

AP STATISTICS STUDY GUIDE

180 High-Yield Concepts

Advanced Placement Course | College Board CED-aligned

- All 9 CED units: data exploration, correlation, collecting data, probability, sampling distributions, and all four inference units
- Every formula: z-scores, the Empirical Rule, CLT, confidence intervals, significance tests, chi-square, and slope inference
 - Exam weights on every unit; real exam format on the how-to page
 - Perfect for the AP Statistics exam in May

TemplateForge | \$19

PDF Workbook | templateforge.com | (c) TemplateForge

How to Use This Guide

This guide condenses the AP Statistics course and exam description into high-yield concepts. It is not a textbook replacement -- it is the fastest way to review everything the exam covers and to target your weak units.

The AP Stats exam is 3 hours: 40 multiple-choice questions (90 min, 50% of the score) and 6 free-response questions (90 min, 50% of the score). The 9 units below map to the College Board Course and Exam Description, with unit weights on each section title.

How to use the concepts

- Each concept = term + concise explanation + high-yield exam tip (+ example where useful).
- First pass: read every concept in order. Underline anything unfamiliar.
- Second pass: focus only on the concepts you underlined.
- Two weeks before the exam: rapid-review every tip line (they are the highest-yield).
- Use the Table of Contents to jump straight to your weakest units.

Table of Contents

1	Unit 1: Exploring One-Variable Data	20 concepts
2	Unit 2: Exploring Two-Variable Data	20 concepts
3	Unit 3: Collecting Data	20 concepts
4	Unit 4: Probability, Random Variables & Probability Distributions	20 concepts
5	Unit 5: Sampling Distributions	20 concepts
6	Unit 6: Inference for Categorical Data: Proportions	20 concepts
7	Unit 7: Inference for Quantitative Data: Means	20 concepts
8	Unit 8: Inference for Categorical Data: Chi-Square	20 concepts
9	Unit 9: Inference for Quantitative Data: Slopes	20 concepts

Total: 180 high-yield concepts

How to use the TOC: work through the units in order on your first pass, then jump back to your weakest units for review. Each concept includes an Explanation, a high-yield exam tip, and where useful an Example. Mark the concepts you miss with a small dot in the margin and revisit them 48 hours later.

Unit 1: Unit 1: Exploring One-Variable Data

15-23% of the AP Stats exam. Describing distributions: shape, center, spread, and outliers.

1 Categorical vs Quantitative Variables

Explanation: Categorical variables place individuals into groups (color, state, yes/no). Quantitative variables take numerical values where arithmetic makes sense (height, income, test score).

AP Stats Tip: Identify the variable type *FIRST* -- it determines every graph and statistic you can use.

Example: Zip codes are categorical (the numbers are labels); age in years is quantitative.

2 Frequency Tables & Relative Frequency

Explanation: A frequency table counts how many individuals fall in each category; relative frequency shows the proportion or percentage. Both summarize categorical data.

AP Stats Tip: $\text{Relative frequency} = \text{count} / \text{total}$. Convert counts to percentages for comparisons across groups.

Example: If 12 of 40 students prefer dogs, the relative frequency is 0.30, or 30%.

3 Bar Charts & Pie Charts

Explanation: Bar charts show counts or proportions for categorical data with separated bars (order is arbitrary). Pie charts show parts of a whole; they mislead when categories overlap or are close in size.

AP Stats Tip: Use bar charts for categorical data; do NOT use histograms or line graphs for categories.

Example: A bar chart comparing favorite-movie genres across 3 classes is easy to read; a pie chart with 12 similar slices is not.

4 Dotplots & Stem Plots

Explanation: Dotplots show each observation as a dot above a number line -- best for small data sets. Stem-and-leaf plots retain the actual data values while showing shape.

AP Stats Tip: Dotplots and stem plots are for small quantitative data sets; they reveal gaps and clusters.

Example: A dotplot of 20 quiz scores instantly shows the mode (the tallest stack) and any gaps.

5 Histograms

Explanation: Histograms group quantitative data into intervals (bins) and draw bars with height equal to the count in each bin. The choice of bin width changes the picture.

AP Stats Tip: Histograms: no gaps between bars (quantitative, continuous intervals). Change bins to change the story.

Example: Income data with \$10,000 bins looks different from the same data with \$50,000 bins.

6 Shape: Symmetric, Skewed, Unimodal, Bimodal

Explanation: A distribution is symmetric if the halves mirror. Skewed right: long tail to the right (mean > median). Skewed left: long tail to the left (mean < median). Unimodal has one peak; bimodal has two.

AP Stats Tip: Skew direction = the direction of the LONG TAIL. In skewed-right data, the mean is pulled toward the tail.

Example: Income distributions are skewed right: most earn little, a few earn a lot, pulling the mean above the median.

7 Center: The Mean

Explanation: The mean ($\bar{x}$) is the arithmetic average: sum of values divided by n . It is the balance point of the distribution and is sensitive to outliers.

AP Stats Tip: $\text{Mean} = \text{sum}/n$. Use it for roughly symmetric data without outliers.

Example: The mean of 2, 4, 6 is 4; add a 100 and the mean jumps to 28.

8 Center: The Median

Explanation: The median is the middle value when data are ordered (the 50th percentile). It resists outliers -- a robust measure of center.

AP Stats Tip: *Median = the middle; resistant to outliers. Use it for skewed data.*

Example: In 2, 4, 6, 100, the median is 5 while the mean is 28 -- very different stories.

9 Center: The Mode

Explanation: The mode is the most frequent value. It is the only measure of center usable for categorical data and for describing peaks of a distribution.

AP Stats Tip: *Mode = most common value; for categorical or bimodal data.*

Example: The mode of 3, 3, 4, 5, 5, 5, 6 is 5.

10 Spread: Range & Interquartile Range (IQR)

Explanation: Range = max - min (sensitive to outliers). IQR = Q3 - Q1, the middle 50% of the data (resistant to outliers).

AP Stats Tip: *Range is simple but outlier-prone; IQR is the robust spread measure.*

Example: Data 1, 2, 3, 4, 100: range = 99, IQR = 4 - 2 = 2.

11 Spread: Standard Deviation & Variance

Explanation: Standard deviation (s) measures typical distance from the mean. Variance = s^2 . Both are nonresistant to outliers.

AP Stats Tip: *$s = \sqrt{\text{sum}((x - \bar{x})^2)/(n-1)}$. Larger s = more spread.*

Example: Quiz scores 80, 80, 80 have s = 0; adding one 50 makes s jump.

12 Outliers & the 1.5 x IQR Rule

Explanation: An outlier lies below $Q1 - 1.5 \times \text{IQR}$ or above $Q3 + 1.5 \times \text{IQR}$. The rule flags values the boxplot would draw separately.

AP Stats Tip: *Outlier check: boundaries = $Q1 - 1.5(\text{IQR})$ and $Q3 + 1.5(\text{IQR})$. Compute IQR first.*

Example: With $Q1 = 10$ and $Q3 = 20$ (IQR = 10), any value above 35 or below -5 is an outlier.

13 The Five-Number Summary

Explanation: The five-number summary: minimum, Q1, median, Q3, maximum. It gives a complete picture of center and spread and feeds the boxplot.

AP Stats Tip: *Five numbers: min, Q1, median, Q3, max. Always list them in order.*

Example: For 1, 3, 5, 7, 9, 11, 13: min=1, Q1=3, med=7, Q3=11, max=13.

14 Boxplots

Explanation: A boxplot displays the five-number summary: a box from Q1 to Q3 with a median line, whiskers to the most extreme non-outliers, and dots for outliers.

AP Stats Tip: *Boxplots compare distributions side by side; whiskers stop at the last non-outlier.*

Example: A boxplot with a long right whisker and dots beyond it signals right skew and outliers.

15 Comparing Distributions

Explanation: When comparing distributions, always address shape, center, spread, and outliers -- in context, with units.

AP Stats Tip: *SOCS: Shape, Outliers, Center, Spread. Compare each, then summarize in context.*

Example: 'The test scores are skewed left with a median of 82 and IQR of 12, with no outliers.'

16 Percentiles & Quartiles

Explanation: The pth percentile is the value below which p% of the data fall. Quartiles split data into fourths: Q1 = 25th percentile, median = 50th, Q3 = 75th.

AP Stats Tip: *Percentile = 'what percent of data are below this value.' Quartiles are specific percentiles.*

Example: Scoring at the 90th percentile means you beat 90% of test takers.

17 Standardized Scores (z-scores)

- Explanation:** A z-score measures how many standard deviations a value is from the mean: $z = (x - \mu)/\sigma$. Positive = above mean; negative = below.
- AP Stats Tip:** $z = (x - \mu)/\sigma$. Z-scores allow comparing values from different distributions.
- Example:** A 90 on a test with mean 80 and $s = 5$ gives $z = 2$ -- two SDs above average.
-

18 The Empirical Rule (68-95-99.7)

- Explanation:** For mound-shaped (approximately normal) distributions: about 68% of data within 1 SD of the mean, 95% within 2 SDs, 99.7% within 3 SDs.
- AP Stats Tip:** Use the Empirical Rule ONLY for approximately normal data; it converts SDs to percentages.
- Example:** Mean 100, SD 15: about 95% of scores fall between 70 and 130.
-

19 Density Curves

- Explanation:** A density curve models the distribution of a variable: the total area under it is 1, and areas represent proportions. The normal curve is the best-known density curve.
- AP Stats Tip:** Density curve: area = proportion. The mean is the balance point; the median splits the area in half.
- Example:** In a skewed-right density curve, the mean lies to the right of the median.
-

20 Choosing the Right Display & Summary

- Explanation:** Match displays and summaries to variable type and shape: categorical -> bar/pie; quantitative -> histogram/boxplot; skewed -> median + IQR; symmetric -> mean + SD.
- AP Stats Tip:** The AP exam rewards correct choices: skewed data -> median/IQR; symmetric -> mean/SD.
- Example:** Report median and IQR for home prices (skewed right), mean and SD for heights (symmetric).
-

Unit 2: Unit 2: Exploring Two-Variable Data

5-7% of the AP Stats exam. Correlation, regression, and categorical associations.

