent random samples or a randomized experiment. 3.14.D.1.ii The 10% condition—when sampling without replacement, check that $n \leq 10\%N$, where N is the size of the population and n is the sample size. (Note: This condition is unnecessary when the data are from a randomized experiment.) 3.14.D.1.iii The expected counts condition—all expected counts should be greater than 5.	
Day 17 3.15 <i>Carrying Out a Chi-Square Test for Homogeneity or Independence</i>	3.C Calculate and estimate expected counts, percentages, probabilities, and intervals.	3.15.A Calculate expected counts for two-way tables of categorical data	3.15.A.1 The expected counts (under the null hypothesis) in a particular cell of a two-way table of categorical data can be calculated using the formula $\text{expected count} = \frac{(\text{row total})(\text{column total})}{\text{total table}}$	Sentence Starters (see Topic 3.7)
	3.E Calculate appropriate statistical inference method results.	3.15.B Calculate the appropriate test statistic and p-value for a chisquare test for homogeneity or independence.	3.15.B.1 The appropriate test statistic for a chi-square test for homogeneity or independence is the chi-square statistic $\chi^2 = \sum \frac{(\text{Observed Count} - \text{Expected Count})^2}{\text{Expected Count}}$ where the sum is taken over all cells of the two-way table. The chi-square statistics have a chi-square distribution with degrees of freedom equal to (number of rows - 1)(number of columns - 1) when the null hypothesis is true. 3.15.B.2 The p-value for a chi-square test for independence or homogeneity is found from a chi-square distribution using a table or technology.	
	4.F Interpret results of statistical inference methods.	3.15.C Interpret the p-value for the chi-square test for homogeneity or independence.	3.15.C.1 The p-value is the probability of obtaining a test statistic as extreme or more extreme than the test statistic that was observed (i.e., in the direction of the alternative hypothesis) given that the null hypothesis is true. An interpretation of the p-value for the chi-square test for homogeneity or independence should include a statement that the p-value is computed by assuming the null hypothesis is true in context.	

	4.G Justify a claim based on statistical inference method results.	3.15.D Justify a claim about the population(s) based on the results of a chi-square test for homogeneity or independence.	3.15.D.1 A formal decision in a hypothesis test explicitly compares the p-value to the significance level, α. If the p-value $\leq \alpha$, then reject the null hypothesis for the appropriate chi-square test. If the p-value $> \alpha$, then fail to reject the null hypothesis. 3.15.D.2 The results of a chi-square test for homogeneity or independence can serve as the statistical reasoning to support the answer to an investigative question about the population that was sampled (independence) or the populations that were sampled (homogeneity). 3.15.D.3 A conclusion for a chi-square test for homogeneity or independence is stated in context consistent with, and in terms of, the alternative hypothesis using non-definitive language. The conclusion should contain a reference to the population(s).	
Day 18	Unit Test: AP Progress Check 3			

Unit Title		
Unit 4: Inference for Quantitative: Means		
Relevant Standards: Bold indicates priority		
2.C, 2.D, 2.E, 3.D, 3.E, 4.C, 4.D, 4.E, 4.F, 4.G		
Essential Question(s)		Enduring Understanding(s)
- How can we compare the average salaries of different professions across different regions of a country? - How can we evaluate the effectiveness of a new medicine in lowering cholesterol for older adults? - How can a pre-test and a post-test be used to evaluate student learning?		- Sampling distributions provide the foundation for making inferences about population means. - Confidence intervals estimate population means while quantifying uncertainty. - Hypothesis tests provide evidence for evaluating claims about population means. - The validity of statistical conclusions depends on appropriate assumptions, conditions, and context.
Demonstration of Learning		Pacing for Unit
Progress Check 4 - Multiple-choice: ~49 questions - Free-response: 3 questions - Question 1: Inference - Question 2: Inference - Question 3: Inference The additional online question bank will be used to create a mid-unit quiz and other checks for understanding throughout the unit. College Board AP Exam Section I MCQ Weighting: 10–20% - t-Procedures for Population Means: Performing 1-sample t-intervals/tests, matched-pairs t-procedures, and 2-sample t-intervals/tests. - Condition Verification: Checking independence and the Normal/Large Sample condition ($n \leq 30$ or checking sample plots for severe skewness/outliers) when population standard deviation is unknown. - Interpretation: Interpreting confidence levels, confidence intervals, p-values, and Type I / Type II errors in context See AP Exam Overview for more information		12 Block Periods <i>NOTE:</i> Midterm & AP Exam Review / Mock Exams - 11 Blocks Flexibility/Buffer/Post-AP Exam Project - 6 Blocks
Family Overview (link below)		Integration of Technology:
AP Statistics FAMILY MATERIALS		- Construct t-intervals and t-tests using calculators and statistical software. - Interpret technology-generated output. - Explore sampling distributions through simulation. - Analyze real-world quantitative datasets using technology.
Unit-specific Vocabulary		Aligned Unit Materials, Resources, and Technology (beyond core resources)
Mean	Two-sample t-interval	www.mathmedic.com www.desmos.com www.stapplet.com https://skewthescrypt.org/
Population mean	Two-sample t-test	
Sample mean	Confidence interval	
Sampling distribution	Hypothesis test	
Standard error	P-value	
Degrees of freedom	Type I error	
t-distribution	Type II error	
One-sample t-interval	Power	

