

Statistics

A TI-84-focused statistics companion with concise notes, worked examples, and quick references for class, study, and review.

Module 1: Introduction, Data Collection, Sampling, and Experimental Design

Students learn the language of statistics, study design, sampling methods, bias, and experiment design.

Section 1.1 Introduction to the Practice of Statistics

Use the basic vocabulary of statistics and classify variables correctly.

QUICK REFERENCE

- Statistics is the science of collecting, organizing, summarizing, and analyzing information to answer questions.
- Data are the information collected from individuals. A variable is the characteristic being recorded, and the data are the observed values of that variable.
- The process of statistics usually starts with a research objective, then data are collected, summarized, and used for inference when appropriate.
- Statistics has two main branches.
 - Descriptive statistics organizes and summarizes data using tables, graphs, or numerical summaries.
 - Inferential statistics uses sample data to draw conclusions about a population and reports uncertainty.
- Three terms identify who or what is being studied.
 - A population is the entire group being studied.
 - A sample is a subset of that population.
 - An individual is one person or object in the group being studied.
- A numerical summary is named according to whether it describes a population or a sample.
 - A parameter is a numerical summary of a population.
 - A statistic is a numerical summary of a sample.
- Variables can be qualitative or quantitative.
 - Qualitative variables describe categories.
 - Quantitative variables are numerical values where arithmetic has meaning, so number labels such as ZIP codes or coded weekdays are not treated as quantitative.
- Quantitative variables can be discrete or continuous.
 - Discrete quantitative variables come from counting.
 - Continuous quantitative variables come from measurement and can take values between two measured values.

- The four levels of measurement describe how the values of a variable can be compared.
 - Nominal means names only.
 - Ordinal means categories with an order.
 - Interval means differences can be compared, but zero is not a true none.
 - Ratio means differences and ratios make sense because zero is a true none.

Example 1. Fill in each blank using statistics, descriptive statistics, inferential statistics, statistic, parameter, sample, population, or individual. Use each term once.

1. A numerical summary of a sample is a _____.
2. Organizing and summarizing data with tables, graphs, or numerical summaries is _____.
3. The science of collecting, organizing, summarizing, and analyzing information to answer questions is _____.
4. A subset of the group being studied is a _____.
5. Using sample data to draw conclusions about a population is _____.
6. One person or object in the group being studied is an _____.
7. A numerical summary of a population is a _____.
8. The entire group being studied is the _____.

Answer: 1. Statistic. 2. Descriptive statistics. 3. Statistics. 4. Sample. 5. Inferential statistics. 6. Individual. 7. Parameter. 8. Population.

Example 2. A researcher surveys 420 adults from a city to estimate the percent of all adults in the city who support a new park. Identify the population, sample, individual, variable, and whether 61% from the survey is a statistic or parameter.

Answer: Population: all adults in the city. Sample: the 420 adults surveyed. Individual: one adult in the city. Variable: whether the adult supports the new park. 61% is a statistic because it summarizes the sample.

Example 3. The table shows information collected from five campus parking spaces. Identify the individuals, variables, and one row of data.

Space	Payment method	Minutes parked	Side of campus	Space number
A	Card	42	North	18
B	Cash	25	South	11
C	App	60	North	22
D	Card	15	West	7
E	App	38	South	14

Answer: Individuals: parking spaces A through E. Variables: payment method, minutes parked, side of campus, and space number. One row of data for Space A is card, 42, north, and 18.

Example 4. Classify each variable as qualitative, quantitative discrete, or quantitative continuous: eye color, number of siblings, commute time, phone brand, ZIP code, day of the week coded as 1 through 7, number of text messages sent yesterday, and height.

Answer: Eye color: qualitative. Number of siblings: quantitative discrete. Commute time: quantitative continuous. Phone brand: qualitative. ZIP code: qualitative, because the number is only a label. Day coded 1 through 7: qualitative, because the numbers are category labels rather than measured amounts. Number of text messages: quantitative discrete. Height: quantitative continuous.

Example 5. Determine the level of measurement: hair color, movie rating from one to five stars, year of birth, and weight.

Answer: Hair color: nominal, because it names categories with no ranking. Movie rating: ordinal, because the stars have an order but the gaps between ratings are not measured amounts. Year of birth: interval, because subtracting years gives a meaningful difference, but year 0 is not a true zero amount. Weight: ratio, because 0 means no weight and ratios such as twice as heavy make sense.

Example 6. A summary gives the ages of students in one class. A separate statement says the average age from that class estimates the average age of all students at the college. Which part is descriptive and which part is inferential?

Answer: Summarizing the ages from one class is descriptive. Using that class average to estimate the average for all students at the college is inferential.

Example 7. A news organization surveys 1007 adults and finds that 38% believe a government agency wastes at least half of its budget. The organization reports a margin of error of 6% at a 90% confidence level. Identify the research objective, the descriptive statistic, and a reasonable inference.

Answer: Research objective: estimate the percent of adults who believe the agency wastes at least half of its budget. Descriptive statistic: 38% of the 1007 adults surveyed gave that response. Reasonable inference: the organization is 90% confident that between 32% and 44% of all adults believe the agency wastes at least half of its budget.

Section 1.2 Observational Studies versus Designed Experiments

Distinguish observational studies from experiments and avoid unsupported causal claims.

QUICK REFERENCE

- Two common ways to collect data are observational studies and designed experiments.
 - An observational study measures characteristics without trying to change the subjects.
 - A designed experiment assigns treatments and observes the response, while controlling other conditions when possible.
- Two variables describe a possible relationship in a study.
 - The explanatory variable may help explain changes in the response variable.
 - The response variable is the outcome being measured.
- Experiments use treatments and experimental units.
 - A treatment is the condition applied in an experiment.
 - An experimental unit is the person or object receiving the treatment.

- Common types of observational studies differ in when and how information is collected.
 - Cross-sectional studies collect information at one point in time or over a short time period.
 - Case-control studies are retrospective: researchers look back or use existing records to compare individuals with a condition to similar individuals without it.
 - Cohort studies are prospective: researchers follow a group over time and record what happens.
- Other variables can make a relationship difficult to interpret.
 - Confounding happens when the effects of two or more explanatory variables are mixed together.
 - A lurking variable was not included in the study but may affect the response and may be related to an explanatory variable in the study.
 - A confounding variable was included in the study, but its effect cannot be separated from the effect of another explanatory variable.
- No type of observational study is always best. The research question, available data, time, and cost help determine which type is most appropriate.
- Observational studies can show association, but they do not prove cause and effect by themselves. A well-designed experiment gives stronger evidence for cause and effect.

Example 1. Scenario A: A researcher records coffee intake and sleep hours for 200 adults without assigning coffee amounts. Scenario B: Students are randomly assigned to either use a new study app or not use it, then their exam scores are compared. Which scenario is observational, and which is an experiment?

Answer: Scenario A is observational because the researcher only records what already happens. Scenario B is an experiment because the researcher assigns the treatment.

Example 2. Researchers study whether study time is related to exam score. Identify the explanatory variable and response variable.

Answer: Explanatory variable: study time. Response variable: exam score.

Example 3. Classify each observational study as cross-sectional, case-control, or cohort: (a) A survey asks adults about their current exercise habits today. (b) Researchers compare patients with kidney failure to similar patients without kidney failure and check past diabetes records. (c) Researchers follow 1806 adults for 7 years to see who develops high blood pressure.

Answer: (a) Cross-sectional. (b) Case-control. (c) Cohort.

Example 4. Which study is retrospective, and which is prospective? Study A uses past medical records to compare people with and without a disease. Study B follows a group of people for 10 years and records who develops the disease. Is one study type always better than the other?

Answer: Study A is retrospective because it looks back at existing records. Study B is prospective because it follows people forward over time. Neither type is always better; the better choice depends on the research question, available data, time, and cost.

Example 5. Researchers study the relationship between exercise and stress. They include sleep in the study, but people who sleep more also tend to exercise more, so the separate effects of sleep and exercise cannot be determined. The researchers do not record workload, which may affect both exercise habits and stress. Identify the confounding variable, the lurking variable, and the safest conclusion.

Answer: Confounding variable: sleep, because it was included but its effect cannot be separated from the effect of exercise. Lurking variable: workload, because it was not included but may affect the response and may be related to exercise. Safest conclusion: exercise and stress are associated, but this observational study does not prove that exercise caused the difference in stress.

Example 6. Researchers find that children living near high-tension wires have a higher rate of leukemia than children who do not live near high-tension wires. Can the researchers conclude that the wires cause leukemia?

Answer: No. This is an observational study, so it can show an association but cannot prove cause and effect by itself.

Example 7. Office workers are randomly assigned to take either meditation breaks or walking breaks every 2 hours for one month. Their stress levels are compared. Identify the treatment, experimental units, explanatory variable, and response variable.

Answer: Treatments: meditation breaks and walking breaks. Experimental units: the office workers. Explanatory variable: type of break. Response variable: stress level.

Section 1.3 Simple Random Sampling

Recognize and create simple random samples.

QUICK REFERENCE

- Data may be collected from a census or a sample.
 - A census collects data from every individual in the population.
 - A sample collects data from only part of the population.
- Random sampling uses chance to select individuals. A simple random sample of size n goes further: every possible group of n individuals has the same chance of being selected.
- A sampling frame is the list of all individuals available for selection. Give each individual one unique number; for a random number table, use the same number of digits, such as 01 through 40.
- Sampling may be done with or without replacement.
 - Sampling without replacement means a selected individual cannot be chosen again.
 - Sampling with replacement means the individual is returned and may be chosen again.
- A random number table can be used to select a simple random sample.
 - Choose a starting position and direction, read digits in groups that match the assigned numbers, and keep valid numbers until the sample is complete.
 - Skip 00, numbers outside the population range, and repeated numbers when sampling without replacement.
- Random selection helps reduce selection bias.

- TI-84: press MATH → PRB to access random-integer commands.
 - Use randInt(lower bound,upper bound,number of integers). Both bounds are included, and repeated values are possible. For sampling without replacement, ignore repeated values.
 - Use randIntNoRep(lower bound,upper bound) when available to generate different values without replacement.

Example 1. A school has 1200 students. Plan A surveys all 1200 students. Plan B numbers the students from 1 to 1200 and uses random integers to select 80 different students. Identify the census, the sample, and the sampling method used in Plan B.

Answer: Plan A is a census because every student is surveyed. The 80 students in Plan B are a sample. Plan B uses simple random sampling, assuming every possible group of 80 students has the same chance of being selected.

Example 2. Which methods could produce a simple random sample of 3 books from a list of 9? (a) Number the books and use a random number generator. (b) Put one slip for each book in a container and draw three without looking. (c) Choose the first three books alphabetically. (d) Ask someone to choose the three best books.

Answer: Methods (a) and (b) could produce a simple random sample because chance determines the selections. Methods (c) and (d) are not random.

Example 3. Ava, Ben, Cora, and Diego are candidates for a two-person committee. List all possible samples of size 2 when sampling without replacement. What is the probability of selecting Ava and Cora? If Ava's name is returned before the second draw, is that with or without replacement, and could Ava be selected twice?

Answer: The six possible samples are Ava-Ben, Ava-Cora, Ava-Diego, Ben-Cora, Ben-Diego, and Cora-Diego. Each has probability $1/6$, so the probability of Ava-Cora is $1/6$. Returning Ava's name is sampling with replacement, so Ava could be selected again.

Example 4. A club numbers its 40 members from 01 through 40. Starting with the upper-left entry and reading across each row, use the random number pairs to select a simple random sample of 5 members without replacement. Explain which entries are skipped.

Random number pairs		
07	44	12
12	00	36
41	25	03

Answer: Select members 07, 12, 36, 25, and 03. Skip 44 and 41 because they are outside the range 01 through 40, skip the second 12 because sampling is without replacement, and skip 00 because no member has that number.

Example 5. A company numbers its 50 employees from 1 to 50 and wants a simple random sample of 6 employees without replacement. The TI-84 command randInt(1,50,8) returns {31, 12, 4, 12, 48, 19, 27, 6}. Which employees are selected? Which command can avoid repeated values when it is available?

Answer: Select employees 31, 12, 4, 48, 19, and 27. Ignore the second 12 because the sample is without replacement, and stop after six different employees have been selected. To avoid repeats, use randIntNoRep(1,50), which randomizes the integers from 1 through 50, and take the first six values from the list.

Example 6. A college wants the opinions of all enrolled students. Why is randomly selecting students from the complete enrollment list preferred over asking for volunteers, and what is the sampling frame?

Answer: The complete enrollment list is the sampling frame. Random selection reduces selection bias, while volunteers may differ from the full student population in important ways.

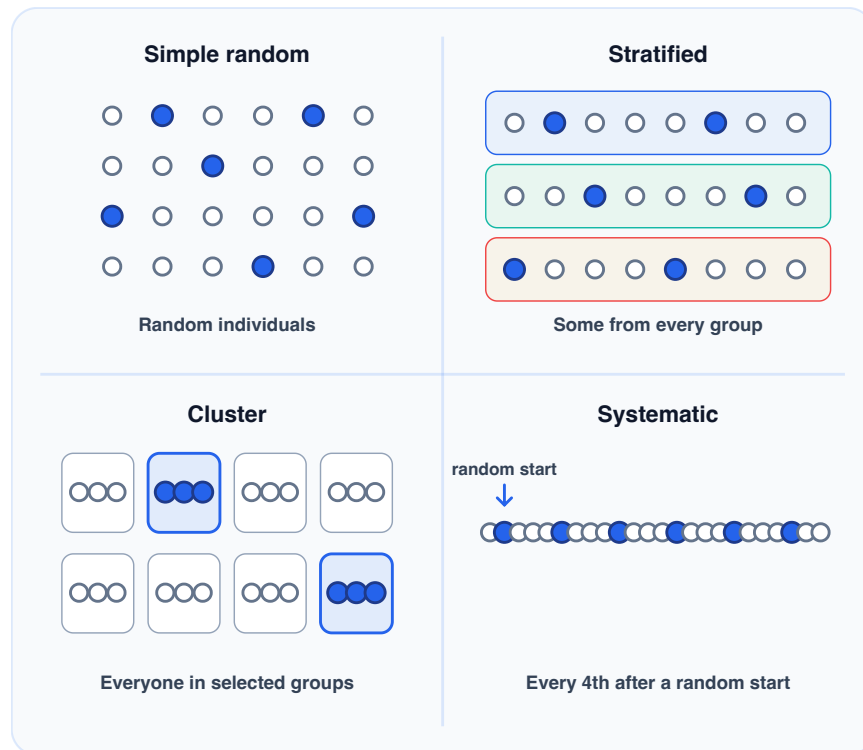
Section 1.4 Other Effective Sampling Methods

Classify and compare stratified, cluster, systematic, and convenience sampling.

QUICK REFERENCE

- Stratified sampling separates the population into nonoverlapping groups, called strata, based on a shared characteristic. A random sample is selected from every stratum.
 - The number selected from each stratum may be the same or may be proportional to the stratum's size.
 - In a proportional sample, a group containing 30% of the population supplies 30% of the sample.
 - Stratified sampling guarantees representation from every stratum.
- Cluster sampling separates the population into groups, randomly selects some entire groups, then includes everyone in the selected groups.
 - Cluster sampling can reduce travel, time, or cost when the groups are spread out.
- To distinguish the methods, remember: stratified sampling selects some individuals from every group; cluster sampling selects everyone from some groups.
- Systematic sampling selects every k th individual after a random starting position.
 - To obtain a sample of size n from a population of size N , use $k = \lfloor N/n \rfloor$, randomly choose p from 1 through k , and select positions $p, p + k, p + 2k$, and so on through $p + (n - 1)k$.
 - A systematic sample uses a random start, but it is not a simple random sample because only samples following the every- k th pattern can be selected.
- Convenience sampling uses individuals who are easy to reach and is not a random sampling method.
 - A voluntary response sample is a common convenience sample in which people choose whether to participate.

Sampling Methods at a Glance



Example 1. Identify the sampling method in each situation. (a) A college numbers all students and randomly selects 160 different numbers. (b) The college randomly selects 40 freshmen, 40 sophomores, 40 juniors, and 40 seniors. (c) It randomly selects 8 classrooms and surveys every student in those rooms. (d) A quality checker inspects every 25th item after a random start. (e) A researcher surveys students who are eating in the nearest dining hall.