21 Scatterplots

Explanation: Scatterplots show the relationship between two quantitative variables: one point per individual, with explanatory variable on the x-axis and response on the y-axis.

AP Stats Tip: *Always describe direction, form, strength, and outliers on a scatterplot.*

Example: Plotting years of education vs income shows the association's direction and strength.

22 Explanatory & Response Variables

Explanation: The explanatory (independent) variable explains or predicts; the response (dependent) variable is what we measure the outcome on. Association does NOT prove causation.

AP Stats Tip: *Identify explanatory vs response from the question's wording -- it drives the graph and analysis.*

Example: Study time (explanatory) vs exam score (response) -- we predict score from study time.

23 Correlation (r)

Explanation: Correlation r measures the strength and direction of a LINEAR relationship. r ranges from -1 to +1; sign = direction, $|r|$ = strength. r has no units.

AP Stats Tip: *r is unitless and only measures linear association. $r = 0$ means no linear relationship.*

Example: $r = 0.9$ is a strong positive linear association; $r = -0.4$ is a weak negative one.

24 Properties of Correlation

Explanation: r is unaffected by changing units or by which variable is x or y . r is NOT resistant: one outlier can inflate or deflate it. r does not capture curved relationships.

AP Stats Tip: *Correlation is not resistant; always check for outliers and curvature before trusting r .*

Example: Add one extreme point to a tight cluster and r can drop from 0.9 to 0.3.

25 Correlation vs Causation

Explanation: Association does NOT imply causation. A lurking variable may drive both variables, or the relationship may be coincidence or reversed causality.

AP Stats Tip: *The exam loves this: state 'association does not prove causation' and suggest lurking variables.*

Example: Ice cream sales and drowning both rise in summer -- a lurking variable (heat), not causation.

26 The Least-Squares Regression Line

Explanation: The LSRL minimizes the sum of squared vertical distances from the data to the line: $\hat{y} = a + bx$. It is the best linear predictor.

AP Stats Tip: *LSRL = minimize squared residuals. Report the equation in context with units.*

Example: $\hat{y} = 30 + 5x$: each additional x predicts 5 more units of y .

27 Interpreting Slope & Intercept

Explanation: The slope b is the predicted change in y for a one-unit increase in x . The intercept a is the predicted y when $x = 0$ (may be meaningless).

AP Stats Tip: *Interpret slope and intercept in CONTEXT with units -- never just 'the slope is 5.'*

Example: 'For each extra hour studied, predicted score rises by 5 points.'

28 Residuals & Residual Plots

Explanation: A residual = observed y - predicted $\hat{y}$. Residual plots show leftover patterns: random scatter = linear model fits; a curve or funnel = model is wrong.

AP Stats Tip: *Residual = actual - predicted. A residual plot with no pattern supports the linear model.*

Example: If residual plots curve, a linear model is inappropriate -- try a transformation.

29 The Coefficient of Determination (r^2)

Explanation: r^2 is the fraction of the variation in y explained by the linear relationship with x . It ranges 0 to 1; $r^2 = 0.81$ means 81% of the variation in y is explained by x .

AP Stats Tip: *r^2 = the proportion of variation explained. State it in context: ' r^2 of the model explains X% of the variability in [y].'*

Example: If $r = 0.9$, then $r^2 = 0.81$ -- 81% of score variation is explained by study time.

30 Extrapolation

Explanation: Predicting y for x -values outside the range of the data is risky: the relationship may change beyond the observed data.

AP Stats Tip: *Never extrapolate beyond the data range -- the model may not hold there.*

Example: Predicting a 200-year-old's height from a growth model fit to children is absurd extrapolation.

31 Influential Points & Outliers

Explanation: An outlier has a large residual (does not follow the pattern). An influential point is far in x and pulls the line toward it. Influential points may or may not be outliers.

AP Stats Tip: *Influential = extreme in x AND affects the slope. Removing it changes the line a lot.*

Example: One billionaire in an income study pulls the regression line up -- highly influential.

32 Lurking Variables & Confounding

Explanation: A lurking variable is related to both x and y , creating a spurious association. Confounding occurs when its effects are mixed with the explanatory variable's.

AP Stats Tip: *Always consider lurking variables before claiming causation from observational data.*

Example: Coffee drinking and heart disease may both relate to stress -- stress is the lurking variable.

33 Transformations to Achieve Linearity

Explanation: When the relationship is curved, transform a variable ($\log y$, $\log x$, square root) to make the data linear, then fit the LSRL on the transformed scale.

AP Stats Tip: *Curved data $\rightarrow$ transform $\rightarrow$ linear. Exponential growth becomes linear under $\log(y)$.*

Example: Population growth over time (exponential) becomes a straight line when $\log(\text{population})$ is plotted.

34 Nonlinear Models: Exponential & Power

Explanation: Exponential growth: $y = a(b^x)$ -- constant percent change. Power models: $y = a(x^b)$. Both become linear with appropriate logs.

AP Stats Tip: *Exponential = constant multiplier; $\log(y)$ vs x is linear. Power = $\log(y)$ vs $\log(x)$ is linear.*

Example: A population doubling every 10 years is exponential; $\log(\text{population})$ vs time is a line.

35 Two-Way Tables

Explanation: Two-way tables summarize the relationship between two categorical variables. Rows and columns are the categories; entries are counts.

AP Stats Tip: *Two-way tables = categorical $\times$ categorical. Compute marginal and conditional distributions from them.*

Example: A table of gender (rows) by pass/fail (columns) shows the association between the two.

36 Marginal & Conditional Distributions

Explanation: Marginal distributions are the totals for one variable alone (row or column totals). Conditional distributions fix the value of one variable and look at the other.

AP Stats Tip: *Marginal = totals; conditional = one variable given a category of the other. Compare conditionals to find association.*

Example: Comparing the pass rates of men and women (conditional distributions) reveals association.

37 Simpson's Paradox

Explanation: A trend that appears in several groups can reverse when the groups are combined. It happens when a lurking variable (group size) drives the overall result.

AP Stats Tip: *Simpson's paradox = direction reverses after combining groups. Identify the lurking grouping variable.*

Example: Hospital A has a lower overall survival rate but higher survival in every department -- its sicker patients skew the total.

38 Segmented Bar Charts & Mosaic Plots

Explanation: Segmented (stacked) bar charts show conditional distributions as bars divided into segments. Mosaic plots add width proportional to group size.

AP Stats Tip: *Use segmented bars to compare conditional distributions visually.*

Example: A stacked bar of pass rates by gender shows at a glance whether association exists.

39 Choosing a Model for Two-Variable Data

Explanation: Match the tool to the variable types: two quantitative -> scatterplot + correlation + regression; two categorical -> two-way table + conditional distributions.

AP Stats Tip: *Variable types determine the entire analysis path -- start there.*

Example: Quantitative x and y: report r, the LSRL, and r^2 . Categorical x and y: report conditional percentages.

40 Outliers in Regression

Explanation: Outliers in regression are points with large residuals; influential points have high leverage (extreme x). Both can distort the LSRL.

AP Stats Tip: *Check both residual size and leverage; report their effect on slope and r^2 .*

Example: A data-entry error at x = 100 in data ranging 0-10 can flip the slope's sign.

Unit 3: Unit 3: Collecting Data

12-15% of the AP Stats exam. Sampling, experiments, and bias.

41 The Purpose of Sampling

Explanation: Sampling aims to estimate population parameters from sample statistics. A good sample is representative; a bad one is biased.

AP Stats Tip: *Sample -> statistics; population -> parameters. Representativeness is everything.*

Example: Estimating the average height of all US adults from a sample of 1,000 requires a representative sample.

42 Census vs Sample

Explanation: A census attempts to measure every individual in the population. It is accurate in principle but expensive, slow, and prone to nonresponse -- samples are usually better.

AP Stats Tip: *Census = everyone; sample = a subset. The exam asks why censuses are rarely practical.*

Example: The US census misses millions each decade despite enormous effort.

43 Convenience Sampling & Bias

Explanation: Convenience sampling picks whoever is easy to reach (mall intercepts, online polls). It is biased because the easy-to-reach are not representative.

AP Stats Tip: *Convenience = easy but biased. Name the specific bias: undercoverage of hard-to-reach groups.*

Example: A mall survey overrepresents shoppers and misses people who never visit malls.

44 Voluntary Response Sampling

Explanation: Voluntary response lets individuals self-select (call-in polls, online surveys). It overrepresents people with strong opinions -- always biased.

AP Stats Tip: *Voluntary response = self-selection bias; strong-opinion people respond.*

Example: A radio call-in poll on gun control attracts activists, not the public.

45 Simple Random Samples (SRS)

Explanation: An SRS gives every group of n individuals an equal chance of selection (e.g., random digits, random number generator). It eliminates selection bias.

AP Stats Tip: *SRS = every possible sample of size n is equally likely. Use random digits to select.*

Example: Numbering all 400 students and using a random-number table to pick 40 is an SRS.