One-sample t-test Matched-pairs t-test	Assumptions Conditions		
Opportunities for Interdisciplinary Connections:		Anticipated misconceptions	
• Biology: treatment and control group comparisons. • Medicine: clinical trial outcomes. • Environmental Science: comparing environmental measurements. • Economics: income and spending analyses. • Physical Education and Sports Science: performance comparisons.		1. Analyzing paired data with a two-sample independent t-procedure instead of a matched-pairs t-procedure. 2. Believing t-procedures require raw sample data to be strictly Normally distributed regardless of sample size. 3. Confusing sample standard deviation (s) with standard error (SE). 4. Using z-critical values instead of t-critical values when analyzing sample means with an unknown population standard deviation. 5. Equating statistical significance (low p-value) with real-world practical significance.	
Connections to Prior Units		Connections to Future Units	
• Extends inference concepts developed in Unit 3. • Relies on probability and sampling distributions from Unit 2. • Uses study-design principles introduced in Unit 1.		• Prepares students for regression-based inference. • Reinforces interpretation of confidence intervals, hypothesis tests, and statistical conclusions.	
Differentiation through <i>Universal Design for Learning</i>			
	Engagement	Representation	Action & Expression
2.C	"Casino & Game Show Audit" evaluating game rules to decide whether outcomes are governed by Discrete, Binomial B(n,p), Geometric G(p), or Continuous Normal probability decision models.	Flowchart decision tree mapping random variable features (fixed trial count n vs. count until 1st success, discrete counts vs. continuous measurements) to the appropriate probability model.	Build an interactive digital model selector or record an audio walkthrough classifying three complex random variable scenarios.
2.D	"Quality Control Inspection Lab" setting decision boundaries for product testing and calculating the probability of false alarms (Type I error, \alpha) versus missing a defect (Type II error, \beta).	Dual Normal/Binomial distribution curves illustrating overlapping null and alternative distributions, visually shading \alpha (Type I area) and \beta (Type II area).	Construct a visual error probability slide deck or create an interactive spreadsheet calculating how shifting the decision cutoff changes \alpha or \beta
2.E	"Is the Game Rigged?" challenge translating probability claims (e.g., "This spinner lands on win 50% of the time") into baseline parameter hypotheses (H0: p = 0.50 vs. Ha: p > 0.50).	Side-by-side reference chart translating verbal claims ("unfair advantage," "yields higher average payout") into symbolic parameters (p, \mu) and H_0/H_a statements.	Write symbolic hypotheses for four probability distribution prompts and record a video explaining how parameter assumptions define H0.
3.E	"Probability Escape Room" computing exact tail probabilities and cumulative probabilities under specified probability models to unlock progressive stages.	Step-by-step formula execution guides pairing exact Binomial/Normal formulas with TI-84/calculator syntax (binomcdf , normalcdf) side-by-side.	Submit a probability calculation portfolio featuring both handwritten formula work and calculator or computational script outputs.
4.E	"BINS Inspector" auditing probability scenarios to spot violated conditions (e.g., checking Binary outcomes, Independence, Fixed n, Constant p, 10% condition, or np \ge 10, n(1-p) \ge 10.	Visual checklist for BINS (Binomial settings) and Normal approximation criteria with visual status indicators (satisfied vs. violated).	Produce a "Model Verification Report" evaluating whether a real-world scenario satisfies Binomial or Normal distribution conditions.
4.F	"Probability Translator" converting raw cumulative probability outputs into consumer-facing explanations regarding luck versus evidence of model bias.	Sentence-template bank providing standardized AP phrasing: "Assuming [baseline model parameter], the probability of observing [result] or more extreme by chance is [probability]."	Record a video or podcast segment explaining what a calculated probability result implies about a casino game or manufacturing standard.
4.G	"Fair Game Courtroom" using probability calculations as evidence to deliver a binding	Decision tree graphic linking probability thresholds (P(X) < \alpha) to final conclusions	Write a complete CER justification essay or deliver a closing argument video evaluating a baseline probability

	decision on whether a process or game deviates from its claimed parameter.	regarding the validity of a baseline probability claim.	claim based on calculated tail probabilities.	
Supporting Multilingual/English Learners (<i>CELP standards</i>)				
	Emerging	Expanding	Bridging	
2.C	Match quantitative inference cards to visual study descriptions.	Select the appropriate t-inference procedure (distinguishing 2-sample independent vs. paired t-tests) using a graphic organizer.	Differentiate between independent 2-sample t-procedures and paired t-procedures based on subtle experimental/sampling design descriptions.	
2.D	Match Type I / Type II error definitions to t-test scenario diagrams with picture cues.	Describe consequences of Type I and Type II errors in a quantitative study using structured cause-and-effect templates.	Evaluate how changing sample size n or population standard deviation \sigma impacts t-distribution shape, degrees of freedom, and test power.	
2.E	Select H0 and Ha statements using mean parameter symbols from visual choice cards.	Write H0 and Ha hypotheses for means in symbols and defined words using structured sentence frames..	Formulate hypotheses for 1-sample, paired, and 2-sample t-tests, precisely defining population means and mean differences in context.	
3.E	Locate degrees of freedom, t-statistic, and p-value on a computer output screenshot using a visual key.	Calculate t-statistics, degrees of freedom, and confidence intervals for means using annotated formula guides and t-tables.	Compute 1-sample, 2-sample, and paired t-test statistics, p-values, and confidence intervals using technology and mathematical formulas.	
4.E	Check off conditions (<i>Random, 10% condition, Normal/Large Sample or no strong skew/outliers</i>) on a visual checklist.	Check and describe t-inference conditions, sketching sample graphs to verify the Normal/Large Sample condition using guided templates.	Evaluate t-procedure conditions, assessing sample plot skewness/outliers when n < 30, and justify whether inference remains robust.	
4.F	Complete a standard confidence interval interpretation frame for means: "We are 95% confident that the true mean [context] is..."	Interpret t-test p-values and confidence intervals for means in context using structured sentence frames.	Interpret confidence intervals for mean differences and p-values, highlighting practical vs. statistical significance.	
4.G	Circle "Reject H0" or "Fail to Reject H0" based on p-value vs. $\alpha = 0.05$ comparison cards.	Draft formal t-test conclusions linked to claims using structured decision templates ("There is [sufficient / insufficient] evidence...").	Synthesize test results and confidence intervals to evaluate real-world claims about population means, addressing directional vs. non-directional claims.	
Unit Outline				
Lesson Sequence	Skills	Learning Objective	Essential Knowledge	AP Activities (if available)
Day 1 4.1 Sampling Distributions for Sample Means	3.D Calculate means, standard deviations, and parameters for probability distributions.	4.1.A Calculate the mean and standard deviation of a sampling distribution of a sample mean.	4.1.A.1 For a population with population mean μ and population standard deviation σ , when the sampled values are independent, the sampling distribution of the sample mean has mean $\mu_{\bar{x}} = \mu$ and standard deviation $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$	Quick Write Provide students with a prompt such as "Describe how to construct a sampling distribution." Then allow 1 or 2 minutes for them to write a response (e.g., "obtain sample, collect responses, calculate statistic, plot statistic, repeat"). Allow more writing time for more difficult prompts.
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	4.1.B Justify the appropriateness of conditions for the sampling distribution of a sample mean.	4.1.B.1 Sampling without replacement requires that two conditions must be met: 4.1.B.1.i The randomization condition—the data should be collected using a random sample. 4.1.B.1.ii	

			The 10% condition—the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size. 4.1.B.2 For a quantitative variable, if the population distribution can be modeled by a normal distribution, the sampling distribution of the sample mean, $\bar{x}$, can be modeled with a normal distribution regardless of the sample size. 4.1.B.3 For a quantitative variable, if the population distribution cannot be modeled by a normal distribution, the sampling distribution of the sample mean, $\bar{x}$, can be modeled approximately by a normal distribution, provided $n \geq 30$. If the population distribution is extremely skewed, a sample size much larger than 30 may be needed to ensure the sampling distribution is approximately normal.	
	4.D Interpret statistical calculations and results to assess meaning or a claim.	4.1.C Interpret the mean, standard deviation, and probabilities for the sampling distribution of a sample mean.	4.1.C.1 The mean, standard deviation, and probabilities for a sampling distribution of a sample mean should be interpreted within the context of a specific population.	
Day 2 4.2 <i>Constructing a Confidence Interval for a Population Mean or Population Mean Difference</i>	4.C Describe distributions and compare relative positions of points within a distribution.	4.2.A Describe t-distributions.	4.2.A.1 t-distributions, also called Student's t-distributions, form a family of symmetric, bell-shaped, standardized distributions with wider tails than that of the standard normal distribution. Specific t-distributions are identified using a parameter known as the number of degrees of freedom (df), which is based on the sample size(s). When the degrees of freedom are small, the t-distribution has a much narrower peak and fatter tails than a normal distribution. As the degrees of freedom increase, the t-distribution more closely resembles the standard normal distribution (mean $\mu = 0$ and standard deviation $\sigma = 1$). 4.2.A.2 t-distributions are used for finding critical values and test statistics for inferences about a population mean, μ, when the population standard deviation, σ, is unknown and the sample standard deviation, s, must be used instead.	Discussion Groups Place students in groups of three or four and ask each group to identify the conditions for performing a test about a population mean. For each condition, have them explain why the condition is required and what would go wrong with the test if the condition were violated. Have groups pair up and compare answers. Round Table (see Topic 4.3) Team Challenge: Marking the Text (see Topic 4.4) Graphic Organizer (see Topic 4.7)
	2.C Identify appropriate statistical inference methods.	4.2.B Identify an appropriate confidence interval procedure including the parameter for a	4.2.B.1 The appropriate confidence interval procedure for estimating the population mean of a quantitative variable for one sample is a one-sample t-interval for a population mean. (The	

