

AP Statistics Unit 3 - Statistical Inference

Content Area: **Math**
Course(s):
Time Period: **MP3**
Length: **45**
Status: **Published**

Unit Overview

Unit Summary	Unit Rationale
Unit 3 of AP Statistics focuses on formal statistical inference, helping students move from simply describing data to making well-informed decisions based on analysis. Over 45 days, students learn two main approaches: interval estimation using the PANIC template and hypothesis testing with the PHANTOMS method. They start by building and interpreting confidence intervals and running significance tests for a single population proportion. Along the way, they learn about Type I and Type II errors and the concept of statistical power. The unit then moves into comparing two independent proportions and introduces Chi-Square tests, such as Goodness-of-Fit, Independence, and Homogeneity, to help students analyze two-way tables and discover patterns in categorical data.	This complete Statistical Inference unit is designed as the main instructional focus of the AP Statistics curriculum. Descriptive analysis helps students summarize past events, while probability explores what could happen in ideal situations. Inference connects these areas to answer a key question: Is what we see in our sample data a real, general pattern, or just random chance? By grouping all the main inference topics—one proportion, two proportions, one mean, two means, categorical tables, and regression slopes—into a single continuous unit, the curriculum avoids the confusion that can arise when these topics are taught separately. This procedure helps students understand inference as a single, logical system with consistent rules, whether they are looking at medical treatment success rates or the average lifespan of a machine part.
In the second half of the unit, students shift from working with categorical data to focusing on quantitative measurements. They learn how to use inference methods for population means and bivariate slopes. Students are introduced to the Student's t-distribution and see how it accounts for extra uncertainty when estimating a population parameter from sample data. They practice t-procedures for single-sample means, learn how to analyze paired data with paired t-tests, and compare two independent groups. The unit concludes with regression inference, in which students check the LINER conditions and use computer outputs to test the significance of the population slope. Throughout the unit, students practice explicit communication, learn to check normality visually, and avoid making overly certain conclusions, all to help them succeed on the AP Exam's multiple-choice and free-response questions.	Breaking this large unit into 47-minute lessons using Math Medic's EFFL model helps students remember key ideas by focusing on patterns rather than memorizing formulas. High school students often struggle when they have to learn many different formulas, but through active inquiry, they see that every confidence interval follows the identical structure: a point estimate plus or minus a margin of error. Every test statistic simply measures how far a sample result is from a null claim in standard errors. This method helps students become more flexible thinkers, moving from z-procedures to t-distributions and Chi-Square curves without confusion. In the end, this super-unit gives students the time and practice they need to master the language, safety checks, and careful interpretations that turn basic calculations into strong statistical arguments.

- 3.1.A.1 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.
- 3.1.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 $\hat{p}$ is a point estimator for the population proportion p .
- 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 $\sigma_{\hat{p}} = \sqrt{p(1-p)/n}$.
- 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 $\hat{p}$ 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.
- 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.
- 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.
- 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.
- 3.3.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 $\pm$ (margin of error). For a population proportion, the one-sample z -interval estimate is $\hat{p} \pm z^* \sqrt{\hat{p}(1-\hat{p})/n}$.
- 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 $SE_{\hat{p}} = \sqrt{\hat{p}(1-\hat{p})/n}$.
- 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 $z^* \sqrt{\hat{p}(1-\hat{p})/n}$.
- 3.3.D.4 The formula for the margin of error (MOE) can be rearranged to solve for n , $n = [(z^*)^2(\hat{p})(1-\hat{p})]/(MOE)^2$, 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.