46 Stratified Random Sampling

Explanation: Divide the population into homogeneous groups (strata), then take an SRS within each stratum. It guarantees representation of each group and reduces variability.

AP Stats Tip: *Stratified = group then sample each group. Use when subgroups differ.*

Example: Split voters by state, then randomly sample within each state -- every state is represented.

47 Cluster Sampling

Explanation: Divide the population into clusters (naturally occurring groups), randomly select clusters, and measure EVERYONE in the chosen clusters. Useful when a list of individuals is hard to build.

AP Stats Tip: *Cluster = select groups, include all members. Stratified samples within groups; clusters sample whole groups.*

Example: Randomly select 5 of 50 schools and survey every student in those schools.

48 Systematic Sampling

Explanation: Select every k th individual from an ordered list after a random start. It is convenient but can be biased if the list has a pattern.

AP Stats Tip: *Systematic = every k th after a random start; watch for periodicity in the list.*

Example: Sampling every 10th name from an alphabetized roster is systematic sampling.

49 Multistage Sampling

Explanation: Real surveys combine methods: random counties, then random schools within, then random students. Each stage narrows the population.

AP Stats Tip: *Multistage = several random stages, often mixing cluster and stratified ideas.*

Example: The National Assessment samples states, then districts, then schools, then students.

50 Bias: Undercoverage

Explanation: Undercoverage occurs when some groups are left out of the sampling frame (no cell phones, no internet, homeless). The sample cannot represent them.

AP Stats Tip: *Undercoverage = the frame misses part of the population. Identify who is missing.*

Example: A landline-only poll undercovers young adults who only use cell phones.

51 Bias: Nonresponse

Explanation: Nonresponse occurs when selected individuals do not participate. If nonresponders differ from responders, the sample is biased.

AP Stats Tip: *Nonresponse = chosen people refuse or cannot be reached. Higher response = less bias.*

Example: A survey with an 8% response rate likely overrepresents engaged, opinionated people.

52 Bias: Response & Wording

Explanation: Response bias: leading questions, sensitive topics, or interviewer effects distort answers. Question wording can push respondents to a particular answer.

AP Stats Tip: *Watch for leading questions and sensitive topics -- both create response bias.*

Example: 'Do you support wasteful government spending?' invites 'no' regardless of the program.

53 Observational Studies

Explanation: Observational studies observe and measure without assigning treatments. They can identify association but cannot establish causation (confounding lurks).

AP Stats Tip: *Observational = no assignment; association only. This is a core limitation.*

Example: Comparing health of joggers vs non-joggers is observational -- fit people may jog, not the reverse.

54 Experiments

Explanation: Experiments actively impose a treatment to observe its effect. Random assignment of subjects to treatments is what allows causal conclusions.

AP Stats Tip: *Experiments + random assignment = causation. This is the gold standard.*

Example: Randomly assigning patients to drug or placebo lets the researcher claim the drug caused the difference.

55 Random Assignment & Control

Explanation: Random assignment balances lurking variables across groups; control keeps other conditions constant. Together they isolate the treatment's effect.

AP Stats Tip: *Random assignment = balanced groups; control = constant conditions. Both are required.*

Example: Giving some plants fertilizer and others none, with identical light and water, controls the environment.

56 Blinding & Placebos

Explanation: Blinding hides the treatment from subjects (single-blind) or from subjects and researchers (double-blind). Placebos control for the psychological effect of receiving treatment.

AP Stats Tip: *Double-blind = neither subjects nor evaluators know the treatment. Placebo = inert treatment.*

Example: A double-blind drug trial gives half the patients sugar pills and no one knows which is which.

57 Blocking & Matched Pairs

Explanation: Blocks are groups of similar subjects (by gender, age) within which treatments are randomly assigned -- reducing variability. Matched pairs = each subject gets both treatments, or subjects are paired.

AP Stats Tip: *Blocking controls for known sources of variation; matched pairs are a special case.*

Example: Testing two fertilizers on paired plots of soil, or each subject on both diets, uses blocking.

58 Replication

Explanation: Replication = applying each treatment to many subjects so that chance variation averages out. More subjects = more precise estimates and more reliable conclusions.

AP Stats Tip: *Replication = enough subjects per treatment. Without it, results are anecdotal.*

Example: A study of 5 subjects cannot distinguish the treatment from luck; 500 can.

59 Scope of Inference

Explanation: If subjects are randomly assigned, we can infer causation (for those subjects). If subjects are randomly sampled, we can generalize to the population. Both are needed for full inference.

AP Stats Tip: *Random assignment -> causation; random sample -> generalization. State which applies.*

Example: A convenience sample with random assignment allows causal claims for the sample but not the population.

60 Designing an Experiment: The Components

Explanation: A well-designed experiment has: comparison of treatments, random assignment, control of lurking variables, and replication. These four components make results trustworthy.

AP Stats Tip: *Comparison + randomization + control + replication = valid experiment. Name all four in design questions.*

Example: A clinical trial comparing drug vs placebo with 1,000 randomized patients has all four components.

Unit 4: Unit 4: Probability, Random Variables & Probability Distributions

10-20% of the AP Stats exam. Probability rules, random variables, and distributions.

61 The Law of Large Numbers

Explanation: As the number of trials increases, the relative frequency of an outcome approaches its theoretical probability.

AP Stats Tip: *LLN: more trials → relative frequency converges to probability.*

Example: Tossing a coin 10 times may give 70% heads; 10,000 tosses will be near 50%.

62 Probability Models & Sample Spaces

Explanation: A probability model lists all possible outcomes (the sample space) and assigns each a probability between 0 and 1, summing to 1.

AP Stats Tip: *Probabilities must be between 0 and 1 and sum to 1 for the sample space.*

Example: Rolling a die: each face has probability $1/6$; the six probabilities sum to 1.

63 Basic Probability Rules

Explanation: $P(A)$ is between 0 and 1. $P(A \text{ or } B)$ uses the addition rule; $P(A \text{ and } B)$ uses the multiplication rule. The probability of the empty event is 0; of the whole space, 1.

AP Stats Tip: *Always check: $0 \leq P \leq 1$, sum = 1. Then choose the correct rule.*

Example: If $P(A) = 0.3$ and $P(B) = 0.4$ with A, B disjoint, $P(A \text{ or } B) = 0.7$.

64 The Complement Rule

Explanation: $P(A^c) = 1 - P(A)$. The complement is everything NOT in A. Often easier to compute 'at least one' via the complement.

AP Stats Tip: *$P(\text{at least one}) = 1 - P(\text{none})$. This is the most used complement trick.*

Example: $P(\text{at least one head in 3 tosses}) = 1 - P(\text{no heads}) = 1 - (1/2)^3 = 7/8$.

65 The Addition Rule

Explanation: For any events: $P(A \text{ or } B) = P(A) + P(B) - P(A \text{ and } B)$. If A and B are mutually exclusive, $P(A \text{ and } B) = 0$.

AP Stats Tip: *General addition rule subtracts the overlap; disjoint events need no subtraction.*

Example: $P(\text{ace}) = 4/52$, $P(\text{heart}) = 13/52$, $P(\text{ace of hearts}) = 1/52$; $P(\text{ace or heart}) = 16/52$.

66 The Multiplication Rule & Independence

Explanation: For independent events: $P(A \text{ and } B) = P(A) \times P(B)$. Independence means knowing A happened does not change $P(B)$.

AP Stats Tip: *Independence: $P(A \text{ and } B) = P(A)P(B)$. Check independence with this rule or by definition.*

Example: Two fair dice: $P(\text{double six}) = (1/6)(1/6) = 1/36$ -- the rolls are independent.

67 Conditional Probability

Explanation: $P(A | B) = P(A \text{ and } B) / P(B)$: the probability of A given that B occurred. It conditions on new information.

AP Stats Tip: *Conditional = restrict the sample space to B. Formula: $P(A|B) = P(A \text{ and } B)/P(B)$.*

Example: $P(\text{female} | \text{honor student})$ uses only the honor-student row of a two-way table.

68 Independence vs Disjoint Events

Explanation: Disjoint events cannot both happen ($P(A \text{ and } B) = 0$). Independent events do not affect each other. Disjoint events are NEVER independent (unless one has probability 0).

AP Stats Tip: *Disjoint = no overlap; independent = no influence. Getting hit by a car and struck by lightning are not independent of each other's probabilities in a Venn sense.*

Example: Drawing one card: 'ace' and 'king' are disjoint; if you drew an ace, it is certainly not a king.

69 Tree Diagrams

Explanation: Tree diagrams organize multi-stage probabilities: branches show conditional probabilities at each stage; multiply along paths and add across paths.

AP Stats Tip: *Multiply along branches (and), add across branches (or).*

Example: A two-branch tree for a test's sensitivity and specificity computes $P(\text{positive} \mid \text{disease})$ paths.

70 Random Variables & Expected Value

Explanation: A random variable assigns a number to each outcome. Its expected value $E(X)$ = sum of $x P(x)$ -- the long-run average. It is the mean of the distribution.

AP Stats Tip: *$E(X)$ = sum of (value $\times$ probability). Interpret as the long-run average.*

Example: A fair game pays \$2 on heads and \$0 on tails: $E(X) = 0.5(2) + 0.5(0) = \1 .

71 Variance & Standard Deviation of a Random Variable

Explanation: $\text{Var}(X)$ = sum of $(x - \mu)^2 P(x)$; $\text{SD}(X) = \sqrt{\text{Var}}$. They measure the spread of the distribution around the mean.