Answer: (a) Simple random sampling; chance selects individuals from the full student list. (b) Stratified sampling; some students are selected from every class year. (c) Cluster sampling; everyone in the selected classrooms is surveyed. (d) Systematic sampling; every 25th item is inspected. (e) Convenience sampling; the nearest available students are surveyed.

Example 2. A company has 120 day-shift employees, 80 evening-shift employees, and 40 night-shift employees. It wants a proportional stratified sample of 24 employees. Complete the table, then determine how many employees should be randomly selected from each shift.

Shift	Population size	Percent of population	Number sampled
Day	120		
Evening	80		
Night	40		

Answer: The sample is 10% of the 240 employees. Day shift is 50% of the population, so select 12 employees. Evening shift is 33.3%, so select 8 employees. Night shift is 16.7%, so select 4 employees. The percentages total 100%, and the sample counts total 24.

Example 3. A university wants information about students living in its 12 residence halls. Plan A randomly selects 10 students from every residence hall. Plan B randomly selects 3 residence halls and surveys every student in those halls. Identify each method and explain the key difference.

Answer: Plan A is stratified sampling because it selects some students from every residence hall. Plan B is cluster sampling because it selects some residence halls and surveys everyone in them.

Example 4. A population list contains $N = 96$ people, and a systematic sample of $n = 8$ people is needed. Find k . If the random starting position is $p = 5$, list the people selected.

Answer: $k = \lfloor 96/8 \rfloor = 12$. Starting at 5 and adding 12 each time gives 5, 17, 29, 41, 53, 65, 77, and 89. On a TI-84, `randInt(1,12)` can be used to choose the random starting position.

Example 5. Plan A randomly selects 8 different identification numbers from 1 through 96. Plan B randomly chooses a starting position from 1 through 12, then selects every 12th person. Which plan is a simple random sample, and which is a systematic sample?

Answer: Plan A is a simple random sample because every possible group of 8 people has the same chance of being selected. Plan B is a systematic sample because only groups that follow the every-12th pattern can be selected.

Example 6. A website posts a poll and lets anyone answer. What sampling concern is present?

Answer: This is convenience sampling, specifically a voluntary response sample. People choose whether to participate, so those with strong opinions may be overrepresented.

Section 1.5 Bias in Sampling

Identify common sources of bias and improve sampling plans.

QUICK REFERENCE

- Bias occurs when a method systematically produces results that do not accurately represent the population.
- The three main sources of bias are sampling bias, nonresponse bias, and response bias.
- Sampling bias occurs when the method used to select individuals tends to favor one part of the population over another.
 - Undercoverage occurs when a group in the population is left out or represented too little in the sample.
 - Voluntary response bias occurs when people choose whether to participate. People with strong opinions may be more likely to respond.
- Nonresponse bias occurs when selected individuals who do not respond differ in an important way from those who do respond. A low response rate creates concern, but nonresponse causes bias when the two groups differ.

- Response bias occurs when people respond, but their answers do not reflect their true beliefs or behavior.
 - Wording bias is caused by leading, loaded, or unbalanced wording. Survey questions should be written in a neutral, balanced way.
 - Interviewer influence or pressure can affect how a person answers.
 - Sensitive questions can lead people to misrepresent their beliefs or behavior.
 - The order of questions or answer choices can influence responses.
- To distinguish the main sources, ask: Who had a chance to be selected? Who was selected but did not respond? Did the responses reflect the truth? These questions point to sampling bias, nonresponse bias, and response bias, respectively.
- Callbacks, alternate contact methods, and rewards or incentives can reduce nonresponse. Simply increasing the number contacted does not correct a biased response pattern.
- Neutral wording, impartial administration, anonymity when appropriate, and rotating question or answer order can reduce response bias.
- Sampling error is the natural sample-to-sample difference that occurs because only part of the population is observed; it is not the same as sampling bias. Undercoverage, nonresponse, response bias, and data-entry mistakes are nonsampling errors and can occur even in a census.

Example 1. Fill in each blank using bias, sampling bias, nonresponse bias, or response bias. Use each term once.

1. _____ occurs when the method used to select individuals tends to favor one part of the population over another.
2. _____ occurs when survey answers do not reflect the respondents' true beliefs or behavior.
3. A sample has _____ when its results systematically fail to represent the population accurately.
4. _____ occurs when selected individuals who do not respond differ from those who respond.

Answer: (1) Sampling bias. (2) Response bias. (3) Bias. (4) Nonresponse bias.

Example 2. Identify the most specific source of bias in each situation.

- a. A store manager surveys the first 60 customers who enter on Saturday morning.
- b. A phone survey calls only landlines to estimate the opinions of all adults.
- c. A news website posts a poll that any visitor may choose to answer.
- d. A questionnaire is mailed to 1117 randomly selected households, but only 14 respond.
- e. A supervisor asks employees face-to-face whether they have violated company rules.
- f. A survey asks, "Do you support the wasteful new tax increase?"

Answer: (a) Sampling bias, because the first customers on one morning may not represent all customers. (b) Undercoverage, a source of sampling bias, because adults without landlines are missed. (c) Voluntary response bias, a source of sampling bias, because people choose whether to respond. (d) Nonresponse bias is a concern because the selected nonrespondents may differ from the respondents. (e) Response bias is likely because the supervisor's presence may influence the answers. (f) Wording bias, a source of response bias, because the word wasteful pushes respondents toward an answer.

Example 3. A food-court manager surveys the first 100 customers who arrive on weekday mornings. Identify the bias, explain how the method could affect the results, and improve the plan.

Answer: This is sampling bias because weekday-morning customers may have different preferences from customers who visit later or on weekends. Randomly select customers at different times throughout weekdays and weekends.

Example 4. A survey asks, "Do you support the wasteful new tax increase?" Identify the bias and rewrite the question more neutrally.

Answer: This is response bias caused by loaded wording. A neutral version is, "Do you support or oppose the proposed tax increase?"

Example 5. Distinguish nonresponse bias from response bias. Scenario A: A mail survey is sent to 2000 randomly selected people, but only 180 respond. Scenario B: College students answer a question about study hours in front of their instructor, and some exaggerate. Identify the concern and give one remedy for each scenario.

Answer: Scenario A raises concern about nonresponse bias because the people who did not respond may differ from those who did. Use callbacks, another contact method, or an incentive. Scenario B has response bias because some answers are inaccurate. Use an anonymous survey administered by an impartial party.

Example 6. Scenario A: A hospital administrator asks doctors to report illicit drug use. Scenario B: A poll asks about the nation's most important problem immediately before asking respondents to rate the president's job performance. Explain the likely response bias and give a remedy for each.

Answer: In Scenario A, doctors may not answer truthfully because the interviewer has authority over them. Use an anonymous survey administered by an impartial party. In Scenario B, the first question may influence the answer to the second. Alternate or randomly change the question order across respondents.

Example 7. Classify each as sampling error or nonsampling error. (a) Two properly selected random samples give slightly different estimates because they contain different individuals. (b) A city census misses residents without stable housing. (c) A recorded age of 39 is entered as 93.

Answer: (a) Sampling error, because natural sample-to-sample differences occur when only part of the population is observed. (b) Nonsampling error caused by undercoverage; this can occur even in a census. (c) Nonsampling error caused by a data-entry mistake.

Section 1.6 The Design of Experiments

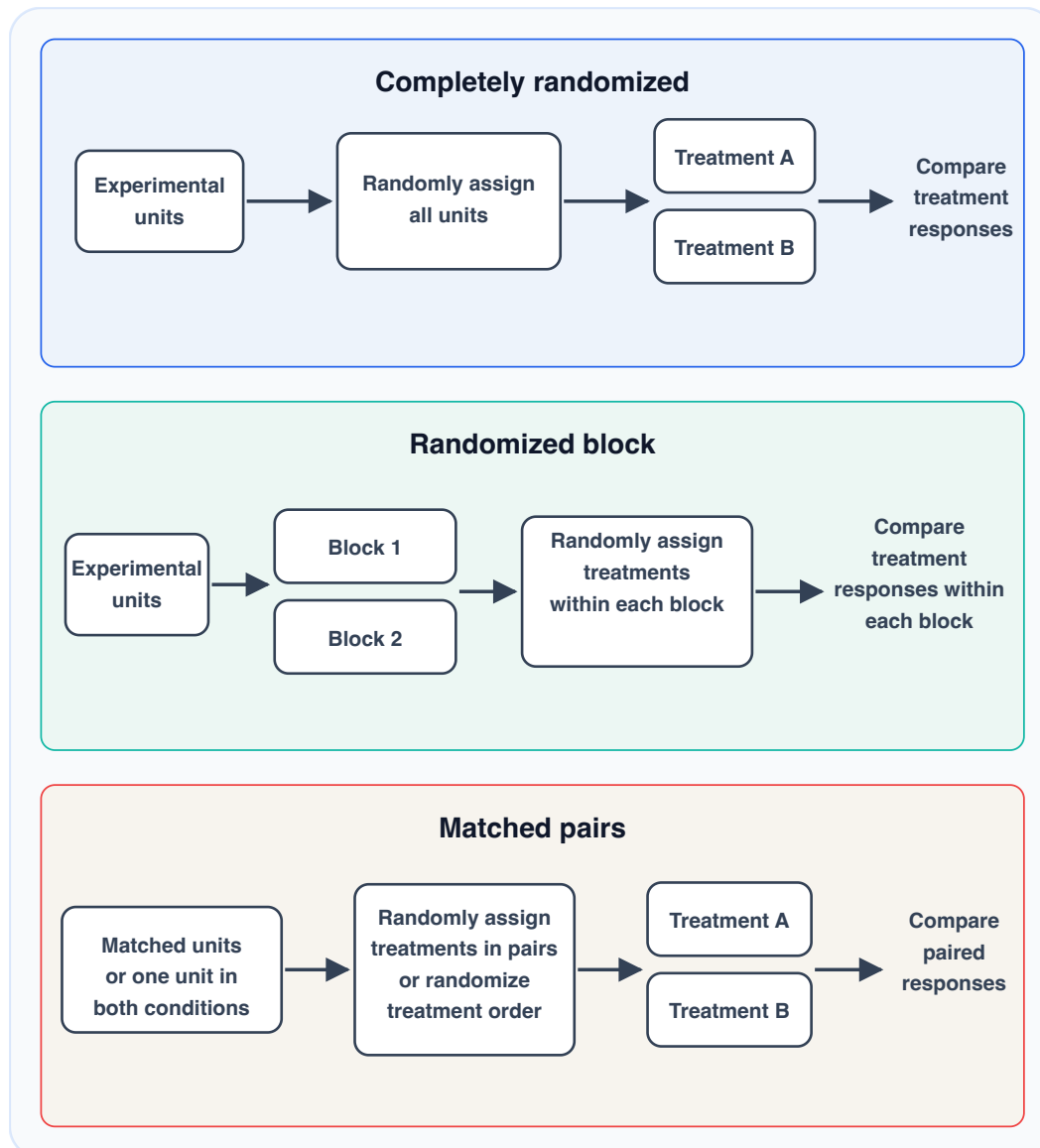
Identify experimental design features and explain their purpose.

QUICK REFERENCE

- An experiment deliberately applies one or more treatments to experimental units and measures a response.
 - An experimental unit is the person or object receiving a treatment. A person in an experiment may also be called a subject.
 - A factor is an explanatory variable whose effect is being studied. Its levels are the values chosen for the experiment.
 - A treatment is a specific condition applied to an experimental unit. When an experiment has several factors, a treatment is a combination of factor levels.
 - The response variable is the outcome measured after the treatment is applied.
- Good experimental designs use control, random assignment, and replication to make treatment comparisons more reliable.
 - Control means keeping other conditions as similar as possible so the treatments are the main planned difference between groups.
 - Some variables can be controlled by holding them constant, such as giving every plant the same amount of water. Other variables, such as natural differences among people or plants, cannot be held constant; random assignment helps distribute their effects across treatment groups.
 - Random assignment uses chance to place experimental units into treatment groups. It helps balance other variables across the groups.
 - Replication means applying each treatment to multiple experimental units. Repeatedly measuring the same unit does not provide the same replication as using additional units.
- Confounding occurs when the effect of a treatment cannot be separated from the effect of another variable. Control and random assignment help reduce confounding.
- A control group gives a baseline for comparison, and a placebo can help separate a treatment effect from the effect of expecting a treatment.
 - A control group may receive a placebo, an existing treatment, or no treatment, depending on the experiment.
 - A placebo looks like the treatment being tested but does not contain its active ingredient.
 - The placebo effect is a response caused by a subject's expectation of receiving a treatment rather than by the treatment itself.
- Blinding reduces the chance that knowledge of the assigned treatment will influence a subject's behavior or an evaluator's judgment.
 - In a single-blind experiment, the subjects usually do not know which treatment they receive.
 - In a double-blind experiment, neither the subjects nor the researchers who interact with them or evaluate the response know the treatment assignments.

- Three common experimental designs differ in how the experimental units are organized before treatments are assigned.
 - Completely randomized design: randomly assign all experimental units directly to the treatments.
 - Randomized block design: first separate similar experimental units into blocks, then randomly assign treatments within every block. Blocking reduces variation caused by the characteristic used to form the blocks.
 - Matched-pairs design: compare two treatments using closely matched units or the same unit under both conditions, then compare responses within each pair.
- Random sampling and random assignment serve different purposes.
 - Random sampling helps produce a sample that can represent the population.
 - Random assignment helps create comparable treatment groups and supports cause-and-effect conclusions from a well-designed experiment.

Experimental Designs at a Glance



Example 1. A study randomly assigns 90 patients to a new medication, an existing medication, or a placebo. After 8 weeks, researchers record each patient's change in systolic blood pressure. Identify the experimental units, factor, factor levels or treatments, response variable, and control group.

Answer: Experimental units: the 90 patients. Factor: medication type. Levels or treatments: new medication, existing medication, and placebo. Response variable: change in systolic blood pressure after 8 weeks. Control group: the patients receiving the placebo.

Example 2. A fertilizer experiment uses 60 similar plants. The plants are randomly assigned to three fertilizer amounts, all plants receive the same amount of water and light, and each fertilizer amount is applied to 20 plants. Explain how the experiment uses control, random assignment, and replication.

Answer: Control: water and light are kept the same. Random assignment: chance places plants into fertilizer groups, helping balance other plant differences across the groups. Replication: each fertilizer amount is applied to 20 plants, so the result does not depend on only one plant.

Example 3. In a drug study, the control group receives a placebo. Neither the patients nor the doctors evaluating the outcomes know who received the real drug. Identify the blinding and explain why the placebo and blinding are useful.

Answer: This is double-blind. The placebo gives the control group a treatment experience similar to the real drug, which helps account for the placebo effect. Blinding helps prevent patient expectations and evaluator judgment from influencing the results.

Example 4. Identify the experimental design in each situation. (a) All 60 volunteers are randomly assigned directly to one of three exercise programs. (b) Volunteers are first grouped by initial fitness level, then randomly assigned to exercise programs within each group. (c) Researchers pair people with similar ages and fitness levels, then randomly assign one person in each pair to Program A and the other to Program B.

Answer: (a) Completely randomized design. (b) Randomized block design, with initial fitness level used to form the blocks. (c) Matched-pairs design.

Example 5. Twelve volunteers are numbered 01 through 12. Starting at the upper left and reading across, use the random number pairs to assign six volunteers to Treatment A. Skip invalid or repeated numbers. Assign the remaining volunteers to Treatment B. What design is this, and is this random sampling or random assignment?

Random number pairs				
08	03	12	03	00
11	14	05	09	02

Answer: Treatment A receives volunteers 08, 03, 12, 11, 05, and 09. Skip the repeated 03, 00, and 14. Treatment B receives 01, 02, 04, 06, 07, and 10. This is a completely randomized design using random assignment, not random sampling, because the volunteers were already selected before treatments were assigned. On a TI-84, use `randIntNoRep(1,12)` and assign the first six values in the randomized list to Treatment A; assign the remaining six volunteers to Treatment B.

Example 6. Researchers expect a medication's effect to differ among adults ages 18-39, 40-64, and 65 or older. They form these three age groups, then randomly assign medication or placebo within every age group. Identify the design, the blocks, where random assignment occurs, and why this design may be better than assigning everyone directly to treatments.