		population mean or population mean difference.	population standard deviation, σ, is not typically known for distributions for quantitative variables.) 4.2.B.2 For a matched pairs design with two dependent samples, the appropriate analysis calculates differences between pairs of values to produce one sample of differences. The confidence interval procedure for the matched pairs design is a one-sample t-interval for a population mean difference. 4.2.B.3 The parameter for a confidence interval for a population mean or population mean difference should reference the population mean or population mean difference and the response variable, in context. For the population mean difference, it is important to state the order of subtraction for the difference	
4.E Justify the use of a chosen statistical inference method by verifying conditions.	4.2.C Justify the appropriateness of constructing a confidence interval for a population mean or population mean difference by verifying conditions.		4.2.C.1 A one-sample t-interval for a population mean or population mean difference requires that three conditions be met: 4.2.C.1.i The randomization condition—the data should be collected using a random sample or a randomized experiment. 4.2.C.1.ii The 10% condition—when sampling without replacement, the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size. 4.2.C.1.iii The sample data condition—it is indicated the population distribution is approximately normal, or $n \geq 30$, or if $n < 30$, the sample data distribution should be free from strong skewness and outliers. For matched pairs, the number of differences should be greater than or equal to 30. If the number of differences is less than 30, the sample of differences should be free from strong skewness and outliers.	
3.E Calculate appropriate statistical inference method results.	4.2.D Calculate an appropriate confidence interval for a population mean or population mean difference.		4.2.D.1 A point estimate for a population mean is the sample mean, $\bar{x}$, or $\bar{x}_d$ for the sample mean difference. 4.2.D.2 To estimate the population mean for one sample or the population mean difference between values in matched pairs, when the population standard deviation is unknown, the confidence interval is $\bar{x} \pm t^* \frac{s}{\sqrt{n}}$, where t* is the	

			critical value for the central C% of a t-distribution with degrees of freedom $n - 1$.	
		4.2.E Calculate the standard error and margin of error for a sample size for a one-sample t-interval.	4.2.E.1 The standard error (SE) for a sample mean is given by $SE_{\bar{x}} = \frac{s}{\sqrt{n}}$. 4.2.E.2 For a one-sample t-interval for a population mean, the margin of error is the critical value (t^*) times the standard error (SE), which equals $(t^*)\left(\frac{s}{\sqrt{n}}\right)$.	
Day 3 4.3 <i>Justifying a Claim Based on a Confidence Interval for a Population Mean or Population Mean Difference</i>	4.F Interpret results of statistical inference methods.	4.3.A Interpret a confidence interval in context for a population mean or population mean difference.	4.3.A.1 Because the confidence interval for a population mean or population mean difference is calculated based on a sample from a population, the computed interval may or may not contain the value of the population mean or population mean difference. 4.3.A.2 The interpretation of the confidence level is as follows: In repeated random sampling with the same sample size from the same population, approximately C% of confidence intervals created will capture the population mean or population mean difference, where C represents the numerical value of the confidence level used. 4.3.A.3 When interpreting a C% confidence interval for a population mean or population mean difference, we say we are C% confident the interval (a, b) contains the value of the population mean or population mean difference, where a represents the lower limit and b represents the upper limit. An interpretation of a confidence interval for a population mean or population mean difference includes a reference to the parameter.	Round Table Arrange students in groups of four and give each student an identical paper containing four different inference procedure problems from Unit 4. Have students complete the first one on their paper and then pass the paper clockwise to another member in their group. That student checks the first problem and then completes the second problem on the paper. Have students rotate again, continuing the process until each student has their original paper back.
	4.G Justify a claim based on statistical inference method results.	4.3.B Justify a claim based on a confidence interval for a population mean or population mean difference.	4.3.B.1 A confidence interval for a population mean or population mean difference provides an interval of values that may serve as convincing evidence to support a particular claim about the population mean or population mean difference.	
	2.D Identify types of errors and relationships among components in statistical inference methods.	4.3.C Identify the relationships among sample size, confidence interval width, confidence level, and margin of error for a	4.3.C.1 For a given sample, increasing the confidence level will result in the following: 4.3.C.1.i The critical value will increase. 4.3.C.1.ii The margin of error will increase.	

		population mean or population mean difference.	4.3.C.1.iii The width of the confidence interval will increase. 4.3.C.2 Increasing the sample size decreases the standard error. Thus, when all other things remain the same, the width of a confidence interval for a population mean or population mean difference tends to decrease as the sample size increases. For a confidence interval for a population mean or population mean difference with a given confidence level, the width of the interval is $\frac{1}{\sqrt{n}}$ approximately proportional to	
Day 4 4.4 <i>Setting Up a Test for a Population Mean or Population Mean Difference</i>	2.C Identify appropriate statistical inference methods.	4.4. A Identify an appropriate testing method and parameter for a population mean or population mean difference with unknown σ .	4.4.A.1 The appropriate test for a population mean with unknown population standard deviation σ is a one-sample t-test for a population mean. 4.4.A.2 For a matched pairs design with two dependent samples, the appropriate analysis calculates differences between pairs of values to produce one sample of differences. The hypothesis testing procedure for the matched pairs design is a one-sample t-test for the population mean difference. 4.4.A.3 The parameter for a hypothesis test for a population mean and population mean difference should reference the population parameter, the response variable, and the population in context.	Team Challenge: Marking the Text Give students examples of statistical situations. Have them highlight or underline the key phrases to help them determine which statistical test to use for each situation. For example, does it require a significance test or a confidence interval? Does it involve means or proportions? Is there only one variable, or are there two or more variables?
	2.E Identify the null and alternative hypotheses.	4.4.B Identify the null and alternative hypotheses for a population mean or population mean difference with unknown σ .	4.4.B.1 The null hypothesis for a one-sample t-test for a population mean is $H_0: \mu = \mu_0$, in which μ_0 is the null hypothesized value for the population mean. A one-sided alternative hypothesis for a one-sample t-test for a population mean is either $H_a: \mu < \mu_0$ or $H_a: \mu > \mu_0$. A two-sided alternative hypothesis is $H_a: \mu \neq \mu_0$. 4.4.B.2 The null hypothesis for a population mean difference is $H_0: \mu_d = 0$. A one-sided alternative hypothesis for a population mean difference is either $H_a: \mu_d < 0$ or $H_a: \mu_d > 0$. A two-sided alternative hypothesis is $H_a: \mu_d \neq 0$.	Round Table (see Topic 4.3)
	4.E Justify the use of a chosen statistical inference method by	4.4.C Justify the appropriateness of a	4.4.C.1 A one-sample t-test for a population mean or a population mean difference	Graphic Organizer (see Topic 4.7)