AP Stats Tip: *Compute variance as the weighted average of squared deviations from the mean.*

Example: A game paying \$0 or \$10 each with $p = 0.5$ has $E(X) = 5$ and $\text{SD} = 5$.

72 Combining Random Variables

Explanation: For independent X and Y : $E(X \pm Y) = E(X) \pm E(Y)$. $\text{Var}(X \pm Y) = \text{Var}(X) + \text{Var}(Y)$. SDs do NOT add -- variances do.

AP Stats Tip: *Variances add for sums and differences of independent variables; SD = sqrt of the sum of variances.*

Example: If X has SD 3 and Y has SD 4, $\text{SD}(X + Y) = 5$, not 7.

73 The Binomial Setting

Explanation: A binomial situation requires: fixed number n of trials, two outcomes per trial (success/failure), constant probability p , and independent trials.

AP Stats Tip: *BINS: Binary outcomes, Independent trials, fixed Number n , constant probability p .*

Example: Counting heads in 10 fair-coin flips is binomial: $n = 10$, $p = 0.5$.

74 The Binomial Distribution

Explanation: $X \sim \text{Binomial}(n, p)$: $P(X = k) = C(n, k) p^k (1-p)^{(n-k)}$. Mean = np . $\text{SD} = \sqrt{np(1-p)}$.

AP Stats Tip: *Know the formula, mean np , and SD $\sqrt{np(1-p)}$.*

Example: $P(\text{exactly 3 heads in 5 flips}) = C(5, 3)(0.5)^3(0.5)^2 = 10/32$.

75 The Geometric Setting

Explanation: A geometric situation counts the number of trials until the FIRST success, with independent trials and constant probability p . $P(X = k) = (1-p)^{(k-1)} p$. Mean = $1/p$.

AP Stats Tip: *Geometric = trials until first success. Mean = $1/p$.*

Example: The probability the first 6 appears on the 3rd die roll: $(5/6)^2(1/6)$.

76 Normal Distributions & Density Curves

Explanation: Normal distributions are symmetric, bell-shaped density curves parameterized by mean μ and standard deviation σ . The total area is 1.

AP Stats Tip: *Normal(μ , σ): symmetric bell. Areas under the curve = probabilities.*

Example: IQ scores $\sim \text{Normal}(100, 15)$: the curve is centered at 100 with SD 15.

77 The Standard Normal Distribution

Explanation: The standard normal has mean 0 and SD 1. Convert any normal value to $z = (x - \mu)/\sigma$, then use the z-table to find areas.

AP Stats Tip: $z = (x - \mu)/\sigma$ is the bridge to the z-table.

Example: A score of 130 with mean 100 and SD 15 gives $z = 2$ -- look up 2.00 for the area below.

78 Normal Calculations

Explanation: To find a percentile: standardize and use the z-table. To find a boundary: look up the z for the area, then $x = \mu + z \sigma$.

AP Stats Tip: Percentile $\rightarrow z \rightarrow x$. Always sketch the normal curve and shade the region.

Example: The 90th percentile of Normal(100, 15): $z = 1.28$, so $x = 100 + 1.28(15) = 119.2$.

79 The Normal Approximation to the Binomial

Explanation: When $np \geq 10$ and $n(1-p) \geq 10$, the binomial can be approximated by Normal(np , $\sqrt{np(1-p)}$).

AP Stats Tip: Check $np \geq 10$ and $n(1-p) \geq 10$ before approximating.

Example: With $n = 200$, $p = 0.5$, $np = 100 \geq 10$ -- the normal approximation is safe.

80 The Mean & Standard Deviation of Discrete Distributions

Explanation: For any discrete random variable, compute $\mu = \sum(x P(x))$ and $\sigma^2 = \sum((x - \mu)^2 P(x))$. These summarize the distribution.

AP Stats Tip: Expected value = the center; SD = the spread. Compare two games by both.

Example: A game with $E(X) = 0$ and SD = 10 is riskier than one with $E(X) = 0$ and SD = 2.

Unit 5: Unit 5: Sampling Distributions

7-12% of the AP Stats exam. The distribution of sample statistics.

81 The Sampling Distribution of the Sample Mean

Explanation: The sampling distribution of $\bar{x}$ describes the distribution of sample means over all possible samples. Its mean equals the population mean μ .

AP Stats Tip: $\bar{x}$ is unbiased: mean of the sampling distribution = μ .

Example: If the population mean is 50, the average of all possible sample means is 50.

82 The Central Limit Theorem (CLT)

Explanation: For large samples ($n \geq 30$, or any n if the population is normal), the sampling distribution of $\bar{x}$ is approximately Normal with mean μ and $SD = \sigma/\sqrt{n}$.

AP Stats Tip: CLT: $\bar{x} \sim \text{Normal}(\mu, \sigma/\sqrt{n})$ for large n , regardless of the population shape.

Example: Sample means from a skewed population become bell-shaped as n reaches 30.

83 The Sampling Distribution of the Sample Proportion

Explanation: The sampling distribution of $\hat{p}$ has mean p and $SD = \sqrt{p(1-p)/n}$. With $np \geq 10$ and $n(1-p) \geq 10$, it is approximately Normal.

AP Stats Tip: $\hat{p} \sim \text{Normal}(p, \sqrt{p(1-p)/n})$ when the large-counts condition holds.

Example: Polling 1,000 voters with true $p = 0.5$: $SD(\hat{p}) = \sqrt{0.25/1000} = 0.0158$.

84 The Standard Deviation of $\bar{x}$

Explanation: $SD(\bar{x}) = \sigma/\sqrt{n}$. It shrinks as the sample size grows -- larger samples give more precise estimates.

AP Stats Tip: $SD(\bar{x}) = \sigma/\sqrt{n}$; quadruple n to halve the SD .

Example: With $\sigma = 20$, $n = 100$ gives $SD = 2$; $n = 400$ gives $SD = 1$.

85 The Standard Error

Explanation: The standard error estimates the SD of a statistic when σ is unknown: $SE = s/\sqrt{n}$ for means, $\sqrt{\hat{p}(1-\hat{p})/n}$ for proportions.

AP Stats Tip: Use SE when σ is unknown (almost always); it plugs in sample values.

Example: With sample SD $s = 12$ and $n = 36$, $SE(\bar{x}) = 2$.

86 When Is the Sampling Distribution Normal?

Explanation: For proportions: $np \geq 10$ and $n(1-p) \geq 10$. For means: $n \geq 30$ (CLT), or the population is normal. State the condition before using Normal calculations.

AP Stats Tip: Always check and STATE the conditions before computing probabilities for statistics.

Example: A poll of 50 with $p = 0.1$ fails $np \geq 10$ -- the normal approximation is not safe.

87 Bias & Variability of Estimators

Explanation: Bias = systematic error (the statistic is centered away from the parameter). Variability = spread of the statistic across samples. A good estimator is unbiased with low variability.

AP Stats Tip: Bias vs variability: a precise but biased statistic can be worse than a variable unbiased one.

Example: A darts player who always hits the same wrong spot is biased but low-variability.

88 The 10% Condition

Explanation: When sampling without replacement, treat draws as independent only if the sample is less than 10% of the population ($n \leq 0.10N$).

AP Stats Tip: $n \leq 0.10N$ justifies the independence assumption in sampling.

Example: Sampling 200 students from a school of 5,000: $200 < 500$, so the 10% condition holds.

89 Why Sampling Distributions Matter

- Explanation:** Sampling distributions connect sample statistics to population parameters -- they are the foundation of confidence intervals and significance tests.
- AP Stats Tip:** *Every inference procedure rests on the sampling distribution of a statistic.*
- Example:** The margin of error of a poll comes straight from $SD(\hat{p})$.
-

90 Sample Proportions vs Sample Means

- Explanation:** $\hat{p}$ applies to categorical data (np , $n(1-p)$ conditions). $\bar{x}$ applies to quantitative data (CLT or normality condition). Their SDs differ by formula.
- AP Stats Tip:** *Match the statistic to the data type and its conditions.*
- Example:** Polling ($\hat{p}$) uses large-counts; measuring heights ($\bar{x}$) uses the CLT.
-

91 The Effect of Sample Size on Variability

- Explanation:** Larger samples shrink the SD of both $\hat{p}$ and $\bar{x}$ (divided by $\sqrt{n}$). Doubling n does NOT halve the SD -- quadrupling it does.
- AP Stats Tip:** *SD falls with $1/\sqrt{n}$. Bigger n = tighter sampling distribution.*
- Example:** A poll of 4,000 has half the margin of error of a poll of 1,000.
-

92 Unbiased Estimators

- Explanation:** An estimator is unbiased if its sampling distribution is centered at the parameter: $E(\bar{x}) = \mu$ and $E(\hat{p}) = p$. Sample mean and proportion are unbiased.
- AP Stats Tip:** *$\bar{x}$ and $\hat{p}$ are unbiased for μ and p . Sample SD is not (but is close for large n).*
- Example:** Repeated samples give sample means that average out to the true population mean.
-

93 The Sampling Distribution of a Difference

- Explanation:** For independent samples, the difference of two means or proportions has a sampling distribution with mean equal to the difference of parameters and variance equal to the sum of variances.
- AP Stats Tip:** *$Var(\hat{p}_1 - \hat{p}_2) = Var_1 + Var_2$ (independent samples).*
- Example:** Comparing two teaching methods: the sampling distribution of $\bar{x}_1 - \bar{x}_2$ centers at $\mu_1 - \mu_2$.
-

94 The Sampling Distribution in Inference

- Explanation:** Confidence intervals and tests use the sampling distribution's shape, center, and spread -- with the SE substituting for the SD.
- AP Stats Tip:** *Inference = sampling distribution + data. Identify the statistic's distribution first.*
- Example:** A 95% CI for a proportion is built around the Normal sampling distribution of $\hat{p}$.
-

95 The CLT in Practice

- Explanation:** The CLT works for means even from badly skewed populations once $n \geq 30$; for very skewed or outlier-prone data, use larger samples.
- AP Stats Tip:** *$n \geq 30$ is a rule of thumb; more skewed data needs larger n .*
- Example:** Sample medians and sums also become normal under the CLT, not just means.
-

96 The t-Distribution Preview

- Explanation:** When σ is unknown and we use s , the statistic $(\bar{x} - \mu)/(s/\sqrt{n})$ follows a t-distribution with $n - 1$ degrees of freedom -- wider tails than the normal.
- AP Stats Tip:** *Unknown $\sigma \rightarrow t$ with $df = n - 1$. t has fatter tails, shrinking toward normal as n grows.*
- Example:** A sample of 20 with unknown σ uses t with 19 degrees of freedom.
-

97 Finite Populations & Sampling Without Replacement

Explanation: Sampling without replacement slightly reduces variability; the 10% condition keeps the correction negligible.