Answer: This is a randomized block design. The three age groups are the blocks, and medication or placebo is randomly assigned within each age group. Blocking allows the treatment comparison to be made among people of similar ages, reducing variation caused by age.

Example 7. Both situations use matched pairs. Explain how each pair is formed and what is compared. (a) For each patient, one randomly selected area of the scalp receives a platelet-rich plasma injection and another area receives saline. Hair-density changes are measured after six months. (b) Bedrooms are paired by size and wall condition. Within each pair, one room is randomly assigned Paint A and the other Paint B, then coverage quality is measured.

Answer: (a) Two areas on the same patient form a pair. Compare the change in hair density after six months for the area receiving platelet-rich plasma with the area receiving saline. (b) Two bedrooms with similar size and wall condition form a pair. Compare the coverage quality of Paint A with Paint B within each pair.

Example 8. A plant-growth experiment compares two fertilizers. Researchers can give every plant the same amount of water, light, and soil, but they cannot make the plants genetically identical. Classify the variables as controllable or uncontrollable, and explain the role of random assignment.

Answer: Water, light, and soil are controllable because researchers can hold them constant. Genetic differences are uncontrollable because they cannot be made identical. Randomly assigning plants to the fertilizers helps spread those natural differences across the treatment groups so one fertilizer group is not deliberately given the stronger plants.

Module 2: Organizing and Displaying Data

Students organize qualitative and quantitative data using tables and graphs, then identify misleading displays.

Section 2.1 Organizing Qualitative Data

Build and interpret frequency tables and categorical graphs.

QUICK REFERENCE

- For qualitative data, a frequency distribution lists each category and its count.
 - Raw data are the original observations before they are organized or summarized.
 - Frequency is the number of observations in a category.
 - Relative frequency is the proportion of observations in a category. Find it by dividing the category frequency by the total frequency.
 - A relative-frequency distribution lists each category with its relative frequency, written as a decimal or percent.
- Check the totals and interpret relative frequencies in context.
 - The frequencies should total the sample size n .
 - The relative frequencies should total 1, or 100%, except for a small difference caused by rounding.
 - A relative frequency of 0.36 means 36% of the observations are in that category; it does not mean 36 observations.
- A bar graph displays qualitative categories using separated bars of equal width.
 - Place the categories on one axis and frequency or relative frequency on the other. The height of each bar gives the category's value.
 - Frequency and relative-frequency bar graphs for the same data have the same overall pattern but different vertical scales.
 - Horizontal bars are useful when category names are long.

- Special bar graphs organize bars for a particular comparison.
 - A Pareto chart orders categories from greatest to least frequency or relative frequency.
 - A side-by-side bar graph compares the same categories across two or more groups.
 - Use relative frequencies when comparing groups of different sizes so the comparison is based on proportions rather than counts.
- A pie chart divides one whole into sectors whose sizes represent category proportions.
 - The categories must represent nonoverlapping parts of the same whole, and their percentages should total 100%.
 - To construct a sector by hand, multiply the relative frequency by 360 degrees to find the sector angle.
 - A pie chart is not appropriate when the percentages overlap or do not represent parts of one whole.
- Category definitions matter. A graph can give a misleading comparison when one category combines several subcategories but the others do not.

Example 1. The table shows the commute method reported by 20 students. Construct a frequency and relative-frequency distribution. Check both totals, then identify the greatest and least frequencies.

Raw commute method responses				
Car	Bus	Walk	Car	Bike
Car	Car	Bus	Walk	Car
Bus	Bike	Car	Walk	Bus
Car	Bus	Bike	Car	Walk

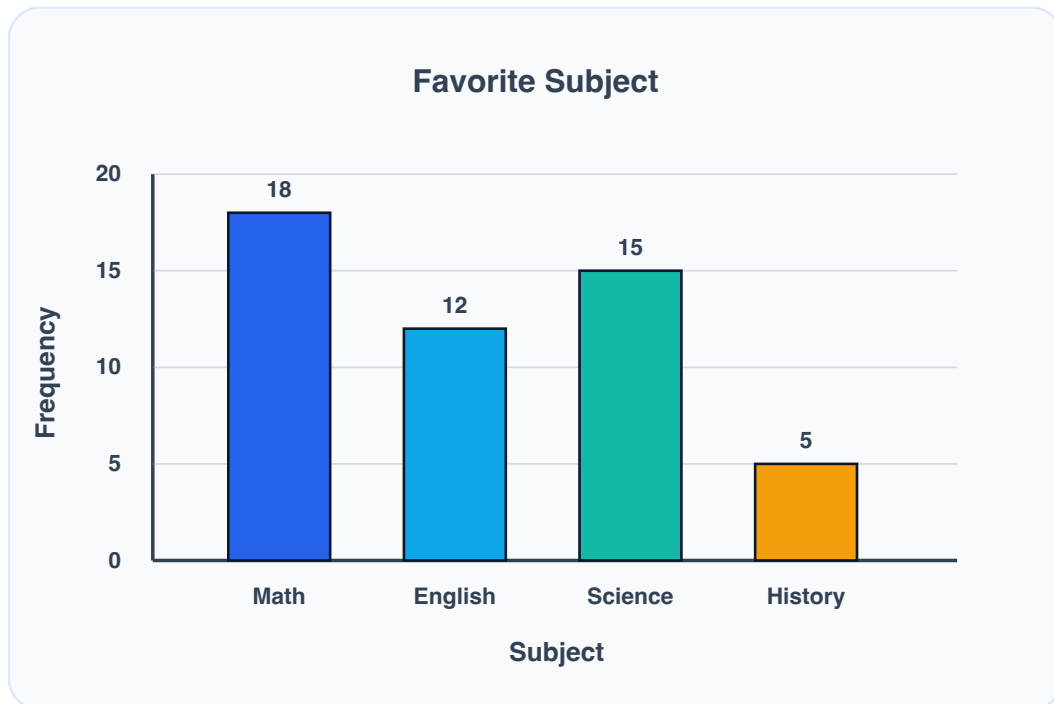
Answer: Car: frequency 8, relative frequency 0.40. Bus: 5, 0.25. Walk: 4, 0.20. Bike: 3, 0.15. The frequencies total 20, and the relative frequencies total 1.00. Car has the greatest frequency, and Bike has the least.

Commute method	Frequency	Relative frequency	Percent
Car	8	0.40	40%
Bus	5	0.25	25%
Walk	4	0.20	20%
Bike	3	0.15	15%
Total	20	1.00	100%

Example 2. Use the table and frequency bar graph to find the relative frequency for each favorite subject. If a relative-frequency bar graph were drawn, how would its pattern compare with the frequency bar graph?

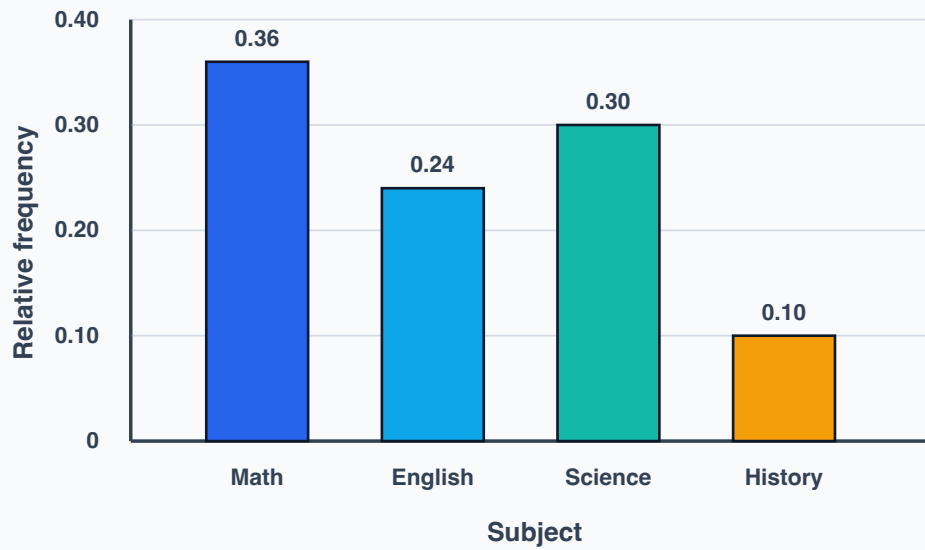
Subject	Math	English	Science	History
Frequency	18	12	15	5

Favorite Subject Bar Graph



Answer: Total: 50. Math: $0.36 = 36\%$; English: $0.24 = 24\%$; Science: $0.30 = 30\%$; History: $0.10 = 10\%$. A relative-frequency bar graph would have the same category order and proportional bar heights, but its vertical axis would show proportions or percentages instead of counts.

Favorite Subject Relative-Frequency Bar Graph



Example 3. Use the Pareto chart. Identify the greatest and least frequencies, find how many more students chose Drama than Chess, and explain why this is a Pareto chart. If the Drama category combined drama club, musical theatre, and stage crew while each other bar represented only one club, what concern would that create?

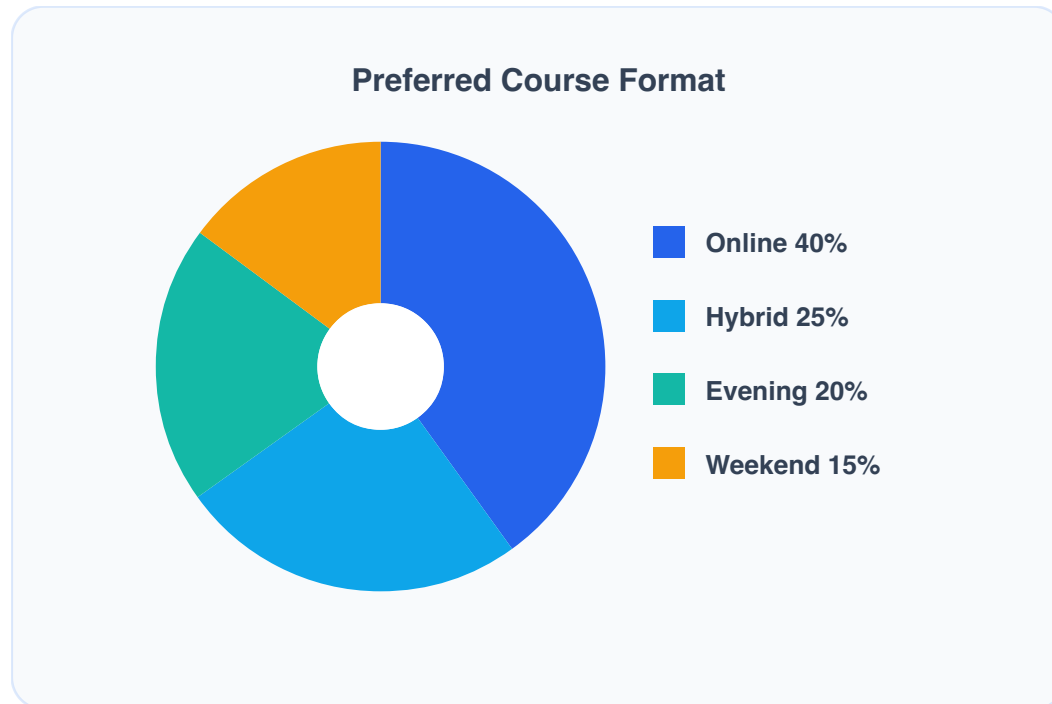
Club Signups Pareto Chart



Answer: Drama has the greatest frequency at 24, and Chess has the least at 5. Drama exceeds Chess by $24 - 5 = 19$. This is a Pareto chart because the bars are ordered from greatest to least frequency. Combining three activities into Drama would make that category broader than the others and could create a misleading comparison.

Example 4. Use the pie chart. Identify the most and least popular formats, find the percent who chose either online or hybrid, and verify that the percentages form one whole. A separate graphic lists overlapping reasons for a choice as 63%, 59%, and 52%. Could those percentages be displayed in a pie chart?

Preferred Course Format



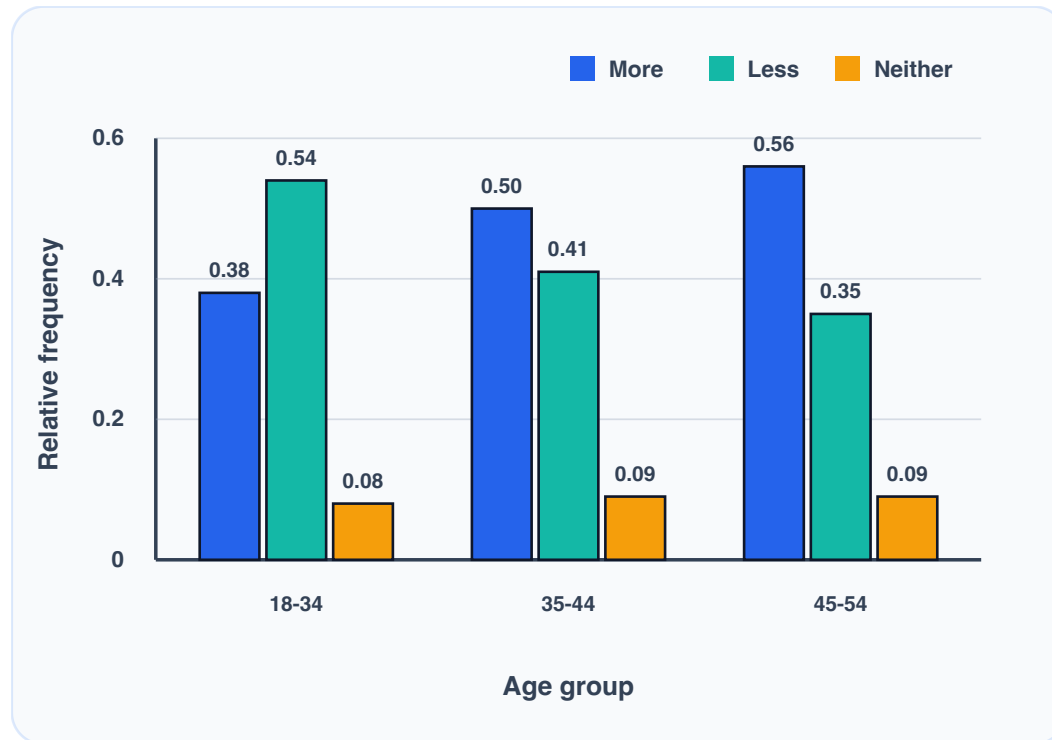
Answer: Online is most popular at 40%, and Weekend is least popular at 15%. Online or hybrid is $40\% + 25\% = 65\%$. The displayed sectors total $40\% + 25\% + 20\% + 15\% = 100\%$, so they form one whole. The overlapping reasons cannot form a pie chart because they total 174% and are not nonoverlapping parts of one whole.

Example 5. Choose the most appropriate display. (a) Show the frequency of each eye-color category. (b) Rank customer complaints from most common to least common. (c) Show how a school's annual budget is divided among nonoverlapping categories.

Answer: (a) Ordinary bar graph, because eye color is qualitative. (b) Pareto chart, because the goal is to rank categories from greatest to least frequency. (c) Pie chart, because the categories are nonoverlapping parts of one budget.

Example 6. Use the side-by-side relative-frequency bar graph. What proportion of the 18-34 group is more likely to buy a product made in the home country? Which age group has the greatest proportion answering More? Describe the apparent pattern as age increases.