	verifying conditions.	hypothesis test for a population mean or population mean difference by verifying conditions.	requires that three conditions be met: 4.4.C.1.i The randomization condition—the data should be collected using a random sample or a randomized experiment. 4.4.C.1.ii The 10% condition—when sampling without replacement, the population size must be at least 10 times larger than the sample size ($n \leq 10\%N$), where N is the size of the population and n is the sample size. 4.4.C.1.iii The sample data condition—it is indicated the population distribution is approximately normal, or $n \geq 30$, or if $n < 30$, the sample data distribution should be free from strong skewness and outliers. For matched pairs, the number of differences should be greater than or equal to 30. If the number of differences is less than 30, the sample of differences should be free from strong skewness and outliers.	
Day 5 4.5 <i>Carrying Out a Test for a Population Mean or Population Mean Difference</i>	3.E Calculate appropriate statistical inference method results.	4.5. A Calculate an appropriate test statistic and p-value for testing a hypothesis about a population mean or population mean difference	4.5.A.1 The test statistic for a one-sample t-test for a population mean or population mean difference is $t = \frac{\bar{x} - \mu_0}{\frac{s}{\sqrt{n}}}$, where t has degrees of freedom $n - 1$. The t -statistic has a t -distribution with degrees of freedom $n - 1$ when the null hypothesis is true. 4.5.A.2 The p-value for a one-sample t-test for a population mean or population mean difference is found using the appropriate t -distribution table or technology.	Round Table (see Topic 4.3)
	4.F Interpret results of statistical inference methods.	4.5.B Interpret the p-value of a hypothesis test for a population mean or population mean difference.	4.5.B.1 The p-value is the probability of obtaining a test statistic as extreme or more extreme than the test statistic that was observed (i.e., in the direction of the alternative hypothesis) given that the null hypothesis is true. An interpretation of the p-value of a hypothesis test for a population mean or population mean difference should include a statement that the p-value is computed by assuming that the null hypothesis is true (i.e., by assuming that the population mean is equal to the particular value stated in the null hypothesis in context).	
	4.G Justify a claim based on statistical inference method results.	4.5.C Justify a claim about the population based on the results of a hypothesis test for a population mean or population	4.5.C.1 A formal decision explicitly compares the p-value to the significance level, α . If the p-value $\leq \alpha$, then reject the null hypothesis, $H_0: \mu = \mu_0$. If the	

		mean difference	p-value > α, then fail to reject the null hypothesis. 4.5.C.2 The results of a hypothesis test for a population mean or population mean difference can serve as the statistical reasoning to support the answer to an investigative question about the population that was sampled. 4.5.C.3 A conclusion for the hypothesis test for a population mean or population mean difference is stated in context consistent with, and in terms of, the alternative hypothesis using nondefinitive language. The conclusion should contain a reference to the parameter and the population.	
Day 6 4.6 <i>Sampling Distributions for the Difference Between Two Sample Means</i>	3.D Calculate means, standard deviations, and parameters for probability distributions.	4.6.A Calculate the mean and standard deviation of a sampling distribution for the difference between two sample means.	4.6.A.1 For two independent populations with population means μ_1 and μ_2 and population standard deviations σ_1 and σ_2 , when the sampled values are independent, the sampling distribution of the difference in sample means $\bar{x}_1 - \bar{x}_2$ has a mean $\mu_{(\bar{x}_1 - \bar{x}_2)} = \mu_1 - \mu_2$ and standard deviation $\sigma_{(\bar{x}_1 - \bar{x}_2)} = \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$.	
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	4.6.B Justify the appropriateness of conditions for the sampling distribution of the difference between two sample means.	4.6.B.1 Sampling without replacement requires that two conditions must be met: 4.6.B.1.i The randomization condition—the data should be collected using two independent random samples. 4.6.B.1.ii The 10% condition—the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N_1$ and $n_2 \leq 10\%N_2$, where N_1 is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2 . 4.6.B.2 If the data come from an experiment, the data only need to meet the randomization condition. The treatments must be randomly assigned to experimental units to meet the randomization condition. 4.6.B.3 The sampling distribution for the difference between sample means, $\bar{x}_1 - \bar{x}_2$, can be modeled with a normal distribution if the two population distributions can each be modeled by a normal distribution. 4.6.B.4 The sampling distribution for the difference between sample means, $\bar{x}_1 - \bar{x}_2$, can be modeled approximately by a normal distribution if the two population distributions cannot be modeled by a normal distribution but $n_1 \geq 30$ and $n_2 \geq 30$.	
	4.D Interpret statistical calculations and results to assess meaning or a claim.	4.6.C Interpret the mean, standard deviation, and probabilities for a sampling	4.6.C.1 The mean, standard deviation, and probabilities for a sampling distribution for the difference between sample means should be interpreted within the context of specific populations.	

		distribution for the difference between sample means.		
Day 7 4.7 <i>Constructing a Confidence Interval for the Difference Between Two Population Means</i>	2.C Identify appropriate statistical inference methods.	4.7.A Identify an appropriate confidence interval procedure including the parameter for the difference between two population means.	4.7.A.1 Based on the sample data, a confidence interval can be calculated to estimate the difference between two population means. The appropriate confidence interval procedure for two independent samples is a two-sample t-interval for the difference between population means. 4.7.A.2 The parameter for a confidence interval for a two-sample t-interval for the difference between population means should reference the difference in the means, the response variable, and the populations in context.	Graphic Organizer Have students work in teams of two or three to develop a flowchart for determining which inference procedure from Units 3 and 4 to use in a given setting. Team Challenge (See Topic 4.4) Round Table (see Topic 4.3)
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	4.7.B Justify the appropriateness of constructing a confidence interval for the difference between two population means by verifying conditions.	4.7.B.1 A two-sample t-interval for a difference between population means requires that three conditions be met: 4.7.B.1.i The randomization condition—the data should be collected using two independent random samples or a randomized experiment. 4.7.B.1.ii The 10% condition—when sampling without replacement, the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N_1$ and $n_2 \leq 10\%N_2$, where N_1 is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2 . (Note: This condition is unnecessary when the data are from a randomized experiment). 4.7.B.1.iii The sample data condition—both samples should have a sample size greater than or equal to 30 or it is indicated that both population distributions are approximately normal. If either sample size is less than 30, both sample data distributions should be free from strong skewness and outliers.	
	3.E Calculate appropriate statistical inference method results.	4.7.C Calculate an appropriate confidence interval for the difference between two population means.	4.7.C.1 A point estimate for the difference between two population means is the difference in sample means, $\bar{x}_1 - \bar{x}_2$. 4.7.C.2 For the difference between population means when the population standard deviations are unknown, the confidence interval can be constructed as point estimate $\pm$ (margin of error). The confidence interval for the difference between population means	