- 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 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.
- 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.
- 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 $1/\sqrt{n}$.
- 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.
- 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 $H_a: p \neq p_0$.
- 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_o , and the expected number of failures, $n(1 - p_o)$, should be at least 10.
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, ($P(z \leq - x) + (P(z \geq x))$), 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.
3.7.A.1	The test statistic for testing a population proportion is $z = (\hat{p} - p_o) / \sqrt{p_o(1 - p_o)/n}$. 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.
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, $H_o : p = p_o$. 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.
- 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.
- 3.8.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 $1 - \text{power}$.
- 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 $\text{power} = 0.80$). The probability of a Type II 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.
- 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.
- 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 $\sigma_{[\hat{p}_1 - \hat{p}_2]} = \sqrt{[(p_1(1 - p_1)/n_1) + (p_2(1 - p_2)/n_2)]}$.
- 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: $n_1 \leq 10\%N_1$ 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: $n_1\hat{p}_1 \geq 10$, $n_1(1 - \hat{p}_1) \geq 10$, $n_2\hat{p}_2 \geq 10$, $n_2(1 - \hat{p}_2) \geq 10$, where $n_1\hat{p}_1$ and $n_2\hat{p}_2$ are the expected number of successes and $n_1(1 - \hat{p}_1)$ and $n_2(1 - \hat{p}_2)$ are the expected number of failures.
- 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.
- 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.
- 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.
- 3.10.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 $\pm$ (margin of error). The interval estimate for the difference between two population proportions is $(\hat{p}_1 - \hat{p}_2) \pm z^* \sqrt{[(\hat{p}_1(1 - \hat{p}_1))/n_1] + [(\hat{p}_2(1 - \hat{p}_2))/n_2]}$.
- 3.10.D.1 The standard error (SE) for the difference between two population proportions is $SE_{[\hat{p}_1 - \hat{p}_2]} = \sqrt{[(\hat{p}_1(1 - \hat{p}_1))/n_1] + [(\hat{p}_2(1 - \hat{p}_2))/n_2]}$.
- 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 is $z^* \sqrt{[(\hat{p}_1(1 - \hat{p}_1))/n_1] + [(\hat{p}_2(1 - \hat{p}_2))/n_2]}$.
- 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.

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.
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 proportions should reference the population parameters, the response variables, and the populations in context.
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 $H_0: p_1 = p_2$ 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 $H_a: p_1 < p_2$ 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 $H_a: p_1 \neq p_2$ or equivalently $H_a: p_1 - p_2 \neq 0$.
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: $n_1 \leq 10\%N_1$ 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— $n_1\hat{p}_c$, $n_1(1 - \hat{p}_c)$, $n_2\hat{p}_c$, $n_2(1 - \hat{p}_c)$, must all be at least 10, with $\hat{p}_c = (n_1\hat{p}_1 + n_2\hat{p}_2)/(n_1 + n_2)$ being the combined (or pooled) proportion assuming that H_0 is true ($H_0: p_1 = p_2$ or $H_0: p_1 - p_2 = 0$).
3.13.A.1	The test statistic for the difference between two population proportions is $z = [(\hat{p}_1 - \hat{p}_2) - 0]/\sqrt{[\hat{p}_c(1 - \hat{p}_c)][(1/n_1) + (1/n_2)]}$, where $\hat{p}_c = (n_1\hat{p}_1 + n_2\hat{p}_2)/(n_1 + n_2)$ 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.
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).
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, 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.

- 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.
- 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.
- 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.
- 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.
- 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 expected count = [(row total) (column total)]/total table
- 3.15.B.1 The appropriate test statistic for a chi-square test for homogeneity or independence is the chi-square statistic $\chi^2 = \sum (\text{Observed Count} - \text{Expected Count})^2 / \text{Expected Count}$, 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.
- 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.