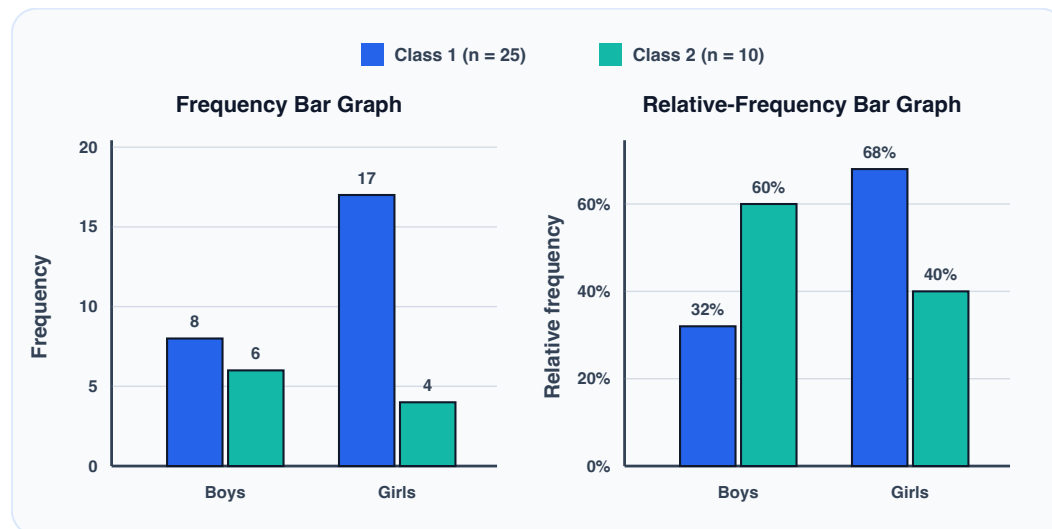
Likelihood to Buy When Made in the Home Country



Answer: The proportion for ages 18-34 answering More is 0.38, or 38%. Ages 45-54 have the greatest More proportion at 0.56. Across these groups, the More proportion increases as age increases, while the Less proportion decreases; the Neither proportion stays about the same.

Example 7. Class 1 has 8 boys and 17 girls. Class 2 has 6 boys and 4 girls. Use the side-by-side frequency and relative-frequency bar graphs to compare the two classes. Which class has the greater frequency and relative frequency for each category? Why is relative frequency the better comparison here?

Class Composition: Counts and Proportions



Answer: Class 1 has 25 students, and Class 2 has 10 students. The frequency graph shows that Class 1 has more boys (8 compared with 6) and more girls (17 compared with 4). For Class 1, the relative frequencies are $8/25 = 0.32 = 32\%$ boys and $17/25 = 0.68 = 68\%$ girls. For Class 2, they are $6/10 = 0.60 = 60\%$ boys and $4/10 = 0.40 = 40\%$ girls. Therefore, Class 2 has the greater proportion of boys, while Class 1 has the greater proportion of girls. Relative frequencies are better for comparing the class compositions because the class sizes differ.

Section 2.2 Organizing Quantitative Data: The Popular Displays

Read and create grouped frequency displays for quantitative data.

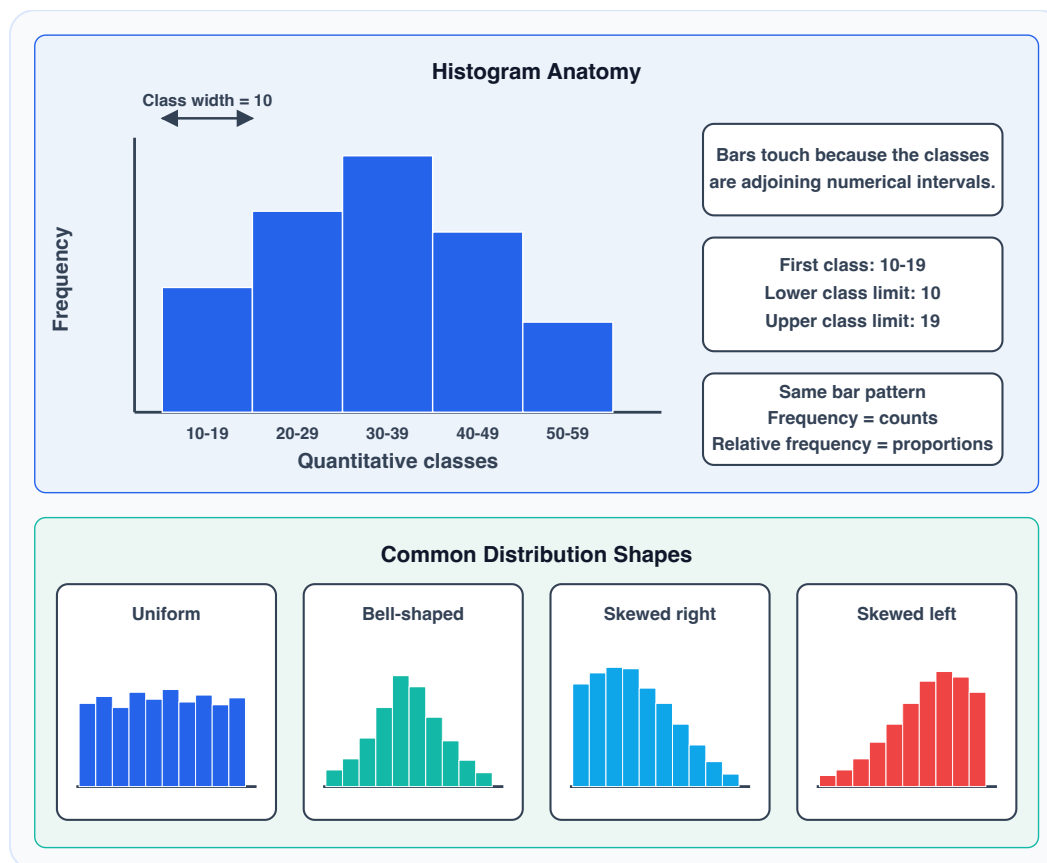
QUICK REFERENCE

- Quantitative frequency distributions may use individual values or grouped intervals.
 - For discrete data with relatively few possible values, each value can be its own class.
 - Discrete data with many values and continuous data are usually grouped into nonoverlapping intervals.
 - Frequency is the count in a class. Relative frequency is the class frequency divided by the total frequency.
 - Frequencies should total the sample size n ; relative frequencies should total 1, or 100%, except for a small difference caused by rounding.

- Class terms describe the intervals in a grouped frequency distribution.
 - The lower class limit is the smallest value that can belong to a class; the upper class limit is the largest value that can belong to that class.
 - Class width is the difference between consecutive lower class limits. For the classes 10-19 and 20-29, the width is $20 - 10 = 10$, not $19 - 10 = 9$.
 - Classes should cover all observations without overlapping. Unless an open-ended class is needed, use equal class widths.
 - If the class setup is not given, begin at the smallest observation or a convenient value below it. To estimate the class width, subtract the minimum from the maximum, divide by the desired number of classes, and round up to a convenient width.
 - Different valid class widths can produce somewhat different-looking summaries. Use enough classes to show the pattern without making the display unnecessarily detailed.
- A histogram displays a quantitative frequency distribution.
 - Place the classes on the horizontal axis and frequency or relative frequency on the vertical axis.
 - The bars have equal widths and touch because the classes represent adjoining numerical intervals. In a bar graph for qualitative data, the bars are separated.
 - Frequency and relative-frequency histograms for the same classes have the same shape and proportional bar heights, but different vertical scales.
 - A histogram shows how many observations fall in each class, but it usually does not show the exact values within a class.
- A dot plot displays every observation in a small quantitative data set.
 - Place each possible value on a number line and draw one dot for each observation. Stack repeated observations above the same value.
 - Dot plots preserve individual values, making clusters, gaps, and repeated values easy to see. They become crowded for large data sets.
- When a question includes several classes, add the relevant frequencies or relative frequencies.
 - At least includes the stated value and all larger values.
 - At most includes the stated value and all smaller values.
- Describe a quantitative distribution using its overall pattern and the direction of any tail.
 - Symmetric: the left and right sides are approximately mirror images.
 - Bell-shaped: symmetric with one central peak and frequencies that taper in both directions.
 - Uniform: the bar heights are approximately equal.
 - Skewed right: the longer tail points to the right.
 - Skewed left: the longer tail points to the left.
 - Real data may only approximate these shapes, so reasonable judgments can differ. Do not use these shape terms for qualitative data.

- TI-84 histogram workflow:
 - Enter raw data in L1 using STAT → EDIT.
 - Open 2nd Y= (STAT PLOT), turn Plot1 on, choose the histogram icon, and set Xlist:L1 and Freq:1. Use ZoomStat for an initial window.
 - For specified classes, set Xmin to the first lower class limit and Xscl to the class width.
 - For a grouped table, enter class midpoints in L1, frequencies in L2, and use Freq:L2. StatCrunch is convenient for larger data sets.

Histograms at a Glance



Example 1. Decide how each quantitative data set should be organized. (a) The number of televisions in a household can be 0, 1, 2, 3, ..., and only a few values occur. (b) Commute time is measured to the nearest tenth of a minute and many different values occur. Identify each variable as discrete or continuous and decide whether individual-value classes or interval classes are more appropriate.

Answer: (a) Number of televisions is discrete. Because only a few values occur, each value can be its own class. (b) Commute time is continuous. Group the values into nonoverlapping intervals.

Example 2. Use the grouped frequency table. How many classes are shown? State the lower and upper class limits of the first class, find the class width, find the sample size, and identify the class with the greatest frequency.

Class	10-19	20-29	30-39	40-49
Frequency	6	11	9	4

Answer: There are 4 classes. The first lower class limit is 10, and the first upper class limit is 19. The class width is $20 - 10 = 10$. The sample size is $6 + 11 + 9 + 4 = 30$. The 20-29 class has the greatest frequency, 11.

Example 3. The table lists the number of customers waiting on 20 Saturdays. Using the classes 1-3, 4-6, 7-9, 10-12, and 13-15, construct frequency and relative-frequency distributions. Then find the percentage of Saturdays with at least 10 customers and the percentage with at most 6 customers.

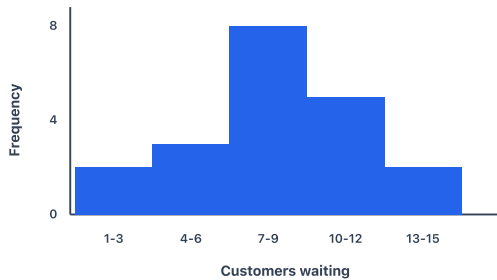
Customers waiting				
1	3	4	5	6
7	7	8	8	8
9	9	9	10	10
11	12	12	13	15

Answer: The class frequencies are 2, 3, 8, 5, 2, and the relative frequencies are 0.10, 0.15, 0.40, 0.25, 0.10. At least 10: $(5 + 2)/20 = 0.35 = 35\%$. At most 6: $(2 + 3)/20 = 0.25 = 25\%$.

Class	Frequency	Relative frequency	Percent
1-3	2	0.10	10%
4-6	3	0.15	15%
7-9	8	0.40	40%
10-12	5	0.25	25%
13-15	2	0.10	10%
Total	20	1.00	100%

Example 4. The two histograms display the customer data from Example 3. Identify the modal class, compare the shapes and vertical scales, and give a TI-84 setup for creating the histogram from either the raw data or the grouped table.

Frequency Histogram



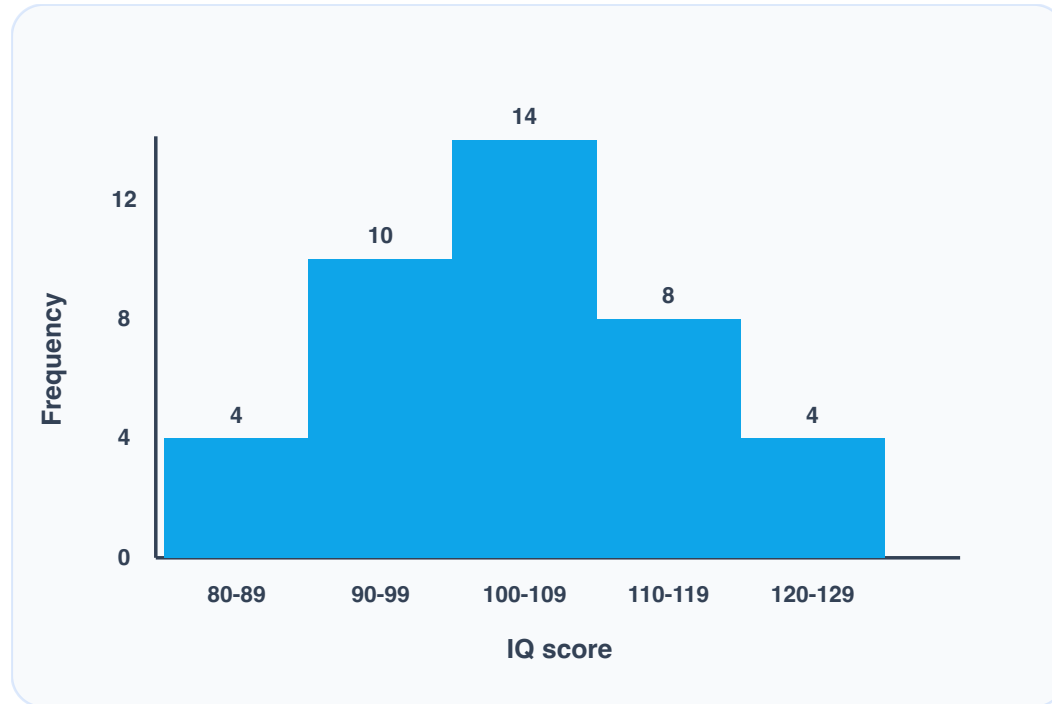
Relative-Frequency Histogram



Answer: The modal class is 7-9. The histograms have the same shape and proportional bar heights. The left vertical scale uses counts, while the right uses proportions. For raw data, enter the observations in L1, use Plot1 with the histogram icon, Xlist:L1, Freq:1, Xmin=1, and Xscl=3. For the grouped table, enter midpoints 2, 5, 8, 11, 14 in L1, frequencies 2, 3, 8, 5, 2 in L2, and use Freq:L2 with the same window settings.

Example 5. Use the IQ-score histogram. Find the sample size, identify the class with the greatest frequency and the class or classes with the least frequency, and find the percentage of students with scores of at least 110. Can the histogram tell us exactly how many students scored 115?

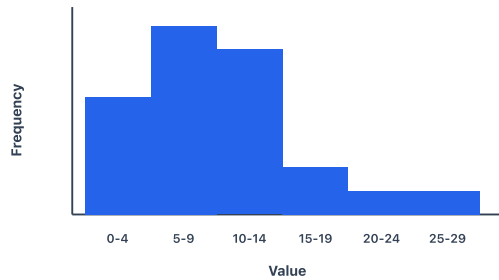
Student IQ Scores



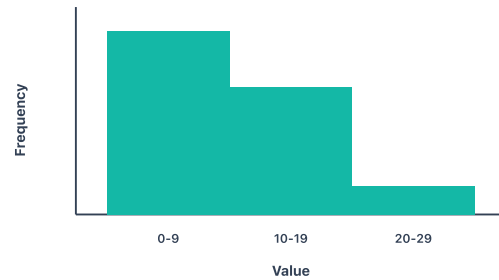
Answer: The sample size is $4 + 10 + 14 + 8 + 4 = 40$. The 100-109 class has the greatest frequency, 14. The 80-89 and 120-129 classes tie for the least frequency, 4. At least 110: $(8 + 4)/40 = 0.30 = 30\%$. No. The histogram shows that 8 scores are between 110 and 119, but it does not show their exact values.

Example 6. The histograms summarize the same 24 observations using class widths of 5 and 10. Describe the overall shape and explain the tradeoff between the two class widths. Is one class width always correct?

Class Width 5



Class Width 10



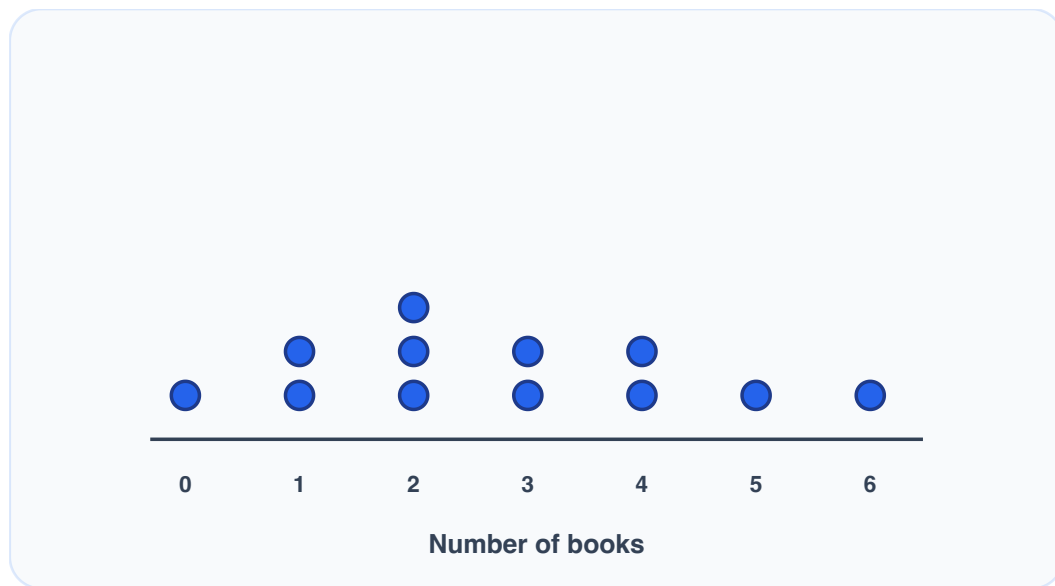
Answer: Both histograms show an overall right-skewed pattern. Width 5 gives more detail about how the observations are distributed, while width 10 makes the overall pattern easier to see. Neither width is automatically the one correct choice; the better choice depends on which display summarizes the data clearly without hiding the pattern or adding unnecessary detail.