			$(\bar{x}_1 - \bar{x}_2) \pm t^* \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$, where t^* are the critical values for the central C% of a t-distribution with appropriate degrees of freedom that can be found using technology. The degrees of freedom fall between $n_1 + n_2 - 2$ and the smaller of $n_1 - 1$ and $n_2 - 1$.	
		4.7.D Calculate the standard error and margin of error for estimating the difference between two population means.	4.7.D.1 The standard error for the difference between two sample means is $SE_{\bar{x}_1 - \bar{x}_2} = \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$, where s_1 and s_2 are the sample standard deviations. 4.7.D.2 For the difference between two sample means, the margin of error is the critical value (t^*) times the standard error (SE) of the difference of two sample means, which equals $t^* \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$.	
Day 8	Mid-Unit Assessment			
Day 9 4.8 Justifying a Claim Based on a Confidence Interval for the Difference Between Two Population Means	4.F Interpret results of statistical inference methods.	4.8.A Interpret a confidence interval in context for the difference between two population means	4.8.A.1 Because the confidence interval for the difference between two population means is calculated based on samples from two populations, the computed interval may or may not contain the value for the difference between the two population means. 4.8.A.2 The interpretation of the confidence level is as follows: In repeated random sampling with the same sample size from the same populations, approximately C% of confidence intervals created will capture the difference between the two population means, where C represents the numerical value of the confidence level used. 4.8.A.3 When interpreting a C% confidence interval for the difference between two population means, we say we are C% confident that the interval (a, b) contains the value of the difference in the population means, where a represents the lower limit and b represents the upper limit. An interpretation of a confidence interval for the difference between two population means includes a reference to the difference in the population	Team Challenge (See Topic 4.4) Round Table (see Topic 4.3) Graphic Organizer (see Topic 4.7)

			means with the details about the populations it represents in the context of the study.	
	4.G Justify a claim based on statistical inference method results	4.8.B Justify a claim based on a confidence interval for the difference between two population means	4.8.B.1 A confidence interval for the difference between two population means provides an interval of values that may serve as convincing evidence to support a particular claim about the difference in two population means. For example, if the interval contains 0, then there is insufficient evidence to conclude there is a difference between the two population means. If the interval does not contain 0, then there is sufficient evidence to conclude there is a difference between the two population means.	
Day 10 4.9 <i>Setting Up a Test for the Difference Between Two Population Means</i>	2.C Identify appropriate statistical inference methods.	4.9.A Identify an appropriate testing method for the difference between two population means including the parameters for the difference between the two population means	4.9.A.1 The appropriate test for the difference between two population means is a two-sample t-test for a difference between two population means. 4.9.A.2 The parameters for a hypothesis test for the difference between two population means should reference the population parameters, the response variables, and the populations in context.	Round Table (see Topic 4.3)
	2.E Identify the null and alternative hypotheses.	4.9.B Identify the null and alternative hypotheses for the difference between two population means.	4.9.B.1 The null hypothesis for a two-sample t-test for the difference between two population means, μ_1 and μ_2 , can be written as either $H_0: \mu_1 - \mu_2 = 0$ or $H_0: \mu_1 = \mu_2$. A one-sided alternative hypothesis for the difference between population means can be written as either $H_a: \mu_1 < \mu_2$ (or equivalently $H_a: \mu_1 - \mu_2 < 0$) or $H_a: \mu_1 > \mu_2$ (or equivalently $H_a: \mu_1 - \mu_2 > 0$). A two-sided alternative hypothesis for the difference between population means can be written as $H_a: \mu_1 \neq \mu_2$ (or equivalently $H_a: \mu_1 - \mu_2 \neq 0$).	
	4.E Justify the use of a chosen statistical inference method by verifying conditions.	4.9.C Justify the appropriateness of a hypothesis test for the difference between two population means by verifying conditions.	4.9.C.1 A two-sample t-test for a difference between population means requires that three conditions be met: 4.9.C.1.i The randomization condition—the data should be collected using two independent random samples or a randomized experiment. 4.9.C.1.ii The 10% condition—when sampling	

			without replacement, the size of each sample should be less than or equal to 10% of the respective population size: $n_1 \leq 10\%N_1$ and $n_2 \leq 10\%N_2$, where N_1 is the size of population 1 and N_2 is the size of population 2. The sample sizes are represented as n_1 and n_2. (Note: This condition is unnecessary when the data are from a randomized experiment.) 4.9.C.1.iii The sample data condition—both samples should have a sample size greater than or equal to 30 or it is indicated that both population distributions are approximately normal. If either sample size is less than 30, both sample data distributions should be free from strong skewness and outliers.	
Day 11 4.10 <i>Carrying Out a Test for the Difference Between Two Population Means</i>	3.E Calculate appropriate statistical inference method results.	4.10.A Calculate an appropriate test statistic and p-value for testing a hypothesis for the difference between two population means.	4.10.A.1 The test statistic for a two-sample t-test for the difference between two population means is $t = \frac{(\bar{x}_1 - \bar{x}_2) - 0}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$ The t-statistic has a t-distribution when the null hypothesis is true. The t-statistics with the degrees of freedom can be found using technology. The degrees of freedom fall between $n_1 + n_2 - 2$ and the smaller of $n_1 - 1$ and $n_2 - 1$. 4.10.A.2 The p-value for a two-sample t-test for the difference between two population means can be found using the appropriate t-distribution table or from the appropriate t-distribution using technology	Round Table (see Topic 4.3)
	4.F Interpret results of statistical inference methods.	4.10.B Interpret the p-value of a hypothesis test for the difference between two population means	4.10.B.1 The p-value is the probability of obtaining a test statistic as extreme or more extreme than the test statistic that was observed (i.e., in the direction of the alternative hypothesis) given that the null hypothesis is true. An interpretation of the p-value of a hypothesis test for a two-sample test for the difference between two population means should include a statement that the p-value is computed by assuming that the null hypothesis is true (i.e., by assuming the population means are equal to each other in context).	
	4.G Justify a claim based	4.10.C Justify a claim about	4.10.C.1 A formal decision explicitly compares	

	on statistical inference method results.	the populations based on the results of a hypothesis test for the difference between two population means.	the p-value to the significance level, α. If the p-value $\leq \alpha$, then reject the null hypothesis, $H_0: \mu_1 - \mu_2 = 0$ or $H_0: \mu_1 = \mu_2$. If the p-value $> \alpha$, then fail to reject the null hypothesis. 4.10.C.2 The results of a hypothesis test for a two-sample t-test for a difference between two population means can serve as the statistical reasoning to support the answer to an investigative question about the two populations that were sampled. 4.10.C.3 A conclusion for the hypothesis test for the difference between two population means is stated in context consistent with, and in terms of, the alternative hypothesis using nondefinitive language. The conclusion should contain a reference to the parameters and the populations.	
Day 12	Unit Test: AP Progress Check 4			