AP Stats Tip: *If $n > 10\%$ of N , the SD shrinks (finite population correction) and independence is violated.*

Example: Surveying 3,000 of 10,000 people exceeds 10% -- the simple SD formula is off.

98 The Shape of the Sampling Distribution of p-hat

Explanation: p-hat's distribution is approximately Normal when np and $n(1-p)$ are both ≥ 10 ; skewed p (near 0 or 1) needs larger samples.

AP Stats Tip: *Check large counts before normal approximation for p-hat.*

Example: $p = 0.95$ with $n = 50$: $np = 47.5$ but $n(1-p) = 2.5$ -- not safe to approximate.

99 Statistics vs Parameters

Explanation: Parameters (μ , p , σ) describe populations and are usually unknown. Statistics ($\bar{x}$, $\hat{p}$, s) describe samples and are known. Inference uses statistics to estimate parameters.

AP Stats Tip: *Greek letters = parameters; Roman letters = statistics. Keep the notation straight.*

Example: The average height of ALL adults (μ) is a parameter; the average in your sample ($\bar{x}$) is a statistic.

100 The Importance of Random Sampling

Explanation: Random sampling makes the sampling distribution valid: without randomness, no probability model describes the statistic.

AP Stats Tip: *Random selection is the prerequisite for inference about a population.*

Example: A biased convenience sample has no known sampling distribution -- inference is impossible.

Unit 6: Unit 6: Inference for Categorical Data: Proportions

12-15% of the AP Stats exam. Confidence intervals and significance tests for p .

101 Confidence Intervals: The Basics

Explanation: A confidence interval estimates an unknown parameter with a range of plausible values: statistic $\pm$ margin of error. The interval is built around the sampling distribution of the statistic.

AP Stats Tip: $CI = \text{statistic} \pm \text{critical value} \times SE$. Interpret the interval, not the confidence level.

Example: A poll finds 54% approve, with a margin of error of $\pm 3\%$: the interval is 51% to 57%.

102 Interpreting a Confidence Interval

Explanation: A 95% confidence interval for a parameter says we are 95% confident the interval captures the true parameter. It does NOT say 95% of the data fall in the interval.

AP Stats Tip: Interpret: 'We are 95% confident that the true proportion is between a and b .'

Example: 'We are 95% confident the true approval rate is between 51% and 57%.'

103 Interpreting a Confidence Level

Explanation: The confidence level describes the method: if we repeated the sampling process many times, about 95% of the intervals would capture the parameter. The parameter is fixed; the intervals vary.

AP Stats Tip: Confidence level = capture rate of the METHOD over many samples, not the probability for one interval.

Example: In 100 polls, about 95 of the 95% CIs will contain the true proportion.

104 The Critical Value z^*

Explanation: z^* is the number of standard errors for the desired confidence level. For 90%: $z^* = 1.645$. For 95%: $z^* = 1.96$. For 99%: $z^* = 2.576$.

AP Stats Tip: Memorize: 90% $\rightarrow 1.645$, 95% $\rightarrow 1.96$, 99% $\rightarrow 2.576$.

Example: A 95% CI uses $z^* = 1.96$: $p\text{-hat} \pm 1.96 \times \sqrt{p\text{-hat}(1-p\text{-hat})/n}$.

105 The Standard Error of $p\text{-hat}$

Explanation: $SE(p\text{-hat}) = \sqrt{p\text{-hat}(1-p\text{-hat})/n}$ estimates the SD of the sample proportion when p is unknown. It shrinks as n grows.

AP Stats Tip: $SE(p\text{-hat}) = \sqrt{p\text{-hat}(1-p\text{-hat})/n}$. Used in intervals and z -tests.

Example: With $p\text{-hat} = 0.6$ and $n = 100$, $SE = \sqrt{0.6 \times 0.4 / 100} = \sqrt{0.0024} = 0.049$.

106 Conditions for a One-Proportion z -Interval

Explanation: Random sample (or experiment), $n \times p\text{-hat} \geq 10$ and $n(1-p\text{-hat}) \geq 10$ (large counts), and independence (10% condition when sampling without replacement).

AP Stats Tip: State RANDOM, LARGE COUNTS, and 10% before every inference procedure.

Example: A poll of 500 registered voters, randomly selected, with 300 yes: $500 \times 0.6 = 300 \geq 10$ -- conditions met.

107 The One-Proportion z -Interval

Explanation: Formula: $p\text{-hat} \pm z^* \times \sqrt{p\text{-hat}(1-p\text{-hat})/n}$. It gives the plausible range for the true proportion p .

AP Stats Tip: Write the formula, plug in, and interpret in context with units.

Example: $p\text{-hat} = 0.54$, $n = 1000$, 95%: $0.54 \pm 1.96 \times 0.0158 = (0.509, 0.571)$.

108 Margin of Error & Sample Size

Explanation: $ME = z^* \times SE$. To halve the margin of error, quadruple the sample size. Solve $ME = z^* \times \sqrt{p^*(1-p)/n}$ for n to plan a study.

AP Stats Tip: $n = (z^*/ME)^2 \times p^*(1-p^*)$. Use $p^* = 0.5$ for the most conservative (largest) n .

Example: For $ME = 0.03$ at 95% with $p^* = 0.5$: $n = (1.96/0.03)^2 \times 0.25 = 1067$.

109 Choosing the Sample Size

Explanation: Before collecting data, pick n so the margin of error is small enough. Use a guessed p^* (often 0.5) when p is unknown.

AP Stats Tip: $n = (z^*/ME)^2 p^*(1-p^*)$. $p^* = 0.5$ maximizes the needed n -- safe but larger.

Example: A 3% margin poll at 95% needs about 1,068 respondents using $p^* = 0.5$.

110 Significance Tests: The Basics

Explanation: A significance test assesses the evidence a sample provides against a null hypothesis H_0 in favor of an alternative H_a . The result is a p -value.

AP Stats Tip: *Test = claim + data $\rightarrow$ p -value $\rightarrow$ decision. Always state hypotheses in context first.*

Example: Testing whether a coin is fair: $H_0: p = 0.5$ vs $H_a: p \neq 0.5$.

111 Null & Alternative Hypotheses

Explanation: H_0 (null) is the claim being tested, usually 'no effect' or 'no difference' ($p = p_0$). H_a (alternative) is what we seek evidence for: $p < p_0$, $p > p_0$, or $p \neq p_0$ (one- or two-sided).

AP Stats Tip: *H_a is always what the researcher suspects; H_0 is the status quo. One-sided vs two-sided matters for the p -value.*

Example: Claim 'new drug beats old': $H_0: p = p_0$ vs $H_a: p > p_0$ (one-sided).

112 P-Values

Explanation: The p -value is the probability of getting a result at least as extreme as the observed one, GIVEN that H_0 is true. Small p -values = strong evidence against H_0 .

AP Stats Tip: *$p\text{-value} = P(\text{data or more extreme} \mid H_0 \text{ true})$. It is NOT the probability H_0 is true.*

Example: $p = 0.03$ means: if H_0 were true, we would see data this extreme only 3% of the time.

113 Statistical Significance

Explanation: A result is statistically significant if the p -value is below the significance level α (usually 0.05). Then we reject H_0 .

AP Stats Tip: *If $p < \alpha$, reject H_0 ; if $p \geq \alpha$, fail to reject H_0 . Significance is NOT practical importance.*

Example: $p = 0.03 < 0.05$: reject H_0 -- the result is statistically significant.

114 Type I & Type II Errors

Explanation: Type I error: rejecting H_0 when it is true (false positive). Type II error: failing to reject H_0 when it is false (false negative). $\alpha = P(\text{Type I})$; $\text{power} = 1 - P(\text{Type II})$.

AP Stats Tip: *Type I = false alarm; Type II = missed signal. α is the Type I rate.*

Example: Convicting an innocent person is a Type I error; freeing a guilty one is Type II.

115 The Power of a Test

Explanation: Power = probability the test rejects H_0 when a specific alternative is true. Power increases with larger n , larger effect size, and higher α .

AP Stats Tip: *$\text{Power} = 1 - P(\text{Type II})$. Bigger samples and larger effects boost power.*

Example: A study with $n = 200$ has more power to detect a 5-point shift than one with $n = 50$.

116 The One-Sample z-Test for a Proportion

Explanation: Test statistic $z = (\hat{p} - p_0) / \sqrt{p_0(1-p_0)/n}$. Compare z to the standard normal to get the p -value; check large counts against p_0 .