Example 7. The table shows the number of books read last month by 12 students. Construct a dot plot and identify the most common value.

Books read last month					
0	1	1	2	2	2
3	3	4	4	5	6

Answer: Place one dot for every observation and stack repeated values. The most common value is 2 books because it has three dots.

Books Read Last Month



Section 2.3 Additional Displays of Quantitative Data

Construct and interpret dot plots, stem-and-leaf plots, frequency polygons, ogives, and time-series plots.

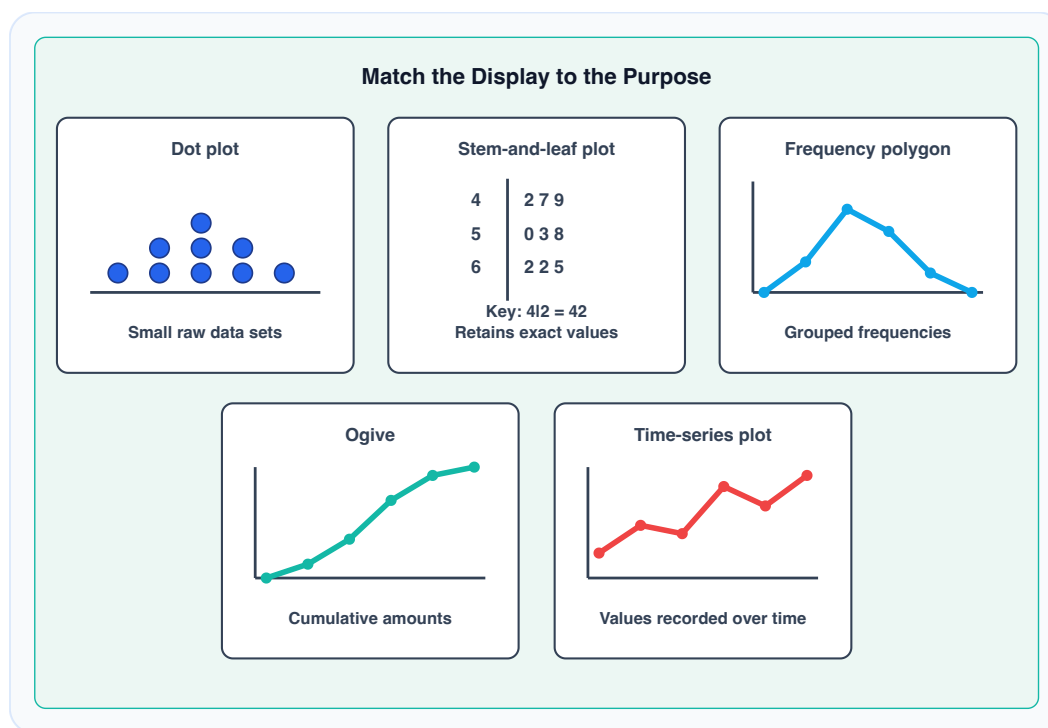
QUICK REFERENCE

- A dot plot displays each observation as a dot above its value.
 - Repeated values form vertical stacks. Count all dots to find the sample size, and use the tallest stack to identify the value that occurs most often.
 - Dot plots work best for relatively small raw data sets because they preserve individual values.
 - Look for the overall shape, clusters, gaps, and observations separated from the rest of the data.

- A stem-and-leaf plot organizes quantitative data while preserving the original values.
 - The stem contains the leading digit or digits, and the leaf contains the final digit. A key such as 13 | 5 = 13.5 explains the place value.
 - The same 13 | 5 notation could represent 135, 13.5, or another value under a different place-value setup, so always use the plot's key.
 - Write stems in increasing order, place every leaf with its stem, and order the leaves from least to greatest.
 - Split stems when one row is crowded. Repeat each stem, using the first row for leaves 0–4 and the second row for leaves 5–9.
 - A back-to-back stem-and-leaf plot uses shared stems to compare the distributions of two groups.
- A frequency polygon displays grouped frequencies using connected points.
 - Find each class midpoint by adding the lower and upper class limits and dividing by 2.
 - Plot each class midpoint on the horizontal axis and its frequency on the vertical axis, then connect consecutive points.
 - Add one midpoint before the first class and one after the last class, each with frequency zero, so the graph meets the horizontal axis.
 - Frequency polygons show the shape of grouped data and can compare multiple distributions on the same axes.
- Cumulative distributions use running totals.
 - Cumulative frequency is the number of observations at or below a class's upper endpoint.
 - Cumulative relative frequency is the proportion or percentage of observations at or below that endpoint.
 - The final cumulative frequency should equal n , and the final cumulative relative frequency should equal 1, or 100%.
- An ogive graphs cumulative frequency or cumulative relative frequency.
 - Plot upper class endpoints on the horizontal axis and cumulative values on the vertical axis. Class boundaries may be used when the classes are written with whole-number limits.
 - Begin at zero using the lower boundary before the first class, then connect the points. Because cumulative totals cannot decrease, an ogive must stay level or rise.
 - Use an ogive to find the amount at or below a cutoff, estimate a cutoff for a given cumulative percentage, or find the amount above a cutoff by subtracting the cumulative amount from the total.
- A time-series plot displays measurements recorded over time.
 - Place time on the horizontal axis and the measured values on the vertical axis, then connect the points chronologically.
 - Describe the overall trend, peaks, lows, changes in direction, and unusual changes.
 - When two series share the same axes, compare their values, trends, and the times when they are closest or farthest apart.

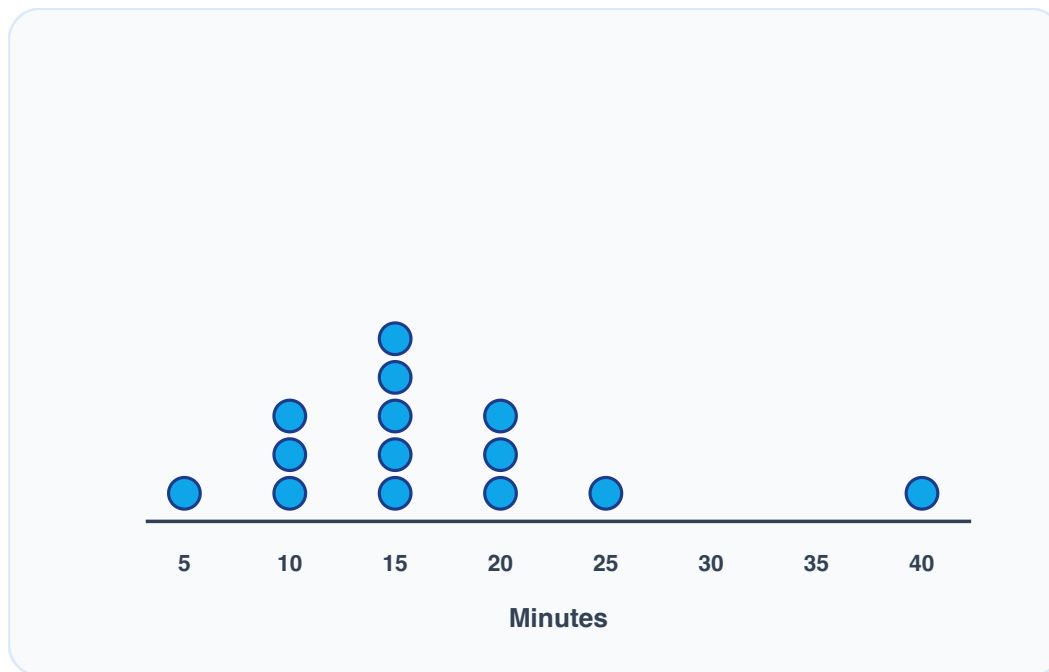
- Choose the display that matches the data and the question.
 - Use a dot plot for a small raw data set and a stem-and-leaf plot when retaining exact values is especially useful.
 - Use a frequency polygon for grouped frequencies, an ogive for cumulative amounts or cutoff questions, and a time-series plot for values recorded over time.
- TI-84 workflows for connected statistical plots:
 - For a frequency polygon, enter class midpoints in L1 and frequencies in L2, including the two zero-frequency endpoints. Open 2nd Y=, choose the xyLine plot, and set Xlist:L1 and Ylist:L2.
 - For cumulative totals, enter frequencies in L2 and calculate $\text{cumSum}(L2) \rightarrow L3$. Use $L3/\text{sum}(L2) \rightarrow L4$ for cumulative relative frequencies.
 - For an ogive, graph upper class endpoints against L3 or L4 with an xyLine plot. Include the starting point with cumulative value zero.
 - For a time series, enter times in L1 and values in L2, then use an xyLine plot. Turn on a second Stat Plot to compare another series.
 - The TI-84 has no dedicated dot-plot or stem-and-leaf command. Construct those displays by hand or use StatCrunch when the data set is large.

Choosing an Additional Display



Example 1. The dot plot shows the one-way commute times of 14 students, rounded to the nearest 5 minutes. Find the sample size and the commute time with the tallest stack. Identify the main cluster, a gap, and a value that may be considered unusual.

Student Commute Times



Answer: There are 14 dots, so $n = 14$. The tallest stack is at 15 minutes, with 5 dots. Most commute times cluster from 10 through 20 minutes. There is a gap at 30 and 35 minutes, and the 40-minute commute is separated from the rest and may be considered unusual for this sample.

Example 2. The stem-and-leaf plot shows the weights, in ounces, of 16 packages. Use the key to reconstruct the original data set. Then state the sample size, minimum weight, and maximum weight.

Stem	Leaves
13	5 7 9 9
14	1 3 6 7 8 9 9
15	1 4 5
16	0 2

Key: 13 | 5 represents 13.5 ounces.

Answer: The package weights are 13.5, 13.7, 13.9, 13.9, 14.1, 14.3, 14.6, 14.7, 14.8, 14.9, 14.9, 15.1, 15.4, 15.5, 16.0, 16.2 ounces. Thus $n = 16$, the minimum weight is 13.5 ounces, and the maximum weight is 16.2 ounces.

Example 3. Construct an ordered stem-and-leaf plot for the data. Then construct a split stem-and-leaf plot, using the first occurrence of each stem for leaves 0-4 and the second for leaves 5-9. Which plot is easier to read, and what is the overall shape?

Delivery time (minutes)				
11	12	13	13	14
15	15	16	17	18
19	21	22	24	27
28	32	36	41	48

Answer: The split plot is easier to read because the values in the teens do not crowd one row. The distribution is right-skewed: most values are in the teens and twenties, with a tail extending through the forties.

Ordered stems

Stem	Leaves
1	1 2 3 3 4 5 5 6 7 8 9
2	1 2 4 7 8
3	2 6
4	1 8

Split stems

Stem	Leaves
1	1 2 3 3 4
1	5 5 6 7 8 9
2	1 2 4
2	7 8
3	2
3	6
4	1
4	8

Key: 1 | 1 represents 11 minutes.

Example 4. The back-to-back stem-and-leaf plot compares training completion times, in minutes, for employees assigned to Methods A and B. Which method generally has longer completion times? Identify the stems where the methods overlap and give evidence from the plot.

Method A leaves	Stem	Method B leaves
9 7 4	2	5 8
8 6 3 1	3	0 2 4 5 9
7 2	4	1 3 6 8
	5	2

Key: 1 | 3 | 0 represents 31 minutes for Method A and 30 minutes for Method B.

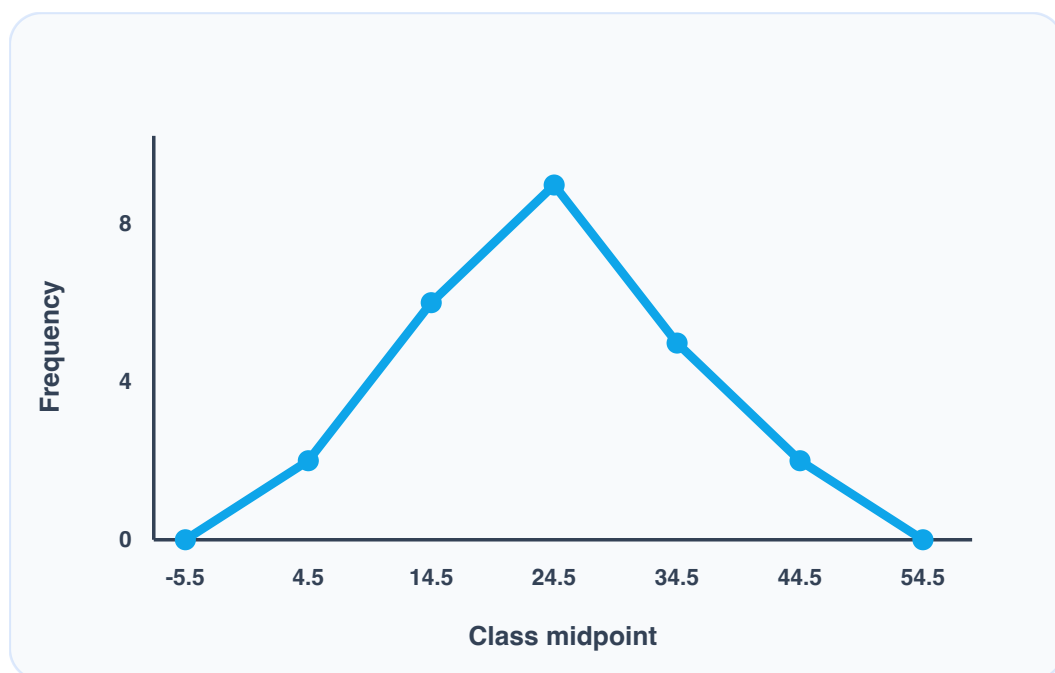
Answer: Method B generally has longer completion times. Both methods have times in the 20s, 30s, and 40s. Method A's times are concentrated in the 20s and 30s, while Method B's times are concentrated in the 30s and 40s and extend into the 50s.

Example 5. Use the grouped frequency table to find the class width and all class midpoints. Identify the class with the greatest frequency, then construct a frequency polygon. Include the two zero-frequency endpoints and give a TI-84 setup.

Wait time	0-9	10-19	20-29	30-39	40-49
Frequency	2	6	9	5	2

Answer: The class width is 10, and the midpoints are 4.5, 14.5, 24.5, 34.5, 44.5. The class 20-29 has the greatest frequency, with 9 observations. Plot $(-5.5, 0)$, $(4.5, 2)$, $(14.5, 6)$, $(24.5, 9)$, $(34.5, 5)$, $(44.5, 2)$, and $(54.5, 0)$. Connect the points. On a TI-84, place the seven x -coordinates in L1, the seven frequencies in L2, and use an xyLine Stat Plot with Xlist:L1 and Ylist:L2.