Unit Title		
Unit 5: Regression Analysis		
Relevant Standards: Bold indicates priority		
3.A, 3.B, 4.A, 4.B, 4.D		
Essential Question(s)		Enduring Understanding(s)
- Does the fact that the number of shark attacks increases as ice cream sales increase necessarily mean that ice cream sales cause shark attacks? - How can we determine the effectiveness of a linear model that uses the number of cricket chirps per minute to predict temperature?		- Regression models describe and quantify relationships between quantitative variables. - Correlation and regression can identify patterns and support prediction but do not establish causation. - Residual analysis helps determine whether a model appropriately represents the data. - Statistical inference can be used to evaluate the strength and significance of relationships.
Demonstration of Learning		Pacing for Unit
Progress Check 5 - Multiple-choice: ~19 questions - Free-response: 2 questions - Question 1: Multi-Focus on Practices 2, 3, and 4 - Question 2: Multi-Focus on Practices 2, 3, and 4 The additional online question bank will be used to create a mid-unit quiz and other checks for understanding throughout the unit. College Board AP Exam Section I MCQ Weighting: 10–20% - Descriptive Bivariate Data: Interpreting scatterplots for direction, form, strength, and unusual features. - Least-Squares Regression Line (LSRL): Calculating and interpreting the slope, y-intercept, standard deviation of residuals, and coefficient of determination. - Residual Analysis: Calculating residuals, interpreting residual plots to assess model linearity, and identifying influential points vs. outliers See AP Exam Overview for more information		7 Block Periods <i>NOTE:</i> Midterm & AP Exam Review / Mock Exams - 11 Blocks Flexibility/Buffer/Post-AP Exam Project - 6 Blocks
Family Overview (link below)		Integration of Technology:
AP Statistics FAMILY MATERIALS		- Create scatterplots and regression models using graphing calculators and Desmos. - Calculate and interpret correlation and regression statistics. - Use regression output generated by calculators and software. - Conduct inference for regression parameters using graphing calculators and Desmos.. - Evaluate model fit using technology.
Unit-specific Vocabulary		Aligned Unit Materials, Resources, and Technology (beyond core resources)
Scatterplot	Residual	www.mathmedic.com www.desmos.com www.stapplet.com https://skewthescrypt.org/
Association	Residual plot	
Linear relationship	Explanatory variable	
Correlation	Response variable	
Correlation coefficient (r)	Coefficient of	

Regression Least-squares regression line Slope Intercept Predicted value	determination (r^2) Extrapolation Influential point Outlier Model fit Inference for slope		
Opportunities for Interdisciplinary Connections:		Anticipated misconceptions	
- Science: modeling relationships among variables. - Economics: forecasting and trend analysis. - Business: sales and marketing analytics. - Environmental Science: climate and ecological modeling. - Sports Analytics: performance prediction and evaluation.		- Assuming correlation implies causation or using correlation to evaluate non-linear relationships. - Using deterministic language rather than estimated/predicted language when interpreting slope and y-intercept. - Relying solely on a high R^2 value to assess model fit instead of checking for randomness in a residual plot. - Confusing high-leverage points, outliers, and influential points in bivariate data. - Reading statistics from the intercept/constant row instead of the predictor variable row on computer software printouts during slope inference.	
Connections to Prior Units		Connections to Future Units	
- Builds on graphical analysis and association concepts from Unit 1. - Uses probability and inference concepts from Units 2–4. - Requires interpretation of variability and statistical significance.		Provides foundational preparation for college statistics, data science, economics, social science research, and STEM fields.	
Differentiation through Universal Design for Learning			
	Engagement	Representation	Action & Expression
2.A	"Parameter vs. Statistic Treasure Hunt" sorting prompt information into population parameters vs. sample statistics across polling and manufacturing scenarios.	Color-coded visual decoder mapping population parameters alongside corresponding sample statistics and sample sizes (n , N).	Annotate AP-style prompts using digital highlighters or tags to isolate parameters, statistics, and sample sizes.
2.C	"Model Matcher" challenge determining whether a scenario requires the sampling distribution of a proportion, difference in proportions, mean, or difference in means.	Flowchart decision tree categorizing sampling distribution models by variable type (categorical vs. quantitative) and group count (1 vs. 2 samples).	Construct a visual selection flowchart or record an audio walkthrough classifying sampling distribution models for 3 unlabelled scenarios.
3.A	"Penny Sampling Lab" physically drawing repeated sample means/proportions to build an empirical sampling distribution dotplot on a wall grid.	Visual progression chart illustrating Population Distribution to Individual Sample Distributions to Sampling Distribution of the Statistic.	Generate empirical sampling distributions using technology and export labeled graphs.
3.B	"Sample Size Effect Lab" calculating how changing n compresses standard error while keeping the mean centered at the true population parameter.	Formula reference card for sampling distribution parameters $(\mu_{\hat{p}} = p, \sigma_{\hat{p}} = \sqrt{\frac{p(1-p)}{n}}, \mu_{\bar{x}} = \mu, \sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}})$ and z-scores $(z = \frac{\text{stat} - \text{param}}{\text{SE}})$.	Calculate means, standard deviations, and z-scores for sampling distributions via worked calculation sheets, calculator steps, or video walkthroughs.
3.E	"Quality Assurance Escape Room" computing tail probabilities under sampling distribution models to pass inspection checkpoints.	Step-by-step workflow displaying z-score transformations and Normal distribution curve shading alongside calculator syntax (normalcdf).	Submit a probability calculation portfolio featuring both handwritten standard error calculations and software/calculator outputs.

4.A	"Three Distributions Sorting Challenge" placing population distributions, single sample distributions, and sampling distributions side-by-side to contrast shape, center, and spread.	Comparative anchor chart illustrating how sampling distribution variability decreases by $\frac{1}{\sqrt{n}}$ compared to population variability.	Deliver a written comparison essay or audio podcast detailing structural differences between a population distribution and its sampling distribution.
4.B	"Unbiased Estimator Debate" testing whether sample medians vs. sample means center at the true population center across simulated samples.	Bullseye graphic displaying high vs. low bias (center alignment) and high vs. low variability (spread width).	Write a Claim-Evidence-Reasoning (CER) justification evaluating whether a given sample statistic serves as an unbiased estimator of a parameter.
4.D	"Auditor Briefing" interpreting standard error to non-technical executives seeking to understand expected sample-to-sample fluctuation.	Sentence-stem bank providing mandatory AP formulaic phrasing for interpreting $\sigma_{\hat{p}}$ and $\sigma_{\bar{x}}$ in context.	Write an executive audit memo or record a client briefing explaining the contextual meaning of standard error for a quality control metric.
4.E	"Condition Auditor" checking whether sample size is large enough for the Central Limit Theorem or Large Counts and verifying 10% independence.	Visual condition checklist featuring icons for Randomness, 10% Condition, and Normality/CLT rules.	Produce a "Sampling Model Verification Report" detailing mathematical proofs verifying or refuting model conditions for a given dataset.
4.F	"Plausible Outcome Investigator" translating calculated sampling probabilities into plain-language assessments of how likely a sample result is under an assumed parameter.	Template guide providing standardized AP sentence frameworks: "Assuming [true parameter], the probability of getting a sample statistic of [observed value] or more extreme is [probability]."	Record an audio segment or write an explanatory analysis explaining what a sampling distribution probability output means in plain English.
4.G	"Is the Manufacturer Lying?" using sampling distribution probabilities as evidence to determine if a company's claimed parameter is credible.	Decision tree linking low sampling probabilities ($P < 0.05$) to rejecting parameter claims due to convincing sample evidence.	Write a formal CER justification essay or deliver a presentation verifying or refuting a parameter claim using sampling distribution calculations.