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).
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 $\sigma_{\bar{x}} = \sigma/\sqrt{n}$.
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.
4.1.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.
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.
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 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.
4.2.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.
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 $\bar{x} \pm t^* (s/\sqrt{n})$, where t^* is the critical value for the central $C\%$ of a t -distribution with degrees of freedom $n - 1$.
4.2.E.1	The standard error (SE) for a sample mean is given by $SE_{\bar{x}} = s/\sqrt{n}$.
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 $(t^*)(s/\sqrt{n})$.
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.
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.
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.
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 approximately proportional to $1/\sqrt{n}$.
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.
- 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 $H_a: \mu_d \neq 0$.
- 4.4.C.1 A one-sample t -test for a population mean or a population mean difference 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.
- 4.5.A.1 The test statistic for a one-sample t -test for a population mean or population mean difference is $t = (\bar{x} - \mu_0) / (s / \sqrt{n})$, 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.
- 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).
- 4.5.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 p -value $> \alpha$, 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 non-definitive language. The conclusion should contain a reference to the parameter and the population.
- 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 $\sigma_{\bar{x}_1 - \bar{x}_2} = \sqrt{(\sigma_1^2/n_1) + (\sigma_2^2/n_2)}$.
- 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.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 $n_2 \geq 30$.
4.6.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.
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.
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.
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 is $(\bar{x}_1 - \bar{x}_2) \pm t^* \sqrt{[(s_1^2/n_1) + (s_2^2/n_2)]}$, 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.1	The standard error for the difference between two sample means is $SE_{(\bar{x}_1 - \bar{x}_2)} = \sqrt{[(s_1^2/n_1) + (s_2^2/n_2)]}$, 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 $t^* \sqrt{[(s_1^2/n_1) + (s_2^2/n_2)]}$.
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 means with the details about the populations it represents in the context of the study.
- 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.
- 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.
- 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 $H_a: \mu_1 - \mu_2 \neq 0$).
- 4.9.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.
- 4.10.A.1 The test statistic for a two-sample t -test for the difference between two population means is $t = [(\bar{x}_1 - \bar{x}_2) - 0] / \sqrt{[(s_1^2/n_1) + (s_2^2/n_2)]}$. 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.
- 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).
- 4.10.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, $H_0: \mu_1 - \mu_2 = 0$ 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.

Standards for Mathematical Practice

- 1 Make sense of problems and persevere in solving them
- 2 Reason abstractly and quantitatively
- 3 Construct viable arguments and critique the reasoning of others
- 4 Model with mathematics
- 5 Use appropriate tools strategically
- 6 Attend to precision
- 7 Look for and make use of structure
- 8 Look for and express regularity in repeated reasoning

Unit Focus

Enduring Understandings	Essential Questions
1. We can leverage sample statistics to construct a confidence interval that estimates an unknown population parameter with a quantified margin of error and a specified level of long-run reliability. 2. Hypothesis testing utilizes the laws of probability to determine if an observed sample result is due to mere random sampling variability or if it represents a statistically significant, generalizable phenomenon. 3. The type of variable being evaluated (categorical vs. quantitative) and the structure of the study design dictate the unique inferential framework (z, t, or χ^2) and safety conditions required to draw mathematically sound conclusions.	1. How can we use data from a single sample to estimate an entire population's hidden characteristics with absolute precision and confidence? 2. How do we decide whether a strange or extreme sample result is a statistical fluke or a real breakthrough? 3. How do the type of data we have and the way we grouped our subjects change the mathematical rules we must use to make an inference?

Instructional Focus

Learning Targets

Part 1: Inference for One Proportion

- I can interpret the meaning of a confidence level and a confidence interval in context.
- I can construct and interpret a one-proportion z -interval, verifying the required conditions (Random, 10%, Large Counts).
- I can explain the relationship between sample size, confidence level, and width of an interval, and compute the minimum sample size required to achieve a target margin of error.
- I can construct a valid set of hypotheses (H_0 and H_a) and interpret a P-value as a conditional probability.
- I can define, detect, and interpret Type I and Type II errors in context, and explain the consequences of each.
- I can perform a full significance test for a population proportion and evaluate the test's power.

Part 2: Inference for Two Proportions and Two-Way Tables

- I can check conditions and construct a confidence interval to estimate the difference between two population proportions.
- I can perform a significance test to compare two population proportions using a pooled sample proportion ($\hat{p}_c$) for the standard error.
- I can use a Chi-Square Goodness-of-Fit Test to evaluate whether an observed categorical distribution matches a hypothesized historical distribution.
- I can construct a two-way contingency table, calculate expected counts, and use a Chi-Square Test for Homogeneity or Independence to determine if an association exists between two categorical variables.