Customer Wait-Time Frequency Polygon



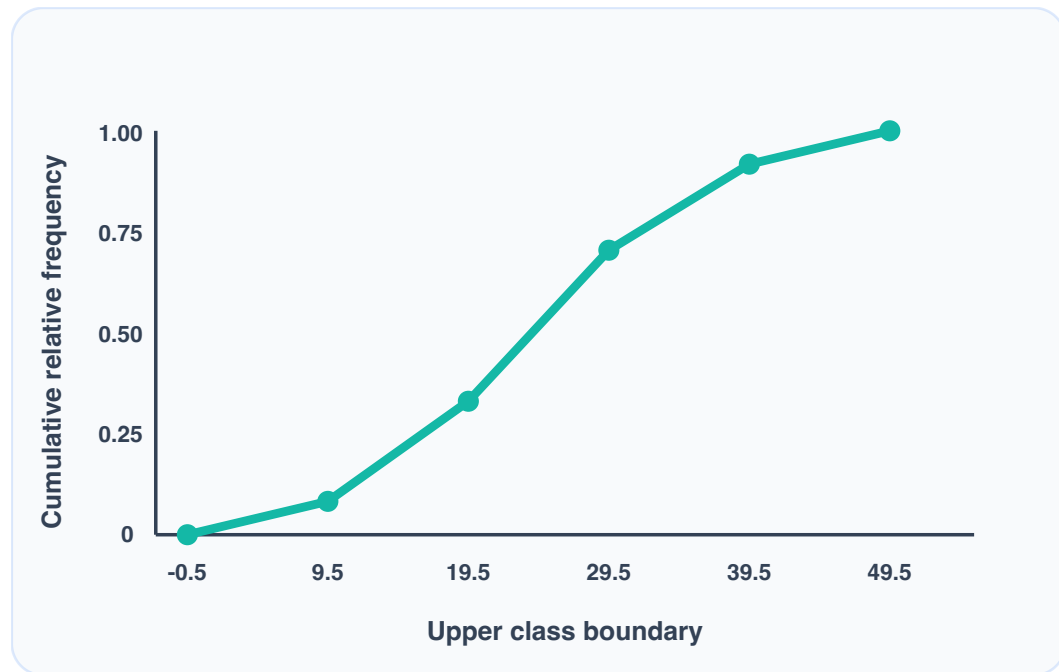
Example 6. Use the frequency distribution from Example 5. Complete cumulative frequency and cumulative relative-frequency columns, then list the points for a relative-frequency ogive using class boundaries. What percentage of waits are at most 29 minutes? Which upper class boundary first reaches at least 90%? What percentage are above 19 minutes?

Wait time	Frequency	Cumulative frequency	Cumulative relative frequency
0-9	2	?	?
10-19	6	?	?
20-29	9	?	?
30-39	5	?	?
40-49	2	?	?

Answer: The cumulative frequencies are 2, 8, 17, 22, 24, and the cumulative relative frequencies are 0.083, 0.333, 0.708, 0.917, 1.000. The ogive points are (-0.5, 0), (9.5, 0.083), (19.5, 0.333), (29.5, 0.708), (39.5, 0.917), and (49.5, 1.000). At most 29 minutes: $17/24$, approximately 70.8%. The first boundary reaching at least 90% is 39.5. Above 19 minutes: $1 - 8/24 = 16/24$, approximately 66.7%.

Wait time	Frequency	Cumulative frequency	Cumulative relative frequency
0-9	2	2	0.083
10-19	6	8	0.333
20-29	9	17	0.708
30-39	5	22	0.917
40-49	2	24	1.000

Relative-Frequency Ogive

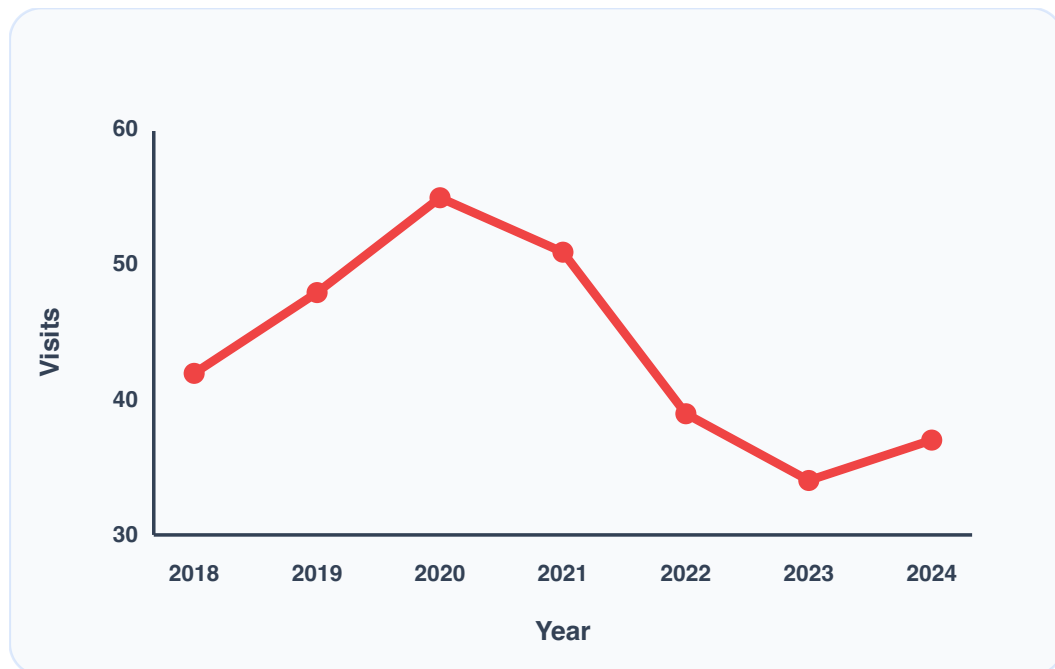


Example 7. The table gives annual tutoring-center visits. Construct a time-series plot, identify the highest and lowest years, describe the changing trend, and give a TI-84 setup.

Year	2018	2019	2020	2021	2022	2023	2024
Visits	42	48	55	51	39	34	37

Answer: Visits are highest in 2020 at 55 and lowest in 2023 at 34. Visits rise from 2018 through 2020, fall through 2023, and rise slightly in 2024. For a TI-84, enter the years in L1, visits in L2, and use an xyLine Stat Plot with Xlist:L1 and Ylist:L2.

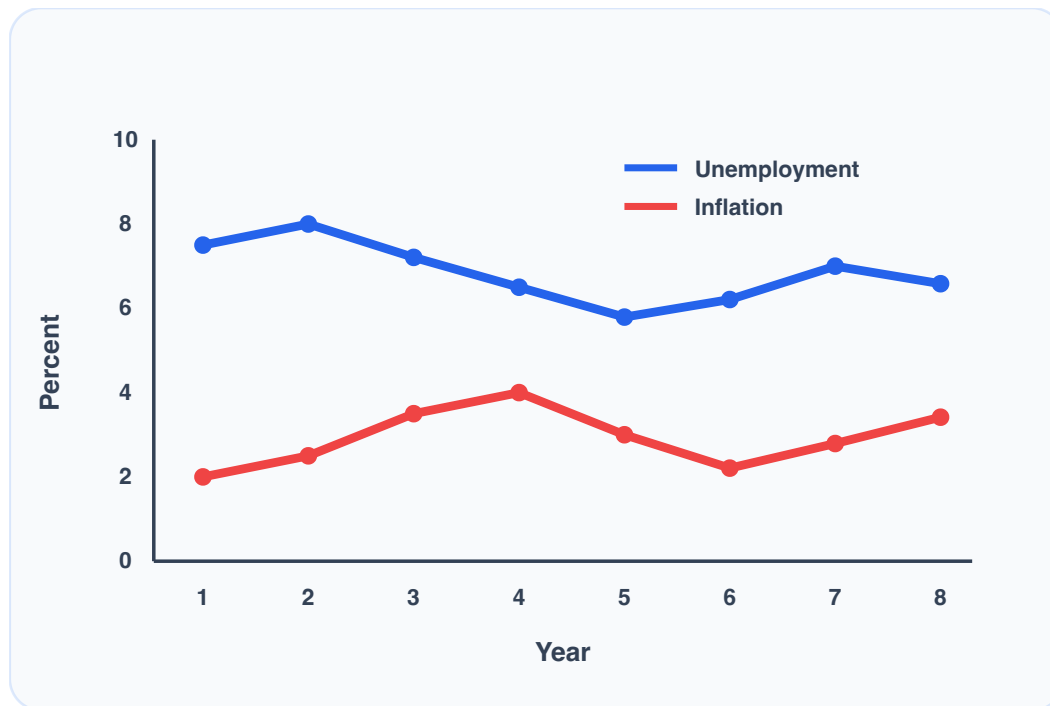
Tutoring Center Visits



Example 8. The table and time-series plot compare a region's annual unemployment and inflation rates over 8 years. The misery index is a simple measure of economic strain found by adding the unemployment rate and inflation rate for the same year. A higher index represents a greater combined burden from unemployment and rising prices. Use the table for exact calculations. In which year is the misery index highest, and in which year is it lowest? List each year and its misery-index value.

Rate (%)	Year 1	Year 2	Year 3	Year 4	Year 5	Year 6	Year 7	Year 8
Unemployment	7.5	8.0	7.2	6.5	5.8	6.2	7.0	6.6
Inflation	2.0	2.5	3.5	4.0	3.0	2.2	2.8	3.4

Unemployment and Inflation Rates



Answer: Highest: Year 3, with a misery index of $7.2\% + 3.5\% = 10.7\%$. Lowest: Year 6, with a misery index of $6.2\% + 2.2\% = 8.4\%$.

Example 9. Choose the most appropriate display for each purpose. (a) Preserve the exact values in a small set of exam scores. (b) Show the grouped frequencies of customer ages and compare the overall shape with another store. (c) Estimate the percentage of scores at or below a cutoff. (d) Show monthly enrollment for the past three years. (e) Display a small raw data set so clusters, gaps, and unusual values are easy to see.

Answer: (a) Stem-and-leaf plot. (b) Frequency polygon. (c) Ogive. (d) Time-series plot. (e) Dot plot.

Section 2.4 Graphical Misrepresentations of Data

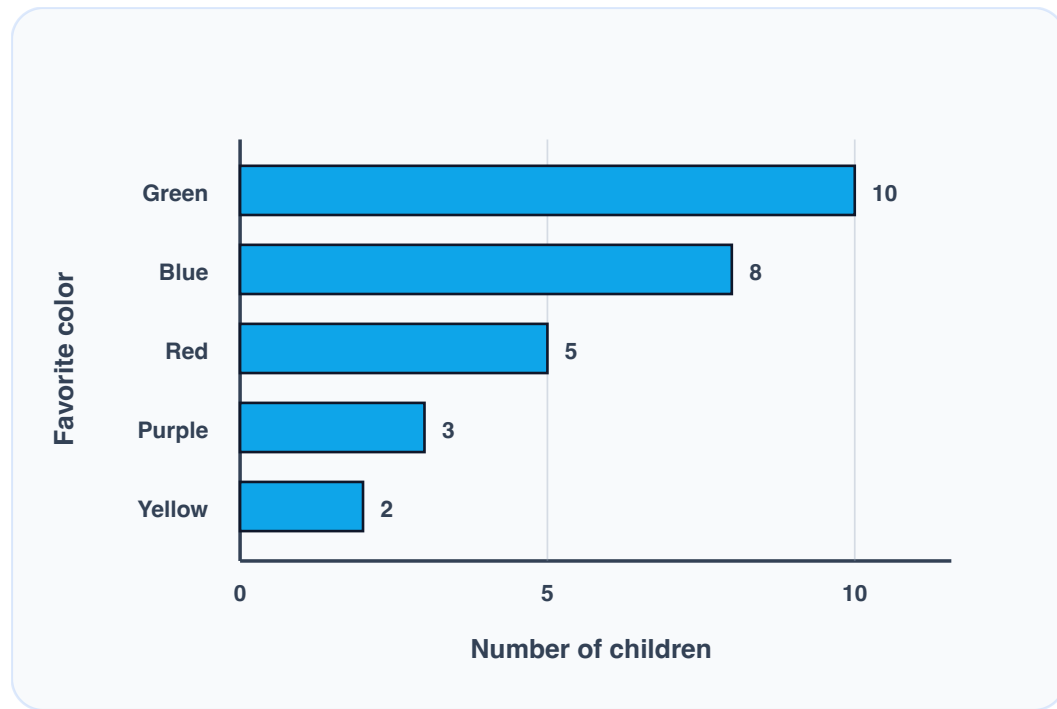
Identify graphs that distort or exaggerate data.

QUICK REFERENCE

- Check the scale before interpreting a graph.
 - Axis values should be clearly labeled and increase by consistent amounts.
 - A truncated vertical axis can exaggerate or hide differences. Bar graphs should generally begin at zero because the length of each bar represents the amount.
 - If a scale is truncated for a valid reason, the break or shortened scale should be clearly shown.
- Check what each bar or class represents.
 - Bars used to compare categories should have equal visual widths.
 - Percentages or counts from time intervals of different lengths do not provide a fair comparison of how often something occurs.
 - For introductory histograms, use equal class widths unless the graph clearly explains how unequal widths are being handled.
- Check pictures and visual effects for distortion.
 - If a picture is enlarged in both height and width, its area grows much faster and can exaggerate the difference.
 - Decorative pictures can make it unclear where a bar begins or ends.
 - Three-dimensional graphs can make categories closer to the viewer appear larger. A clear two-dimensional graph is usually easier to compare.
- Check the labels and context.
 - Use an informative title, clearly labeled axes, units, and enough context to identify the data.
 - Include exact values or a data source when they help the reader judge the display.
- Decide whether the graph is actually misleading.
 - A graph is not misleading simply because one category is much larger or smaller than another. The display is honest when its lengths, areas, scales, and labels accurately represent the data.
 - To correct a misleading graph, preserve the data while using a clear scale, comparable intervals, simple shapes, and complete labels.

Example 1. The bar graph shows the favorite colors of 28 children. Is it misleading? Explain.

Favorite Colors of 28 Children



Answer: No. The bars begin at zero, have equal widths, and have lengths that match the labeled frequencies. Green is correctly shown as the most common response, and yellow is correctly shown as the least common response.

Example 2. Both graphs show the same median wait times: 50, 52, 54, and 55 minutes. Why is the left graph misleading, and what does the right graph communicate more honestly?

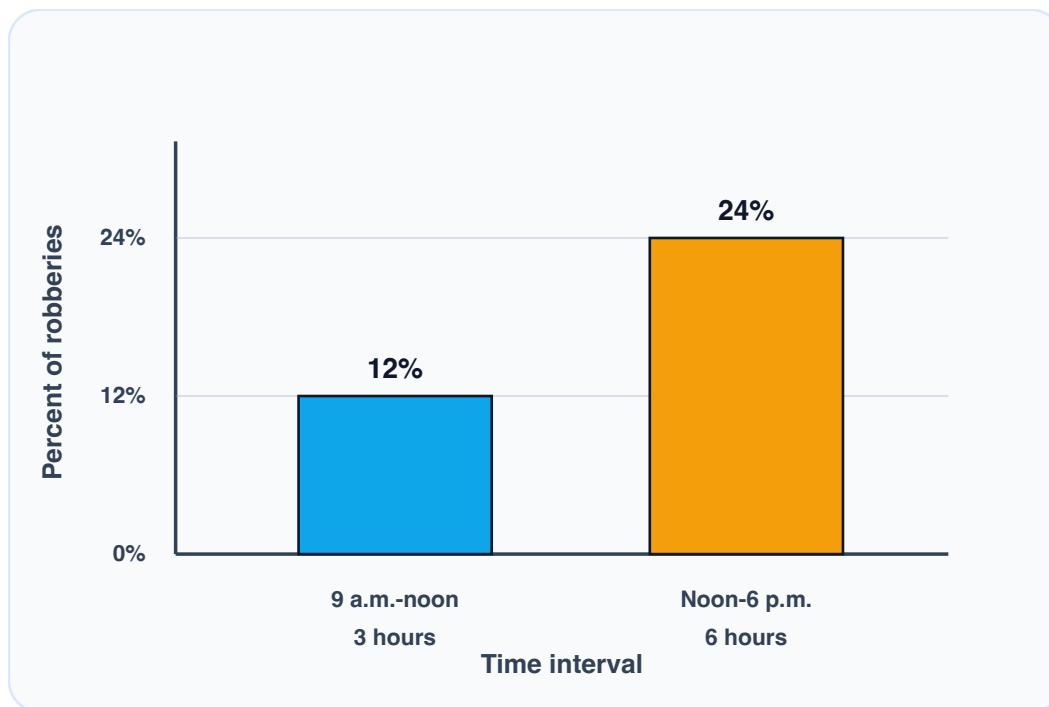
Same Data, Different Vertical Scales



Answer: The left graph starts at 48 minutes, so the bars make a 5-minute increase look dramatic. The zero-based graph shows that the median wait time increased, but only slightly compared with the total bar lengths.