Supporting Multilingual/English Learners ([CELP standards](#))

	Emerging	Expanding	Bridging
2.A	Identify response variable (y) and explanatory variable (x) from a scatterplot using color-coded axis labels.	Extract slope, y-intercept, r , r^2 , and s from a computer regression output table using a graphic key.	Parse bivariate regression problem descriptions, identifying explanatory/response variables, transformation needs, and regression output parameters.
2.C	Match "Scatterplot / Slope" scenario cards to <i>Linear Regression t-test for Slope</i> using visual cards.	Identify when to use inference for slope vs. descriptive regression modeling using a decision flowchart.	Distinguish between descriptive bivariate analysis, linear regression slope t-inference, and non-linear model transformations.
3.A	Plot points (x, y) on a scatterplot grid and draw an estimated line of best fit using visual alignment guides.	Construct scatterplots and residual plots from two-variable quantitative data sets using step-by-step graphing templates.	Construct scatterplots, residual plots, and transformed plots, verifying linear pattern fit.
3.B	Calculate predicted values by plugging given x-values into a provided regression equation,	Calculate predicted values, residuals, slope, and r^2 using formula templates.	Calculate LSRL equations, residuals, standard error of slope, and transformed regression models using mathematical formulas and technology.
3.E	Identify the t-statistic for slope and corresponding p-value from a computer output table using a highlighter key.	Calculate test statistics and confidence intervals for slope using computer output and formula guides.	Independently compute t-test statistics, p-values, and t-confidence intervals for population regression.
4.A	Describe scatterplot direction (positive/negative) and form (linear/non-linear) using visual vocabulary cards.	Describe scatterplots using D.O.F.S. (<i>Direction, Outliers, Form, Strength</i>) in context using sentence frames.	Compare multiple bivariate relationships using D.O.F.S., r , r^2 , and s , explaining context and influential points.
4.B	Complete claim statements about linear association strength based on r values.	Justify whether a linear model is appropriate using residual plot patterns (<i>no pattern = linear model appropriate</i>) and r^2 templates.	Construct rigorous arguments justifying linear model appropriateness or non-linear transformation needs based on residual plots and r^2 .

4.D	Match slope and y-intercept definition cards to their contextual meanings in a sample regression line.	Interpret slope (b), y-intercept (a), r^2 , and s in context using structured sentence frames ("For every 1 unit increase in x...").	Interpret slope, y-intercept, r^2 , standard deviation of residuals s, and SE in complex real-world scientific contexts.
4.E	Check off slope t-inference conditions L.I.N.E. (Linear, Independent, Normal residuals, Equal variance) on a visual condition checklist card.	Verify slope inference conditions (L.I.N.E.) using provided residual plots and histograms with sentence starters.	Evaluate L.I.N.E. conditions using scatterplots, residual plots, and histograms.
4.F	Complete a confidence interval interpretation frame for slope: "We are 95% confident that for each 1 [unit x], the true mean [variable y] increases by..."	Interpret confidence intervals for slope β_1 and p-values for slope t-tests in context using structured sentence frames.	Interpret confidence intervals and hypothesis tests for slope β_1 in context, articulating the relationship between variables.
4.G	Circle "Linear Relationship Exists" or "No Linear Relationship" based on p-value vs. $\alpha = 0.05$ cards.	Write formal conclusions regarding linear relationships using decision templates ("Because p-value < α , we have evidence of a linear relationship...").	Formulate comprehensive statistical arguments regarding linear association between two quantitative variables based on regression slope inference.

Unit Outline

Lesson Sequence	Skills	Learning Objective	Essential Knowledge	AP Activities (if available)
Day 1 5.1 Graphical Representations Between Two Quantitative Variables	3.A Construct tabular and graphical representations of data and distributions.	5.1.A Construct scatterplots depicting the relationship between two quantitative variables.	5.1.A.1 A bivariate quantitative data set consists of observations of ordered pairs from two quantitative variables, collected from the same individuals in a sample or population, and can be used to construct a scatterplot. 5.1.A.2 A scatterplot shows the relationship between two quantitative variables for each observation, one corresponding to the value on the x-axis and one corresponding to the value on the y-axis. The explanatory variable is placed on the x-axis and is the variable whose values are used to explain or predict the corresponding values for the response variable, which is placed on the y-axis.	Sketch and Switch Ask students to create a sketch of a scatterplot using a prompt such as a "weak negative, non-linear association." Then have them switch papers to review each other's work. Team Challenge (see Topic 5.2) Relief Pitcher (see Topic 5.5)
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	5.1.B Describe the characteristics of a scatterplot.	5.1.B.1 A description of the association shown in a scatterplot includes form, direction, strength, and unusual features. 5.1.B.2 The form of the association shown in a scatterplot, if any, can be described as linear or non-linear. 5.1.B.3 The direction of the association shown in a scatterplot, if any, can be described as positive or negative. A positive association means that as values of the explanatory variable increase, the values of the response variable tend to increase. A negative association means that as values of the explanatory variable increase, the values of the response variable tend to decrease. 5.1.B.4 The strength of the association shown in a scatterplot is how closely the points	

			follow the general pattern. Strength can be described as strong, moderate, or weak. 5.1.B.5 Unusual features of a scatterplot include clusters of individual points or points that don't fit in the general pattern of association between the two variables.	
	4.B Justify a claim based on statistical calculations and results.	5.1.C Justify a claim using scatterplots depicting the relationship between two quantitative variables.	5.1.C.1 Scatterplots depicting the relationship between two numeric variables may reveal information that can be used to justify claims about the variable in context.	
Day 2 5.2 <i>Correlation</i>	4.D Interpret statistical calculations and results to assess meaning or a claim.	5.2.A Interpret the correlation for a linear relationship	5.2.A.1 The correlation coefficient, r, summarizes the strength and direction of the linear association between two quantitative variables. The correlation coefficient r is unit-free and always between -1 and 1, inclusive. A negative correlation coefficient value indicates a negative association, and a positive correlation coefficient value indicates a positive association. 5.2.A.2 The strength of the linear association is determined by how close the correlation coefficient is to -1 or 1. A value of $r = 0$ indicates that there is no linear association. A value of $r = -1$ or $r = 1$ indicates that there is a perfect linear association. 5.2.A.3 A correlation coefficient close to -1 or 1 does not necessarily mean that a linear model is appropriate. 5.2.A.4 A perceived or real relationship between two variables does not mean that changes in one variable cause changes in the other. That is, correlation does not necessarily imply causation.	Team Challenge Set up teams of three or four and task students with solving different multiple-part problems in Unit 5 with different datasets to answer investigative questions. Have teams explain their work to other teams. Relief Pitcher (see Topic 5.5)
Day 3 5.3 <i>Linear Regression Models</i>	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	5.3.A Calculate a predicted response value using a linear regression model.	5.3.A.1 If the form of the relationship between x and y appears linear, we can approximate the relationship between x and y using a linear regression model, which is a linear equation that uses an explanatory variable, x, to predict the response variable, y. 5.3.A.2 In a linear regression model, the predicted response value, denoted by $\hat{y}$, is calculated as $\hat{y} = a + bx$, where a is the y-intercept, b is the slope of the regression line, and x is the explanatory variable. 5.3.A.3 Extrapolation is predicting a response value using a value for the explanatory	Think-Aloud Have students toss a handful of coins and record how many land heads up. This is trial 1. Then have them remove the coins that landed heads up. For trial 2, have students toss the remaining coins (i.e., the coins left over after removing those that landed heads up) and record how many land heads up on the second toss. Ask students to think about the trend and make a prediction, discussing their ideas