AP Stats Tip: *Use p_0 (not $\hat{p}$) in the SE for the test: $z = (\hat{p} - p_0) / \sqrt{p_0(1-p_0)/n}$.*

Example: $\hat{p} = 0.62$, $p_0 = 0.5$, $n = 100$: $z = 0.12/0.05 = 2.4 \rightarrow p = 0.016$ (two-sided).

117 Tests & Confidence Intervals (Two-Sided)

Explanation: A two-sided test at level α rejects H_0 if and only if the parameter value in H_0 falls outside the $(1 - \alpha)$ CI. Intervals and two-sided tests agree.

AP Stats Tip: *For two-sided tests: reject H_0 at α if p_0 is outside the $100(1-\alpha)\%$ CI.*

Example: A 95% CI of (0.51, 0.57) excludes 0.5, so a two-sided test of $p = 0.5$ at 0.05 rejects H_0 .

118 The Two-Sample z-Interval for $p_1 - p_2$

Explanation: Interval: $(\hat{p}_1 - \hat{p}_2) \pm z^* \times \sqrt{SE_1^2 + SE_2^2}$, where $SE_i = \sqrt{\hat{p}_i(1-\hat{p}_i)/n_i}$. Conditions: independent random samples, large counts in each.

AP Stats Tip: *SE of a difference = $\sqrt{SE_1^2 + SE_2^2}$. Variances add; SDs do not.*

Example: Comparing approval in two states: difference 0.08 with ME 0.05 gives CI (0.03, 0.13).

119 The Two-Sample z-Test for $p_1 - p_2$

Explanation: $H_0: p_1 = p_2$. Test statistic $z = (\hat{p}_1 - \hat{p}_2) / \sqrt{\hat{p}\text{-hatc}(1-\hat{p}\text{-hatc})(1/n_1 + 1/n_2)}$, using the pooled $\hat{p}\text{-hatc} = (x_1 + x_2)/(n_1 + n_2)$.

AP Stats Tip: *Pool the samples for the two-proportion z-test (SE under H_0); do NOT pool for the interval.*

Example: $x_1 = 40/100$, $x_2 = 30/100$: $\hat{p}\text{-hatc} = 70/200 = 0.35$, then $z = 0.10/\sqrt{0.35 \times 0.65 \times 0.02}$.

120 Choosing Inference for Proportions

Explanation: Match the procedure to the design: one sample $\rightarrow$ one-proportion z; two independent samples $\rightarrow$ two-proportion z; paired or matched $\rightarrow$ not proportions (use means).

AP Stats Tip: *One sample = one-proportion z. Two independent samples = two-proportion z. Everything else, slow down.*

Example: A before/after survey of the same 100 people is PAIRED -- do not use two-sample z.

Unit 7: Unit 7: Inference for Quantitative Data: Means

10-18% of the AP Stats exam. t-procedures for means.

121 The t-Distribution

Explanation: When sigma is unknown and we use s, the statistic $(\bar{x} - \mu)/(s/\sqrt{n})$ follows a t-distribution with $n - 1$ degrees of freedom. t has fatter tails than the normal.

AP Stats Tip: *Unknown sigma $\rightarrow$ t with $df = n - 1$. t $\rightarrow$ normal as df grows.*

Example: With $n = 15$, use t with 14 degrees of freedom for inference about the mean.

122 Degrees of Freedom

Explanation: $df = n - 1$ for one-sample t-procedures. $df = n - 1$ for paired differences. For two samples, use the conservative $\min(n_1 - 1, n_2 - 1)$ or software's formula.

AP Stats Tip: *$df = n - 1$ (one sample); two-sample conservative $df = \min(n_1 - 1, n_2 - 1)$.*

Example: A sample of 25 gives $df = 24$ -- find t^* in the $df = 24$ row of the t-table.

123 The One-Sample t-Interval for a Mean

Explanation: Interval: $\bar{x} \pm t^* \times s/\sqrt{n}$, with $df = n - 1$. Conditions: random sample, normal population or $n \geq 30$ (CLT), independence.

AP Stats Tip: *t-interval: $\bar{x} \pm t^*(df) \times SE$, $SE = s/\sqrt{n}$.*

Example: $\bar{x} = 52$, $s = 8$, $n = 25$, t^* ($df = 24$, 95%) = 2.064: $52 \pm 2.064 \times 1.6 = (48.7, 55.3)$.

124 Conditions for t-Procedures

Explanation: Random: the sample is randomly selected or the experiment is randomized. Normal: population is normal or $n \geq 30$ (larger n for heavy skew/outliers). Independent: 10% condition for sampling without replacement.

AP Stats Tip: *State random, normal ($n \geq 30$ or no strong skew/outliers), and independent before t-procedures.*

Example: $n = 40$ from a slightly skewed population: CLT justifies normality -- proceed.

125 The One-Sample t-Test for a Mean

Explanation: $H_0: \mu = \mu_0$. Test statistic $t = (\bar{x} - \mu_0)/(s/\sqrt{n})$ with $df = n - 1$. Find the p-value from the t-table or technology.

AP Stats Tip: *$t = (\bar{x} - \mu_0)/(s/\sqrt{n})$. Larger $|t|$ = more evidence against H_0 .*

Example: $\bar{x} = 105$, $\mu_0 = 100$, $s = 12$, $n = 36$: $t = 5/2 = 2.5$, $df = 35$, $p = 0.017$ (two-sided).

126 Paired Data & the Paired t-Test

Explanation: When observations come in pairs (before/after, twins, two measurements on the same subject), take the DIFFERENCE for each pair and run a ONE-sample t-test on the differences.

AP Stats Tip: *Paired data $\rightarrow$ analyze the differences with a one-sample t. Never use two-sample procedures.*

Example: Blood pressure before/after a drug on 20 patients: test the mean of the 20 differences.

127 The Two-Sample t-Interval for $\mu_1 - \mu_2$

Explanation: Interval: $(\bar{x}_1 - \bar{x}_2) \pm t^* \times \sqrt{s_1^2/n_1 + s_2^2/n_2}$, $df = \min(n_1 - 1, n_2 - 1)$ (conservative) or software df .

AP Stats Tip: *Two-sample interval: difference $\pm t^* \times \sqrt{SE_1^2 + SE_2^2}$.*

Example: $\bar{x}_1 - \bar{x}_2 = 12$, $SE = 4$, $df = 29$, $t^* = 2.045$: $12 \pm 8.18 = (3.8, 20.2)$.

128 The Two-Sample t-Test for $\mu_1 - \mu_2$

Explanation: $H_0: \mu_1 = \mu_2$. Test statistic $t = (\bar{x}_1 - \bar{x}_2)/\sqrt{s_1^2/n_1 + s_2^2/n_2}$, $df = \min(n_1 - 1, n_2 - 1)$. Compare to the t-table.

AP Stats Tip: *Two-sample t: do NOT pool unless told to. Use the conservative df on the exam.*

Example: Two teaching methods, $n_1 = n_2 = 20$: $t = (82 - 78)/\sqrt{9/20 + 10/20} = 4/0.975 = 4.1$.

129 Pooled vs Unpooled t-Procedures

- Explanation:** The unpooled (Welch) procedure is the default and does not assume equal SDs. The pooled procedure assumes $\sigma_1 = \sigma_2$ -- rarely justified; the AP exam uses the unpooled form.
- AP Stats Tip:** *Default to unpooled. Pooling requires the equal-SD assumption.*
- Example:** If sample SDs differ a lot ($s_1 = 3$ vs $s_2 = 15$), pooling would badly misstate the SE.
-

130 The Robustness of t-Procedures

- Explanation:** t-procedures are robust to moderate non-normality, especially for larger samples and two-sided tests, but NOT robust to outliers or strong skew in small samples.
- AP Stats Tip:** *t tolerates mild skew for $n \geq 15-30$; outliers still wreck small-sample t-tests.*
- Example:** A sample of 12 with one extreme outlier: do not trust the t-interval -- the mean is distorted.
-

131 The Standard Error of the Mean

- Explanation:** $SE(\bar{x}) = s/\sqrt{n}$ estimates $\sigma/\sqrt{n}$ when σ is unknown. It is the denominator of every one-sample t statistic.
- AP Stats Tip:** *$SE(\bar{x}) = s/\sqrt{n}$. Quadruple n to halve the SE.*
- Example:** $s = 10$, $n = 100$: $SE = 1$. At $n = 400$, $SE = 0.5$.
-

132 Interpreting t-Intervals for Means

- Explanation:** A 95% t-interval for a mean says we are 95% confident the true mean lies between the bounds -- in context with units.
- AP Stats Tip:** *Interpret the interval in context: 'We are 95% confident the mean [variable] is between...'*
- Example:** 'We are 95% confident the mean commute time is between 28 and 36 minutes.'
-

133 Means vs Proportions: Choosing the Procedure

- Explanation:** Proportions: categorical data, z-procedures, np conditions. Means: quantitative data, t-procedures, $n \geq 30$ /normal condition. Match the variable type first.
- AP Stats Tip:** *Categorical $\rightarrow$ proportions (z). Quantitative $\rightarrow$ means (t). Never mix them.*
- Example:** 'What fraction of voters approve?' = proportion. 'Average commute?' = mean.
-

134 The Effect of Sample Size on t-Intervals

- Explanation:** Larger n shrinks the SE and the t^* value, so intervals get narrower. The confidence level and data spread also affect width.
- AP Stats Tip:** *Interval width falls with $1/\sqrt{n}$; bigger n = more precise estimate.*
- Example:** Doubling n narrows the interval by about 30%; quadrupling n halves its width.
-