Example 3. A report claims robberies are twice as common from noon to 6 p.m. as from 9 a.m. to noon because the graph shows 24% versus 12%. Why does the graph not support that comparison?

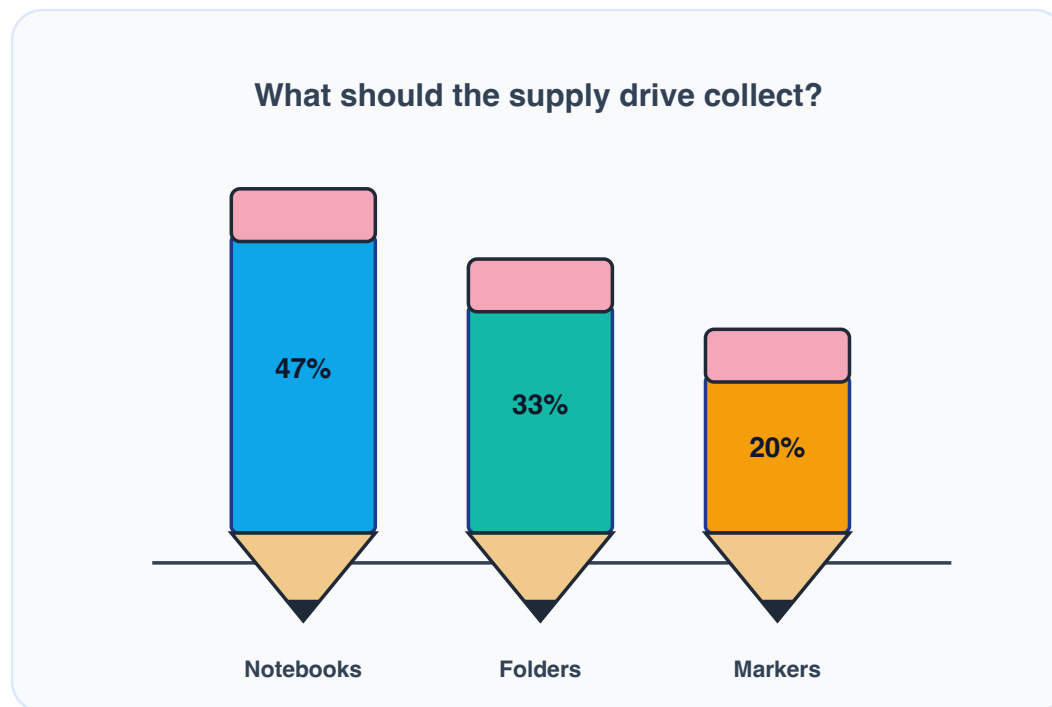
Percent of Robberies by Time Interval



Answer: The noon-to-6 p.m. interval covers 6 hours, while the 9 a.m.-to-noon interval covers only 3 hours. The longer interval has twice the percentage over twice as much time, so the graph does not show a higher hourly occurrence. Equal-length time intervals are needed for a fair comparison.

Example 4. A survey graphic uses pencils as bars. Identify two problems with the display and describe a better graph.

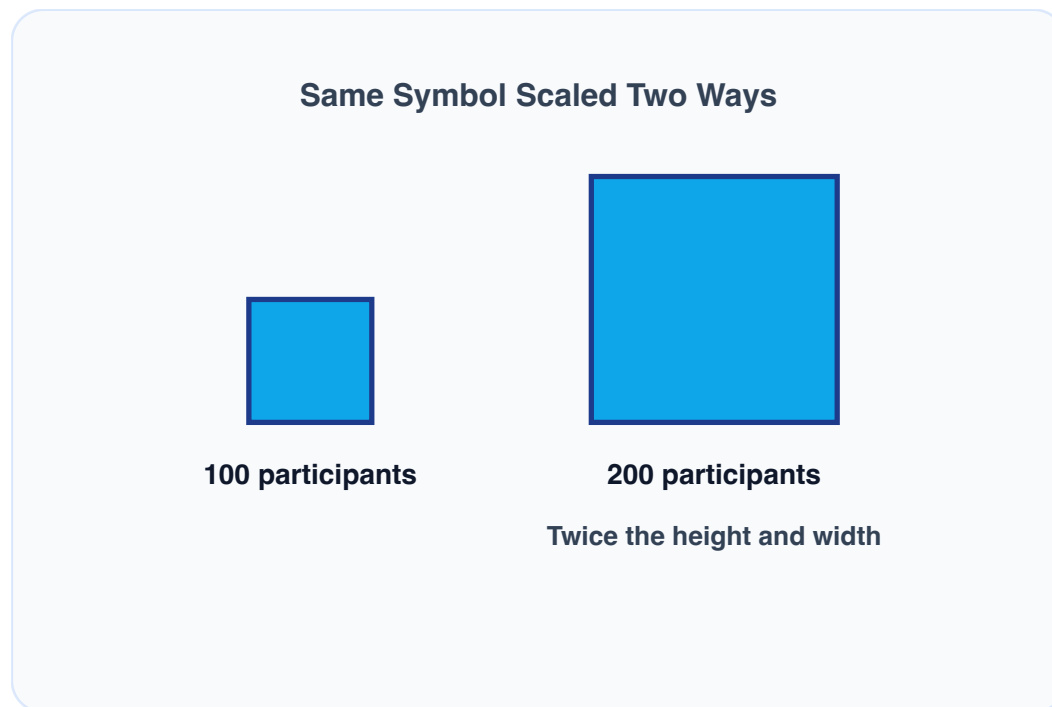
Preferred School-Supply Donation



Answer: It is unclear whether each pencil's length includes the eraser and sharpened tip, and there is no numbered scale for checking the printed percentages. A plain bar graph with a zero-based percentage axis, equal-width bars, and clear endpoints would be easier to compare.

Example 5. A pictograph doubles both the height and width of a symbol to represent twice as many participants. Why is that misleading?

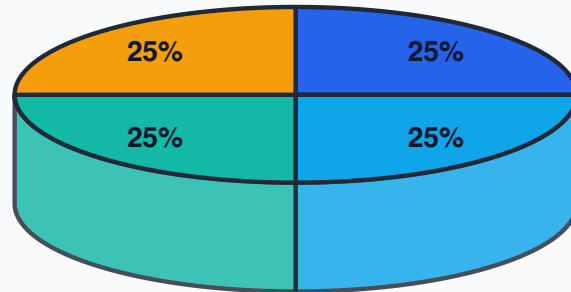
Pictograph Area Distortion



Answer: Doubling both dimensions makes the area four times as large, so the picture makes a twofold increase look like a fourfold increase. Use equally sized symbols with a key or change only one dimension.

Example 6. Each category represents 25% of the responses. Why can the three-dimensional display still be misleading?

Equal Shares Shown in Three Dimensions



All four categories have equal percentages.

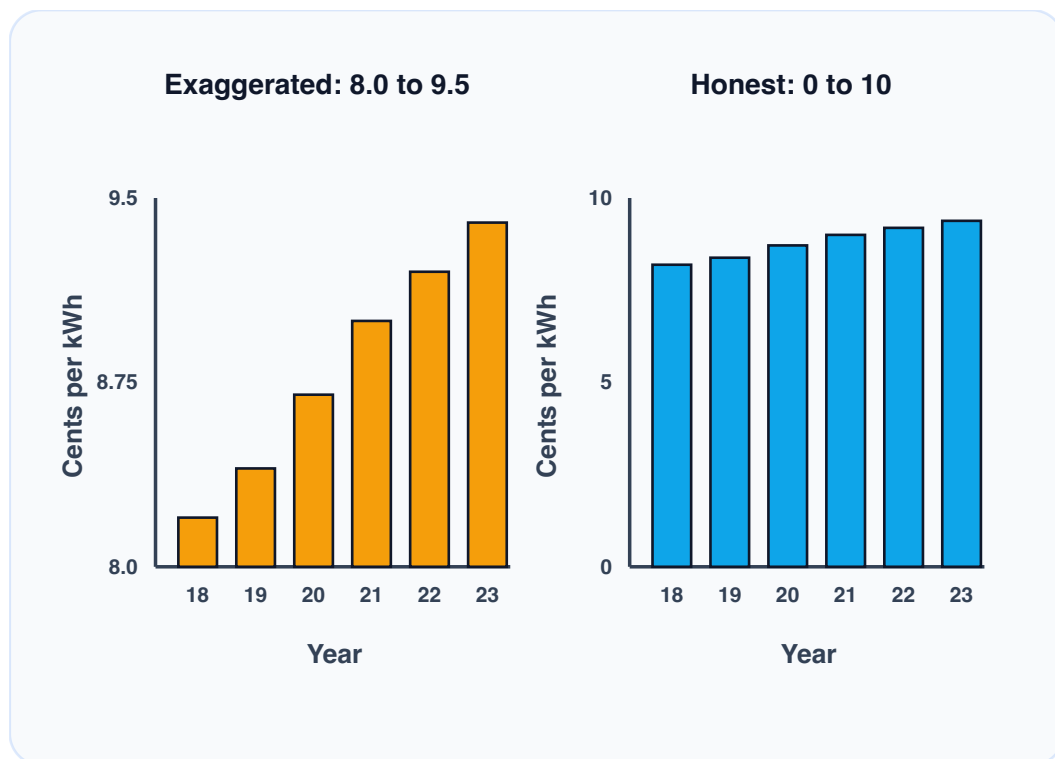
Answer: The front half has a visible side wall and appears closer to the viewer, so its categories seem more prominent even though all four percentages are equal. A flat two-dimensional circle graph would show the equal shares more clearly.

Example 7. Use the table to describe a vertical scale that would exaggerate the increase and a scale that would display the increase more honestly.

Year	2018	2019	2020	2021	2022	2023
Price per kWh (cents)	8.2	8.4	8.7	9.0	9.2	9.4

Answer: A vertical scale from 8.0 to 9.5 cents would exaggerate the increase because the first bar would appear very short while the last would nearly fill the graph. A zero-based scale from 0 to 10 cents shows the same increase without exaggerating the bar lengths.

Electricity Prices with Two Vertical Scales



Module 3: Descriptive Statistics

Students compute and interpret center, spread, position, outliers, five-number summaries, and boxplots.

Section 3.1 Measures of Central Tendency

Find and interpret mean, median, and mode.

QUICK REFERENCE

- Measures of central tendency describe the center or typical value of a distribution.
 - The mean uses every numerical value. Add the values and divide by the number of observations.
 - For a sample, $\bar{x} = \frac{\sum x}{n}$. For a population, $\mu = \frac{\sum x}{N}$.
 - Unless a problem gives another instruction, round the mean to one more decimal place than the original data.
- The median, M , is the middle of the ordered data.
 - First arrange the values from least to greatest.
 - If n is odd, the median is the single middle observation.
 - If n is even, the median is the mean of the two middle observations. The result does not have to be one of the original data values.
- The mode is the value or category that occurs most often.
 - A data set may have no mode, one mode, two modes (bimodal), or more than two modes (multimodal).
 - The mode can describe qualitative data because categories can be counted even when arithmetic is not meaningful.
 - The TI-84 1-Var Stats screen does not report the mode; inspect the list or a frequency table.
- A resistant measure is not changed substantially by an extreme value.
 - The median is resistant, but the mean is pulled toward unusually large or small values.
 - For a roughly symmetric distribution without strong outliers, the mean and median are usually close and the mean is commonly reported.
 - For a strongly skewed distribution or one with extreme values, the median usually better represents a typical observation.
 - Skewed left: mean less than median. Roughly symmetric: mean approximately equal to median. Skewed right: mean greater than median.
- Interpret a measure of center in the context and units of the variable.
 - The word average is sometimes used loosely, so identify whether a report means the mean, median, or mode.
 - Do not calculate a mean for qualitative labels or numerical codes whose arithmetic has no meaning.
- TI-84 1-Var Stats workflow:
 - Press STAT → 1:Edit and enter the raw data in L1.
 - Press STAT → CALC → 1:1-Var Stats, set List:L1, and choose Calculate.
 - Read $\bar{x}$ for the sample mean and n for the number of observations. Scroll down to Med for the median.

Example 1. The number of minutes six students spent completing a quiz were 12, 15, 15, 18, 20, and 22. Find and interpret the mean, median, and mode.

Answer: Mean: 17 minutes. Median: 16.5 minutes. Mode: 15 minutes. The students averaged 17 minutes, half of the times were at or below 16.5 minutes and half were at or above 16.5 minutes, and 15 minutes occurred most often.

Example 2. Find the mode in each data set. (a) 3, 4, 5, 6. (b) 2, 2, 3, 4, 4. (c) Bus, Walk, Bus, Bike, Walk, Bus.

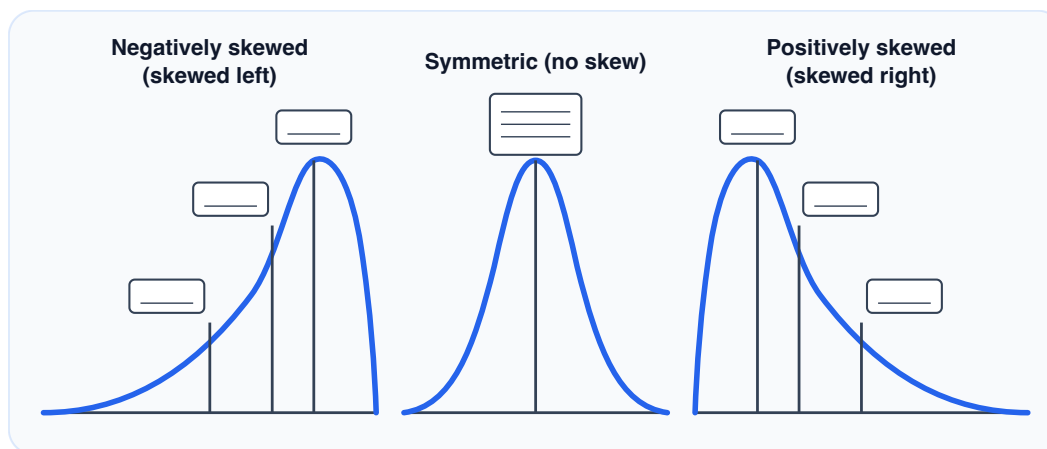
Answer: (a) No mode because no value repeats. (b) The modes are 2 and 4, so the data are bimodal. (c) Bus is the mode because it occurs three times. This also shows that a mode can be a qualitative category.

Example 3. The six ordered values are 7, 12, 21, x , 41, and 49. If the median is 28.5, find x .

Answer: With six values, the median is the mean of the third and fourth values: $(21 + x)/2 = 28.5$. Therefore, $21 + x = 57$, so $x = 36$.

Example 4. Fill each blank box with mean, median, or mode. Then describe the skew of each unimodal distribution and compare the three measures using $<$, $>$, or $\approx$.