			variable that is beyond the interval of x-values used to determine the regression line. The predicted value is less reliable the further the estimate is extrapolated. 5.3.A.4 Interpolation is predicting a response value using a value for the explanatory variable that is within the interval of x-values used to determine the regression line.	with a partner. Will it be linear? A scatterplot of many trials should show a nonlinear relationship. Team Challenge (see Topic 5.2) Relief Pitcher (see Topic 5.5)
Day 4	Mid-Unit Assessment			
Day 5 5.4 <i>Residuals</i>	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	5.4.A Calculate the differences between the observed and predicted values.	5.4.A.1 A residual is the difference between the observed response value and the predicted response value for the given value of the explanatory variable: $residual = y - \hat{y}$ or (residual - observed y - predicted y).	Reversing Interpretations Instead of asking students to interpret a residual, give them the residual and the equation of the least-squares regression line and ask them to make a prediction for a particular observation (e.g., "One wolf in the pack had a length of 14. m and a residual of -98. 7. What does that -98. 7 tell us about that particular wolf?") Team Challenge (see Topic 5.2) Relief Pitcher (see Topic 5.5)
	4.D Interpret statistical calculations and results	5.4.B Interpret the differences between the observed and predicted values.	5.4.B.1 If the residual is positive, the model underpredicts (underestimates) the value of the response variable. If the residual is negative, the model overpredicts (overestimates) the value of the response variable.	
	4.A Describe and compare tabular and graphical representations of data, as well as summary statistics.	5.4.C Describe the form of association of bivariate data using residual plots.	5.4.C.1 A residual plot is a scatterplot of the residuals versus the predicted response values (or the explanatory variable values). 5.4.C.2 Residual plots can be used to investigate the appropriateness of the linear regression model for the observed data. 5.4.C.3 The linear regression model should only be fit to the data if the data exhibit a linear trend. Apparent randomness in a residual plot for a linear regression model is confirmation of a linear form in the association between the two variables and indicates that the simple linear regression model is an appropriate model for the data. 5.4.C.4 Curvature in the residual plot for a linear	

			regression model suggests that the linear model is not the most appropriate model for the data.	
Day 6 5.5 <i>Least-Squares Regression</i>	3.B Calculate summary statistics, relative positions of points within a distribution, and predicted responses.	5.5.A Calculate the coefficients for the least-squares regression line model.	5.5.A.1 The simple linear regression model is fit to the data by minimizing the sum of the squares of the residuals. Because of this, the resulting equation is often called the least-squares regression line (LSRL) and is calculated using technology. This regression line will pass through the point $(\bar{x}, \bar{y})$. 5.5.A.2 The slope of the regression line, b , is calculated using technology. 5.5.A.3 The y-intercept of the regression line, a , is calculated using technology. 5.5.A.4 In simple linear regression, the correlation coefficient, r , is calculated using technology. 5.5.A.5 In simple linear regression, the square of the correlation coefficient, r^2 , is called the coefficient of determination. The value of r^2 is the proportion of variation in the response variable that is explained by the linear relationship with the explanatory variable.	Relief Pitcher Use this strategy to review multiple concepts at the end of a unit or to spiral review at the end of the course. For instance, when reviewing for the Unit 5 test, students can earn 5 points if they get a question correct when up at the “pitcher’s mound” and 3 points if they do not know the answer but their team in the “bullpen” helps them answer it. Team Challenge (see Topic 5.2) Relief Pitcher (see Topic 5.5)
	4.D Interpret statistical calculations and results to assess meaning or a claim.	5.5.B Interpret coefficients for the least-squares regression line model.	5.5.B.1 The coefficients of the least-squares regression line model (line of best fit) are the slope, b , and the y-intercept, a , because they are based on a sample of values. 5.5.B.2 The slope of the least-squares regression line can be interpreted as the predicted increase or decrease in the response variable for a one-unit increase in the explanatory variable, and it should be interpreted in context. 5.5.B.3 The y-intercept in the least-squares regression line is the predicted value of the response variable when the explanatory variable is equal to 0, and it should be interpreted in context. Sometimes, the y-intercept of the line does not have a reasonable interpretation in context because $x = 0$ might be beyond the interval of x -values used to determine the regression line (extrapolation). At other times, the y-intercept of the line does not have a logical interpretation in context because it might be a negative value for a response variable that has no negative values, such as height	
Day 7	Unit Test: AP Progress Check 4			

Section I: Multiple-Choice Questions (MCQs)

42 Questions | 90 Minutes | 50% of Exam Weight

In addition to content units, multiple-choice questions across the exam are explicitly mapped to the four Statistical Practices:

Statistical Practice	Description	MCQ Exam Weighting
Practice 1: Formulate Questions	Identify valid statistical questions requiring data investigation.	5–10%
Practice 2: Collect Data	Identify and justify data collection methods, experimental design, and inference procedures.	20–30%
Practice 3: Analyze Data	Construct graphs, calculate summary statistics, probabilities, standard errors, and test results.	25–35%
Practice 4: Interpret Results	Interpret statistical outputs, evaluate claims, check inference conditions, and communicate conclusions in context.	25–35%

Section II: Free-Response Questions (FRQs)

4 Questions | 90 Minutes | 50% of Exam Weight

Equal-weighted multi-part free-response questions (10 points each):

- Question 1 (Multi-Focus)
Primary focus on Practices 1 & 2 (Formulating Questions & Data Collection / Experimental Design).
- Question 2 (Multi-Focus)
Primary focus on Practices 3 & 4 (Data Analysis & Result Interpretation).
- Question 3 (Inference Focus)
Requires completing a formal inference procedure (a Confidence Interval or Hypothesis Test from Unit 3 or Unit 4), including checking conditions, executing calculations, and interpreting results in context.
- Question 4 (Comprehensive Multi-Focus)
Integrates Practices 2, 3, & 4 across multiple content units to evaluate deep statistical reasoning.