135 The Effect of Confidence Level

- Explanation:** Higher confidence = wider interval (larger t^*). The trade-off: 99% intervals are wider but more likely to capture μ than 90% intervals.
- AP Stats Tip:** *Higher confidence level $\rightarrow$ wider interval, same data.*
- Example:** At $n = 25$, the 99% t^* (2.797) exceeds the 95% t^* (2.064), widening the interval.
-

136 The Structure of an Inference Problem

- Explanation:** The AP exam wants four steps: STATE the parameter and hypotheses, PLAN (procedure + conditions), DO (compute the test statistic/interval), CONCLUDE (in context).
- AP Stats Tip:** *STATE - PLAN - DO - CONCLUDE. Never skip the conditions check or the context.*
- Example:** 'We reject H_0 ; there is convincing evidence the mean differs from 100.'
-

137 t* from the Table vs Technology

- Explanation:** The t-table gives t^* for whole df values; software interpolates. On the exam, use the nearest df row and round conservatively.
- AP Stats Tip:** Use the t-table's df row at or below your df when in doubt (conservative).
- Example:** df = 46: use the df = 40 row ($t^* = 2.021$) -- slightly wider than exact.
-

138 Inference for a Mean Difference vs Two Means

- Explanation:** Paired design (same subjects twice) -> one-sample t on differences, df = n - 1. Independent samples -> two-sample t, df = min(n1-1, n2-1).
- AP Stats Tip:** The DESIGN decides the procedure: paired differences vs independent samples.
- Example:** Testing a diet on the same 30 people before/after is paired; testing two separate diet groups is two-sample.
-

139 Errors in Context for Means

- Explanation:** Type I: concluding the mean changed when it did not. Type II: concluding no change when it did. Alpha and power behave the same as for proportions.
- AP Stats Tip:** Restate Type I/II errors in the context of the actual variable being tested.
- Example:** 'Type I error: we say the new curriculum raises scores when it really does not.'
-

140 Practical vs Statistical Significance

- Explanation:** A statistically significant result may be tiny in real-world terms, especially with large samples. Always judge the SIZE of the effect, not just the p-value.
- AP Stats Tip:** Significance = evidence of an effect; practical importance = size of the effect.
- Example:** n = 5,000 finds a 0.3-point IQ difference significant -- but it is practically meaningless.
-

Unit 8: Unit 8: Inference for Categorical Data: Chi-Square

2-5% of the AP Stats exam. Goodness-of-fit, homogeneity, and independence.

141 The Chi-Square Distribution

Explanation: Chi-square distributions are right-skewed, nonnegative, and indexed by degrees of freedom. The test statistic $\chi^2 = \sum((\text{observed} - \text{expected})^2 / \text{expected})$ follows this distribution under H_0 .

AP Stats Tip: $\chi^2 = \sum((O - E)^2 / E)$. Skewed right; df changes the shape.

Example: With df = 4, the critical value for $\alpha = 0.05$ is 9.49; larger χ^2 lies in the right tail.

142 Expected Counts

Explanation: Expected counts are what we would see if H_0 were true. For a two-way table: $E = (\text{row total} \times \text{column total}) / \text{table total}$. For goodness-of-fit: $E = n \times p$.

AP Stats Tip: $E = (\text{row total} \times \text{column total}) / \text{grand total}$. Every expected count must be ≥ 5 .

Example: In a 2x2 table, cell (1,1) with row 30, column 40, total 100: $E = 30 \times 40 / 100 = 12$.

143 The Chi-Square Statistic

Explanation: $\chi^2 = \sum \text{over all cells of } (\text{observed} - \text{expected})^2 / \text{expected}$. Large values = big deviations from H_0 's expectation $\rightarrow$ small p-value.

AP Stats Tip: Each cell contributes $(O - E)^2 / E$; bigger contributions mean more evidence.

Example: A cell with $O = 25$, $E = 10$ contributes $(15)^2 / 10 = 22.5$ to the total.

144 The Chi-Square Goodness-of-Fit Test

Explanation: Tests whether a categorical variable follows a claimed distribution. H_0 : the distribution is $p_1, p_2, \dots$ H_a : at least one proportion differs. df = number of categories - 1.

AP Stats Tip: GOF: one variable, one claimed distribution. $df = k - 1$.

Example: Testing whether M&M colors occur in the company's claimed 24/20/16/14/13/13 split.

145 Conditions for Goodness-of-Fit

Explanation: Random sample; all expected counts ≥ 5 ; independence (10% condition). State them before computing.

AP Stats Tip: Large counts for GOF: every $E \geq 5$.

Example: A die rolled 60 times with equal expected counts of 10 per face: $E = 10 \geq 5$ -- OK.

146 The Chi-Square Test for Homogeneity

Explanation: Tests whether the distribution of a categorical variable is the SAME across two or more populations or treatments. H_0 : the distributions are identical. $df = (r - 1)(c - 1)$.

AP Stats Tip: Homogeneity: compare distributions across groups; same formula, different question.

Example: Do Democrats, Republicans, and Independents differ in their opinion distribution? (3 x k table).

147 The Chi-Square Test for Independence

Explanation: Tests whether two categorical variables are associated in ONE population. H_0 : the variables are independent. $df = (r - 1)(c - 1)$.

AP Stats Tip: Independence: one sample, two variables. Homogeneity: multiple groups, one variable.

Example: From one random sample of students, test whether gender and major are associated.

148 Homogeneity vs Independence

Explanation: Both use the same χ^2 formula and df, but the design differs: homogeneity has separate random samples (one variable), independence has one random sample with two variables measured.

AP Stats Tip: Homogeneity = several groups, one variable. Independence = one group, two variables.

Example: 3 schools sampled separately $\rightarrow$ homogeneity. One city sampled, asking 2 questions $\rightarrow$ independence.

149 Degrees of Freedom for Chi-Square

Explanation: Two-way tables: $df = (\text{number of rows} - 1) \times (\text{number of columns} - 1)$. Goodness-of-fit: $df = \text{number of categories} - 1$.

AP Stats Tip: $df = (r - 1)(c - 1)$ for tables; $k - 1$ for GOF.

Example: A 3 x 4 table has $df = (3 - 1)(4 - 1) = 6$.

150 The Chi-Square Contribution

Explanation: Each cell's $(O - E)^2/E$ measures how far that cell is from expectation. Cells with large contributions drive a significant result -- examine them after the test.

AP Stats Tip: After a significant χ^2 , look at the largest $(O - E)^2/E$ cells for the story.

Example: If one ethnic group is hugely underrepresented, its cell contributes most of the χ^2 .

151 The Chi-Square P-Value

Explanation: The p-value = the area to the RIGHT of the computed χ^2 in the chi-square distribution with the right df. Small p -> reject H_0 .

AP Stats Tip: chi-square p-values are right-tail areas. Reject when χ^2 exceeds the critical value.

Example: $\chi^2 = 14.2$ with $df = 4$: $p = 0.0067 < 0.05$ -- reject H_0 .

152 Chi-Square and the z-Test for Two Proportions

Explanation: For a 2 x 2 table, the χ^2 test for independence is equivalent to a two-sided two-proportion z-test: $\chi^2 = z^2$ and the p-values match.

AP Stats Tip: $\chi^2(1 df) = z^2$. The two tests agree on 2 x 2 tables.

Example: $z = 2.0$ two-sided gives $p = 0.0456$; $\chi^2 = 4.0$ with $df = 1$ gives the same p.

153 Large Counts Condition for Tables

Explanation: For chi-square tests on tables, ALL expected counts must be at least 5 (some books allow 80% ≥ 5 with all ≥ 1 ; the AP exam expects $E \geq 5$ for all cells).

AP Stats Tip: Check every $E \geq 5$ before trusting the chi-square approximation.

Example: A cell with $E = 3.1$ fails the condition -- combine categories or collect more data.

154 Visualizing Categorical Association

Explanation: Segmented bar charts and mosaic plots display conditional distributions. Differences in segment proportions across groups hint at what χ^2 will find.

AP Stats Tip: Graph the conditional distributions FIRST; χ^2 tests what the graph shows.

Example: A mosaic plot with very different column heights suggests association.

155 What Chi-Square Does NOT Tell You

Explanation: A significant χ^2 says the variables are associated -- it does NOT measure the strength of association, identify causation, or say which cells differ.

AP Stats Tip: Significant χ^2 = association exists, not how strong or why.

Example: A huge sample can make a trivial association 'significant'.

156 Reading the Chi-Square Table

Explanation: The chi-square table gives critical values for right-tail probabilities at each df. Find df in the rows, alpha in the top row, and read the critical value.

AP Stats Tip: Critical values are in the body of the table; p-values usually come from technology.

Example: $df = 5$, $\alpha = 0.05$: critical value 11.07. $\chi^2 = 12$ rejects H_0 .

157 Follow-Up Analysis

Explanation: After a significant test for independence/homogeneity, compare observed vs expected counts cell by cell to describe the pattern of association.

AP Stats Tip: *Describe which cells are above or below expectation -- that is the real answer.*

Example: 'Women were overrepresented in nursing and underrepresented in engineering relative to expectation.'

158 Small Expected Counts

Explanation: Expected counts below 5 make the chi-square approximation unreliable. Remedies: combine categories (with care) or collect a larger sample.

AP Stats Tip: *$E < 5$ -> approximate is poor. Pool categories or gather more data.*

Example: A rare response category with $E = 2$ in three cells needs combining or a bigger n .