Part 3: Inference for One Mean

- I can describe the properties of a t-distribution, explain how it changes as degrees of freedom ($df = n - 1$) increase, and find critical t^* values.
- I can construct and interpret a confidence interval for a population mean (μ), checking for normality using the sample size rules ($n \geq 30$, or examining a graph of the sample data for outliers/skewness).
- I can perform a one-sample significance test for a population mean and state a statistically valid conclusion.
- I can distinguish between independent samples and paired data, and perform paired t-procedures to test for a population mean difference (μ_d).

Part 4: Inference for Two Means

- I can verify conditions, compute standard errors, and construct a confidence interval for the difference between two independent population means ($\mu_1 - \mu_2$).
- I can execute a two-sample t-test to compare the centers of two independent quantitative distributions.
- I can verify the structural conditions (LINER) required for inference in a linear regression model.
- I can construct a confidence interval and perform a significance test for the true population slope (β) using standard linear regression computer output.

Prerequisite Skills

Parts 1 and 2: Inference for Proportions

- Calculating relative frequencies out of subset pools rather than the grand total.
- Instant fluency with the Large Counts condition ($np \geq 10$ and $n(1 - p) \geq 10$) from Unit 2.
- Solving for an unknown variable embedded beneath a radical sign or within a fraction denominator.

Parts 3 and 4: Inference for Means

- Recalling that sample means approach normality when $n \geq 30$, regardless of population shape.
- Quick identification of extreme skew or outliers from dotplots, boxplots, or normal probability plots.
- Locating the intercept, slope coefficient, and standard error lines within a standard software regression readout table.

Common Misconceptions

Proportions:

- Students write: "There is a 95% probability that the true proportion of voters is between 0.42 and 0.48."
 - Once an interval is calculated, its endpoints are fixed. There is no probability involved. The correct interpretation is: "We are 95% confident that the interval from 0.42 to 0.48 captures the true population proportion."
- Students define a P-value as the probability that the null hypothesis is true.
 - A P-value assumes the null hypothesis is true and measures the probability of obtaining our observed sample data (or something more extreme) purely by chance.
- When writing out Chi-Square tests, students plug raw percentages or observed counts into the expected-values rows. They must use the formula $(\text{row total} \times \text{column total}) \div (\text{table total})$ and leave expected entries as decimals.

Means:

- If a dataset has two columns of numbers, students automatically default to a 2-sample t-test.
 - Train them to ask: Are these groups independent, or are these two measurements taken on the same individual (or matched pairs)? If paired, they must subtract the columns first and run a 1-sample test on the single list of differences.
- Students use normal distribution critical values (z^*) instead of Student's t distribution values (t^*) when analyzing quantitative averages.
 - Remind them of the mantra: When you don't know the population standard deviation (σ) and must estimate it using the sample standard deviation (s), you must use t .

Current Unit Content/Skills	Spiral Focus	Activity
<i>Part 1: Inference for One Proportion</i> - Constructing Intervals - Performing Significance Tests - Evaluating Type I/II errors <i>Part 2: Two Proportions and Two-Way Tables</i> - Running 2-proportion z-tests. - Calculating Chi-Square statistics - Assessing independence <i>Part 3: Inference for One Mean</i> - Building t-intervals - Executing 1-sample t-tests - Evaluating paired data <i>Part 4: Inference for Two Means</i> - Performing 2-sample t-tests - Modeling independent groups - Evaluating regression slopes.	<i>Part 1: Inference for One Proportion</i> - Unit 2 - Binomial conditions and sampling distributions - Verifying safety via standard deviations and expected successes. <i>Part 2: Two Proportions and Two-Way Tables</i> - Unit 1 - Contingency tables and relative frequencies - Evaluating conditional percentages visually vs. mathematically. <i>Part 3: Inference for One Mean</i> - Unit 1 - Outliers and Graph Audits (CUSS & BS) - Checking normality using visual data sets. <i>Part 4: Inference for Two Means</i> - Unit 1 - Two-Variable Relationships - Interpreting computer output parameters	<i>Part 1: Inference for One Proportion</i> The Standard Error Connection - When calculating the standard error for a 1-proportion z-interval, give students 3 minutes to justify how this formula stems directly from the standard deviation of a sample proportion distribution. <i>Part 2: Two Proportions and Two-Way Tables</i> From Graph to Test Pivot - Present students with a segmented bar chart showing two categorical variables from Unit 1. Have them write a 4-minute hypothesis set (H_0 / H_a) and use the visual breakdown to predict whether a Chi-Square test will reject or fail to reject the null. <i>Part 3: Inference for One Mean</i> The Normality Audit Drill - When teaching the Normal condition for a 1-sample t-test with a small sample size ($n < 30$), display a boxplot or stemplot of the sample data. Have

		students spend 3 minutes evaluating the distribution's shape, specifically testing for strong skewness or outliers that could invalidate t-procedures. <i>Part 4: Inference for Two Means</i> • The Slope Diagnostic ○ When introducing inference for regression slope, project a standard computer output block from Super-Unit 1. Have students identify the slope estimate (b) and the standard deviation of the residuals (s), and connect them to the calculation of the standard error of the slope coefficient.
--	--	--

Assessment

Formative Assessment	Summative Assessment
• Homework • Lesson Checks • Quizzes • Exit Tickets • Lesson Reflections • Performance Tasks • AP Classroom Progress Checks	• Topic Tests • Unit 3 Benchmark (AP Classroom)

Resources

Key Resources	Supplemental Resources
AP Classroom AP Statistics CED Pacing Guide	iXL Delta Math Desmos Khan Academy Math Medic Skew the Script Teacher-made worksheets

Career Readiness, Life Literacies, and Key Skills

9.4.12.CI.1	Demonstrate the ability to reflect, analyze, and use creative skills and ideas (e.g., 1.1.12.prof.CR3a).
9.4.12.CI.3	Investigate new challenges and opportunities for personal growth, advancement, and transition (e.g., 2.1.12.PGD.1).
9.2.12.CAP.4	Evaluate different careers and develop various plans (e.g., costs of public, private, training schools) and timetables for achieving them, including educational/training requirements, costs, loans, and debt repayment.
9.4.12.CT.2	Explain the potential benefits of collaborating to enhance critical thinking and problem solving (e.g., 1.3E.12.prof.CR3.a).
9.4.12.CT.3	Enlist input from a variety of stakeholders (e.g., community members, experts in the field) to design a service learning activity that addresses a local or global issue (e.g., environmental justice).
9.2.12.CAP.21	Explain low-cost and low-risk ways to start a business.
9.4.12.IML.5	Evaluate, synthesize, and apply information on climate change from various sources appropriately (e.g., 2.1.12.CHSS.6, S.IC.B.4, S.IC.B.6, 8.1.12.DA.1, 6.1.12.GeoHE.14.a, 7.1.AL.PRSNT.2).
9.4.12.IML.7	Develop an argument to support a claim regarding a current workplace or societal/ethical issue such as climate change (e.g., NJSLSA.W1, 7.1.AL.PRSNT.4).

Interdisciplinary Connections

RL.CR.11–12.1	Accurately cite strong and thorough textual evidence and make relevant connections to strongly support a comprehensive analysis of multiple aspects of what a literary text says explicitly and inferentially, as well as interpretations of the text; this may include determining where the text leaves matters uncertain.
---------------	--

W.IW.11–12.2	Write informative/explanatory texts (including the narration of historical events, scientific procedures/experiments, or technical processes) to examine and convey complex ideas, concepts, and information clearly and accurately through the effective selection, organization, and analysis of content.
SL.PI.11–12.4	Present information, findings and supporting evidence clearly, concisely, and logically. The content, organization, development, and style are appropriate to task, purpose, and audience.
6.1.12.EconET.14.b	Analyze economic trends, income distribution, labor participation (i.e., employment, the composition of the work force), and government and consumer debt and their impact on society.
HS-LS2-6	Evaluate the claims, evidence, and reasoning that the complex interactions in ecosystems maintain relatively consistent numbers and types of organisms in stable conditions, but changing conditions may result in a new ecosystem.
HS-ETS1-3	Evaluate a solution to a complex real-world problem based on prioritized criteria and trade-offs that account for a range of constraints, including cost, safety, reliability, and aesthetics, as well as possible social, cultural, and environmental impacts.