159 Choosing the Right Chi-Square Test

Explanation: GOF: one variable vs a claimed distribution. Homogeneity: same variable across several groups. Independence: two variables in one sample. Read the design, then choose.

AP Stats Tip: *Match the test to the data collection design -- the question wording reveals it.*

Example: 'Compare defect rates across 4 factories' = homogeneity. 'Is gender related to major?' = independence.

160 Chi-Square in Context

Explanation: A complete chi-square conclusion states the test, the p-value, the decision, and the association in the context of the study.

AP Stats Tip: *Conclude in context: 'There is convincing evidence that [variable 1] and [variable 2] are associated.'*

Example: 'The chi-square test ($\chi^2 = 21.4$, $df = 3$, $p < 0.001$) shows region and preference are related.'

Unit 9: Unit 9: Inference for Quantitative Data: Slopes

2-5% of the AP Stats exam. Confidence intervals and tests for regression slopes.

161 The Regression Model

Explanation: The population model: $y = \alpha + \beta x + \text{error}$. The observed data estimate α and β with the LSRL's intercept and slope. β is the true change in y per unit of x .

AP Stats Tip: β = true slope (parameter); b = sample slope (statistic). Inference targets β .

Example: If the true model is $y = 3 + 2x$, every +1 in x raises the mean y by 2.

162 The Least-Squares Line as an Estimate

Explanation: The LSRL $\hat{y} = a + bx$ estimates the population line. The sample slope b estimates β ; the sample intercept a estimates α .

AP Stats Tip: b estimates β ; a estimates α . Both have sampling variability.

Example: A study of 40 students finds $b = 5.2$ points per study hour -- an estimate of the true slope.

163 The Sampling Distribution of the Slope

Explanation: If the regression conditions hold, b is unbiased for β and approximately Normal with $SD(b) = \sigma / \sqrt{\sum((x - \bar{x})^2)}$.

AP Stats Tip: b is Normal around β ; its SE shrinks with more spread in x .

Example: Wider x -spread gives a more precise slope estimate -- that is why design matters.

164 The Standard Error of the Slope

Explanation: $SE(b) = s / \sqrt{\sum((x - \bar{x})^2)}$, where s is the SD of the residuals. It estimates the variability of the sample slope.

AP Stats Tip: $SE(b) = s / \sqrt{S_{xx}}$. Smaller residuals and wider x -spread shrink it.

Example: With $s = 4.2$ and $\sum((x - \bar{x})^2) = 225$, $SE(b) = 4.2/15 = 0.28$.

165 Conditions for Inference for Slopes

Explanation: Random: data are a random sample (or randomized experiment). Linear: the true relationship is linear (check residual plot). Independent: observations independent. Normal: residuals approximately Normal for small n .

AP Stats Tip: Check LINEARITY with a residual plot and NORMALITY with a histogram/QQ of residuals.

Example: A random sample of 30 homes with a random residual plot and Normal residual histogram passes all checks.

166 The Confidence Interval for the Slope

Explanation: Interval: $b \pm t^* \times SE(b)$, with $df = n - 2$. Interpret: we are 95% confident the true change in y per unit x lies between the bounds.

AP Stats Tip: CI for β : $b \pm t^*(n-2) \times SE(b)$. $df = n - 2$, not $n - 1$.

Example: $b = 5.2$, $SE = 0.8$, $n = 40$, $t^*(df\ 38) = 2.024$: $5.2 \pm 1.62 = (3.58, 6.82)$.

167 The Significance Test for the Slope

Explanation: $H_0: \beta = 0$ (no linear relationship). $H_a: \beta \neq 0$ (or one-sided). $t = (b - 0)/SE(b)$ with $df = n - 2$. Rejecting says x is a useful linear predictor of y .

AP Stats Tip: $t = b/SE(b)$, $df = n - 2$. Rejecting $\beta = 0$ = evidence of a linear association.

Example: $b = 5.2$, $SE = 0.8$: $t = 6.5$, $df = 38$, $p < 0.001$ -- strong evidence of a slope.

168 Interpreting the Slope Interval

Explanation: The interval for β gives the plausible range of the average change in y for a one-unit increase in x -- in context, with units.

AP Stats Tip: 'We are 95% confident each extra hour of study raises the predicted score by 3.6 to 6.8 points.'

Example: An interval of (3.58, 6.82) for score-per-hour says study time's effect is positive and meaningful.

169 Test vs Interval for the Slope

Explanation: A two-sided test at level α and a $(1 - \alpha)$ CI agree: reject H_0 ($\beta = 0$) exactly when 0 is outside the interval.

AP Stats Tip: *Two-sided test and CI agree for slopes, just like for means and proportions.*

Example: CI (3.58, 6.82) excludes 0, so the two-sided test at 0.05 rejects $\beta = 0$.

170 Residual Plots & Slope Inference

Explanation: The residual plot validates the linearity condition: random scatter around 0 supports the model; a curve or funnel shape invalidates inference.

AP Stats Tip: *Check residual plots for curvature and changing spread BEFORE trusting slope inference.*

Example: A U-shaped residual plot means a linear model is wrong -- the slope test is meaningless.

171 The t-Test for the Slope in Practice

Explanation: Technology (TI-84 LinRegTTest, software) outputs b , $SE(b)$, t , df , and p directly. Report all five in your conclusion.

AP Stats Tip: *Report: $b = \dots$, $SE = \dots$, $t = \dots$, $df = \dots$, $p = \dots$. Then conclude in context.*

Example: 'LinRegTTest gives $t = 6.5$, $df = 38$, $p < 0.001$: there is convincing evidence that GPA predicts salary.'

172 Slope Inference vs Correlation

Explanation: Testing $\beta = 0$ is equivalent to testing $\rho = 0$ (population correlation): both ask whether a linear relationship exists. r^2 describes strength, not significance.

AP Stats Tip: *t-test on slope = test of linear association; r^2 measures how much variation is explained.*

Example: With $r = 0.9$ and $n = 50$, the slope test is highly significant and $r^2 = 0.81$.

173 Extrapolation & Slope Inference

Explanation: Inference for the slope only applies within the observed range of x . Predicting y for x outside the data is extrapolation and invalid.

AP Stats Tip: *Never extend slope conclusions beyond the sampled x -range.*

Example: Height data from ages 5-18 cannot predict adult height at age 40.

174 The Role of Randomness

Explanation: Slope inference requires the data to be a random sample (or randomized experiment). Without randomness, no probability model justifies the p -value.

AP Stats Tip: *Random sampling or random assignment is the non-negotiable first condition.*

Example: A convenience sample of gym members cannot support inference about all adults.

175 Checking Linearity

Explanation: Look at the scatterplot and residual plot: no curved pattern = linearity OK. The residual plot is the primary diagnostic.

AP Stats Tip: *A patternless residual plot is the green light for linearity.*

Example: If residuals fan out (increasing spread), the equal-variance condition fails.

176 The Standard Deviation of the Residuals (s)

Explanation: $s = \sqrt{\text{sum}(\text{residual}^2)/(n - 2)}$ estimates σ , the SD of y about the line. It appears in $SE(b)$ and describes typical prediction error.

AP Stats Tip: *s = typical residual size; smaller s = tighter data around the line.*

Example: $s = 4.2$ means typical predictions miss by about 4.2 units.

177 Degrees of Freedom for Slopes

Explanation: Every regression inference uses $df = n - 2$: two parameters (slope and intercept) are estimated from the data.

AP Stats Tip: *$df = n - 2$ for all slope procedures.*

Example: A regression on 30 observations uses t with 28 degrees of freedom.

178 Prediction vs Estimation Intervals

Explanation: A confidence interval for the MEAN y at a given x is narrower than a prediction interval for an INDIVIDUAL y . Prediction intervals include the extra scatter about the line.

AP Stats Tip: *Prediction intervals (individual y) are wider than confidence intervals (mean y).*

Example: Estimating the average weight of 5'8" men is more precise than predicting one specific man's weight.

179 Common Errors in Slope Inference

Explanation: Common mistakes: using $df = n - 1$, forgetting the conditions, interpreting the interval as a prediction range, and confusing b with β .

AP Stats Tip: *$df = n - 2$; conditions first; interpret the interval as the plausible range of the TRUE slope.*

Example: A student who uses $df = 39$ instead of 38 for $n = 40$ misreads the t -table row.

180 A Complete Slope Example

Explanation: State: β = true change in score per study hour. Plan: t -interval, conditions met. Do: $b = 5.2$, $SE = 0.8$, $t^* = 2.024$, interval (3.58, 6.82). Conclude: strong evidence study time predicts score.

AP Stats Tip: *Full answers earn full credit: state the parameter, verify conditions, show the computation, conclude in context.*

Example: 'We are 95% confident the true slope is between 3.58 and 6.82 points per hour; since 0 is excluded, the association is significant.'

What's Next?

You have now reviewed every unit of the AP Statistics course, including every inference procedure. Go back to the Table of Contents, find the units you marked, and reread just those tip lines. Two weeks out, do a rapid pass over the entire guide; the day before, review only the tip lines. On the exam, structure every inference FRQ as State-Plan-Do-Conclude: name the parameter and hypotheses, identify the procedure, verify the conditions (random, normal/large counts, independent), show the test statistic or interval with its formula, and conclude in context -- graders award full credit for correct structure and conditions even when arithmetic wobbles.

Pair it with the AP Calculus AB Study Guide (\$19) -- both quantitative exams covered

Also available: AP Biology, AP Chemistry, AP Physics 1, AP Psychology, and more study guides

templateforge.com