Mean, Median, Mode, and Distribution Shape



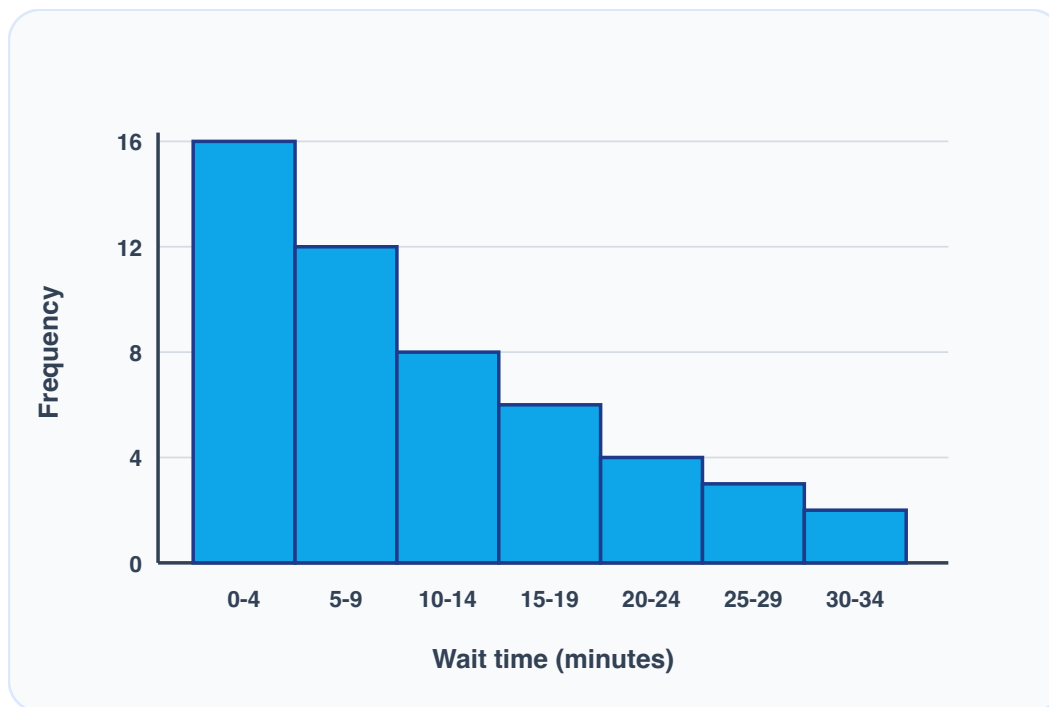
Answer: Negatively skewed (skewed left): from left to right, mean, median, mode; $\text{mean} < \text{median} < \text{mode}$. Symmetric (no skew): all three markers coincide; $\text{mean} \approx \text{median} \approx \text{mode}$. Positively skewed (skewed right): from left to right, mode, median, mean; $\text{mode} < \text{median} < \text{mean}$.

Example 5. Compare the data sets 8, 9, 10, 11 and 8, 9, 10, 11, 52. Find the mean and median of each, then explain what the value 52 demonstrates.

Answer: Without 52: mean = 9.5 and median = 9.5. With 52: mean = 18 and median = 10. The extreme value raises the mean by 8.5 but raises the median by only 0.5, illustrating that the median is resistant and the mean is not.

Example 6. The histogram shows amusement-park wait times. Describe the shape and choose the mean or median as the better measure of a typical wait.

Amusement-Park Wait Times



Answer: The distribution is skewed right. Most waits are short, with a tail toward longer waits, so the median is the better measure of a typical wait.

Example 7. The pulse rates 68, 72, 74, 76, 80, and 86 beats per minute are treated as a population. Find the population mean. Then find the means of Sample A: 68, 72, 74 and Sample B: 76, 80, 86. Does each sample mean under- or overestimate the population mean?

Answer: Population mean: $\mu = 76$ beats per minute. Sample A: $\bar{x} \approx 71.3$, so it underestimates the population mean. Sample B: $\bar{x} \approx 80.7$, so it overestimates the population mean. Different samples need not have the same mean as the population.

Example 8. A class has 30 exam scores with a mean of 82. The mean of the 29 readable scores is 84. Find the unreadable score.

Answer: All 30 scores total $30(82) = 2460$. The 29 readable scores total $29(84) = 2436$. The unreadable score is $2460 - 2436 = 24$.

Example 9. A tornado damage scale uses the category codes 0, 1, 2, 3, 4, and 5. Does it make sense to calculate the mean code? Which measure of center can still be used?

Answer: No. The codes identify ordered damage categories, but arithmetic differences between the codes do not represent measured amounts. The mode can be used to report the most common damage category.

Example 10. Enter 78, 81, 84, 84, 85, and 90 in L1 and run 1-Var Stats. Use the output to identify the mean, median, and sample size. Then identify the mode.

Answer: The output gives $\bar{x} \approx 83.67$, Med=84, and $n = 6$. The mode is 84 because it occurs twice; 1-Var Stats does not report the mode.

$\bar{x}$	83.6667	n	6
minX	78	Med	84
maxX	90	Mode	Not displayed

Section 3.2 Measures of Dispersion

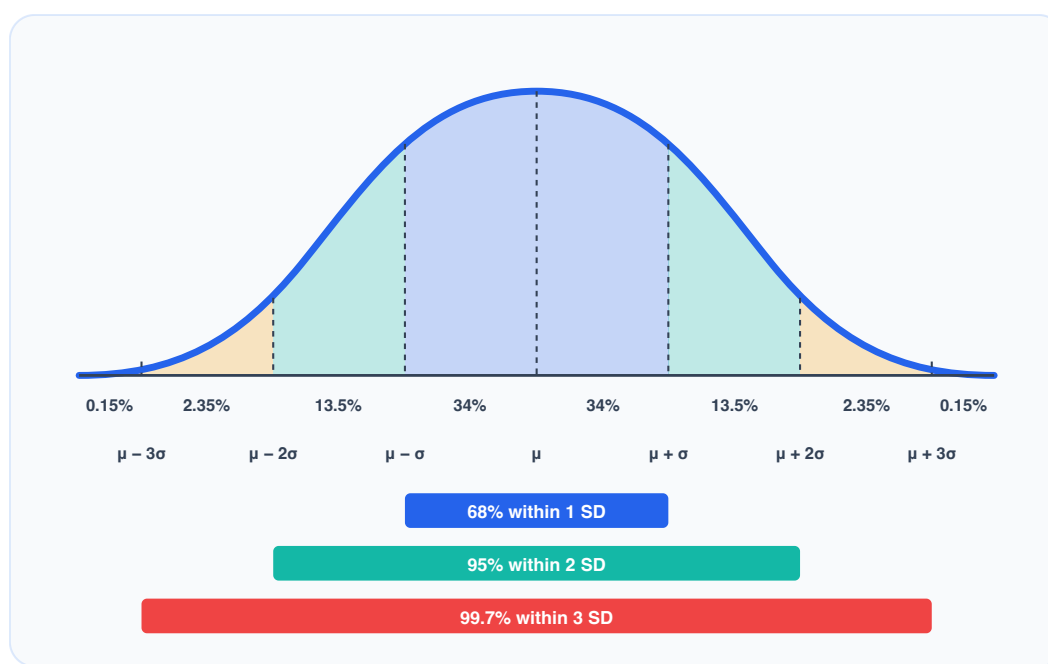
Use range, variance, and standard deviation to describe spread.

QUICK REFERENCE

- Measures of dispersion describe how spread out the data are.
 - The range is maximum - minimum. It uses only the two extreme values and is not resistant to outliers.
 - Standard deviation measures a typical distance of the observations from the mean. It is reported in the same units as the data.
 - Variance is the square of the standard deviation, so its units are squared.
- Standard deviation is never negative.
 - A standard deviation of 0 means every observation is equal.
 - A larger standard deviation means more spread when comparing the same variable measured in the same units.
 - A value more than 2 standard deviations from the mean is often considered far from the mean.
- Use population measures when the data contain the entire population and sample measures when the data are a sample.
 - Population standard deviation is $\sigma = \sqrt{\frac{\sum(x-\mu)^2}{N}}$, and population variance is σ^2 .
 - Sample standard deviation is $s = \sqrt{\frac{\sum(x-\bar{x})^2}{n-1}}$, and sample variance is s^2 .
 - The formulas show how spread is calculated, but use calculator output for most data sets. A sample standard deviation can be smaller or larger than the population standard deviation.
- Range and standard deviation are not resistant measures.
 - An extreme value can increase the range immediately and can pull the mean away from the other values, increasing the standard deviation.
 - Compare both the numerical measures and the distribution's shape when describing spread.

- For a roughly bell-shaped distribution, the Empirical Rule estimates how much data lie near the mean.
 - About 68% lie within 1 standard deviation of the mean, about 95% within 2, and about 99.7% within 3.
 - On either side of the mean, the regions are about 34%, 13.5%, 2.35%, and 0.15% as you move outward.
 - Check that the distribution is roughly bell-shaped before applying the rule.
- TI-84 1-Var Stats workflow:
 - Enter the data in L1, then choose STAT → CALC → 1:1-Var Stats and calculate using L1.
 - Read S_x for the sample standard deviation and σ_x for the population standard deviation.
 - Square the appropriate standard deviation to obtain the variance.

Empirical Rule for a Bell-Shaped Distribution



Example 1. The delivery times 6, 8, 10, and 12 minutes are a sample of a restaurant's deliveries. Find the range, sample variance, and sample standard deviation.

Answer: The range is $12 - 6 = 6$ minutes. From 1-Var Stats, the sample standard deviation is S_x approximately 2.58 minutes and the sample variance is approximately 6.67 square minutes.

Example 2. The same delivery times, 6, 8, 10, and 12 minutes, include every delivery made during a short shift. Find the population variance and population standard deviation.

Answer: Because the data are the entire population, use σ_x . The population standard deviation is $\sqrt{5}$, or approximately 2.24 minutes, and the population variance is 5 square minutes.

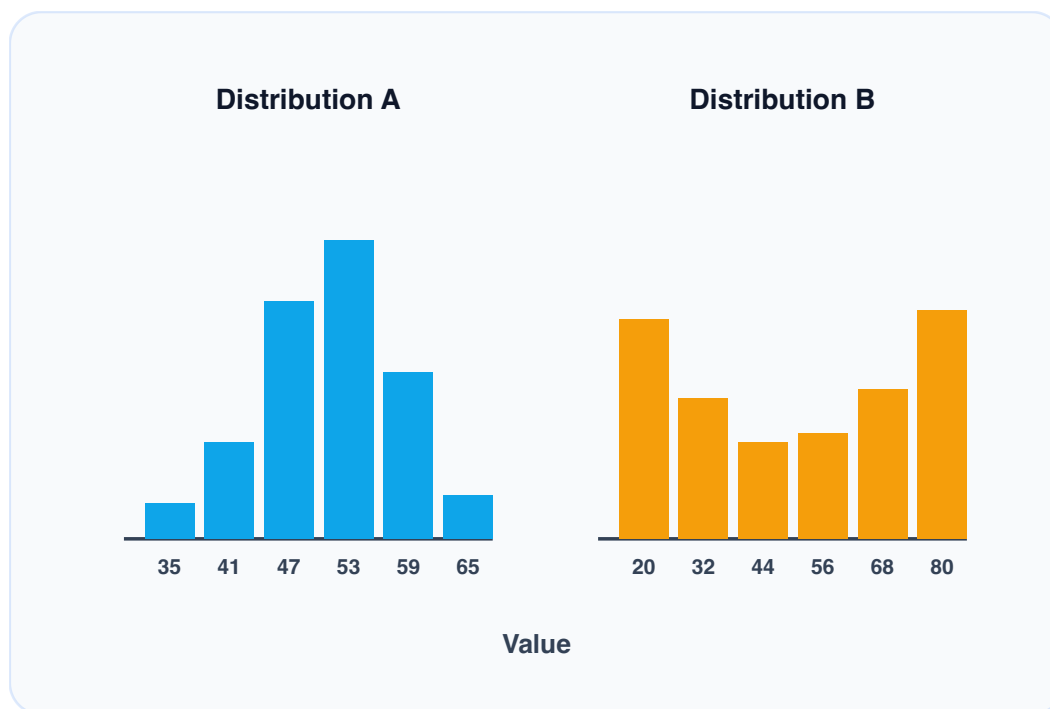
Example 3. A TI-84 1-Var Stats screen reports the output below for quiz scores from one class. The class is treated as a sample of future students. Report the mean, standard deviation, and variance.

x-bar	78.4	n	25
Sx	6.2	σ_x	6.075

Answer: The mean is 78.4 points. Use $Sx=6.2$ because the class is a sample. The sample variance is $6.2^2 = 38.44$ square points.

Example 4. Both distributions have a mean near 50. Which distribution has the larger standard deviation? Explain using the histograms.

Same Center, Different Spread



Answer: Distribution B has the larger standard deviation because its values are spread farther from the mean. Distribution A is more concentrated near 50.

Example 5. A placement-test distribution has mean 72 and standard deviation 6. A student scores 87. How many standard deviations above the mean is the score, and would it be considered far from the mean?

Answer: The score is $(87 - 72)/6 = 2.5$ standard deviations above the mean. Because it is more than 2 standard deviations from the mean, it would be considered far from the mean.

Example 6. The commute times 18, 20, 21, 22, and 24 minutes have range 6 minutes and sample standard deviation about 2.24 minutes. Replace 24 with 60. What happens to the range and standard deviation? Are either measure resistant?

Answer: The range increases from 6 to 42 minutes, and the sample standard deviation increases to about 17.84 minutes. Neither measure is resistant because the extreme value changes both substantially.

Example 7. A population has standard deviation 8. Two random samples from it have standard deviations 6.9 and 9.1. Which sample underestimates the population spread, and which overestimates it?

Answer: The sample with $s = 6.9$ underestimates the population standard deviation. The sample with $s = 9.1$ overestimates it.

Example 8. Decide whether each statement is possible. (a) A data set has standard deviation -3. (b) A data set has standard deviation 0. (c) Find the sample standard deviation of a sample containing only the value 14.

Answer: (a) Impossible; standard deviation cannot be negative. (b) Possible only when all observations are equal. (c) The sample standard deviation is undefined because the formula would divide by $n - 1 = 0$.

Example 9. Adult commute times in a city are roughly bell-shaped with mean 30 minutes and standard deviation 5 minutes. Use the Empirical Rule to estimate the percentages with commute times from 25 to 35 minutes, from 20 to 40 minutes, and from 15 to 45 minutes.

Answer: About 68% are from 25 to 35 minutes, about 95% are from 20 to 40 minutes, and about 99.7% are from 15 to 45 minutes.

Example 10. For the bell-shaped commute-time distribution with mean 30 minutes and standard deviation 5 minutes, estimate each percentage. (a) More than 35 minutes. (b) Less than 20 minutes. (c) From 20 to 35 minutes.

Answer: (a) About 16% because half of the 32% outside 1 standard deviation lies above 35. (b) About 2.5% because half of the 5% outside 2 standard deviations lies below 20. (c) About 81.5% because $13.5\% + 34\% + 34\% = 81.5\%$.

Example 11. A sample of 200 battery lives has mean 10 hours and standard deviation 1 hour. The histogram is strongly skewed right. Should about 190 battery lives be expected between 8 and 12 hours by the Empirical Rule?

Answer: No. Although 8 to 12 hours is within 2 standard deviations of the mean, the Empirical Rule requires a roughly bell-shaped distribution. The strong right skew makes that estimate inappropriate.

Example 12. Treat the 20 values shown as a population. Enter them in L1 and use TI-84 1-Var Stats to find the population mean and population standard deviation. Round each to the nearest whole number. Then find the actual number and percentage of values within 1, 2, and 3 population standard deviations of the mean, and compare the percentages with the Empirical Rule.

Observed values									
30	35	38	40	42	44	46	47	48	49
51	52	53	54	56	58	60	62	65	70

Answer: The population mean is 50, and $\sigma_x \approx 9.995$, so use $\mu = 50$ and $\sigma = 10$ after rounding. Within 1 standard deviation, 40 through 60 contains 14 of 20 values, or 70%, compared with about 68%. Within 2 standard deviations, 30 through 70 contains all 20 values, or 100%, compared with about 95%. Within 3 standard deviations, 20 through 80 also contains all 20 values, or 100%, compared with about 99.7%. The observed percentages are close to the Empirical Rule estimates but do not have to match them exactly.

Section 3.3

Measures of Central Tendency and Dispersion from Grouped Data

Use frequency tables and grouped data to estimate summary measures.

QUICK REFERENCE

- A frequency table summarizes repeated observations.
 - For a table of individual values, treat each value as if it appeared the stated number of times.
 - The total number of observations is the sum of the frequencies: $n = \sum f$ for a sample or $N = \sum f$ for a population.
 - The mean is $\bar{x} = \frac{\sum xf}{\sum f}$ for a sample or $\mu = \frac{\sum xf}{\sum f}$ for a population.
- For data grouped into class intervals, use each class midpoint as a representative value.
 - Class midpoint = (lower class limit + upper class limit) / 2.
 - Classes do not need to have equal widths. Find the midpoint of each class separately.
 - Because the original values within each class are unknown, a grouped-data mean, variance, or standard deviation is an approximation. Use raw data when they are available.
- A weighted mean gives observations different amounts of influence.
 - Use $\bar{x}_w = \frac{\sum xw}{\sum w}$, where x is a value and w is its weight.
 - Common weights include course credit hours, category percentages, quantities purchased, and frequencies.
 - If weights are percentages that total 100%, multiply each value by its decimal weight and add the products.
- Use the grouped summary in context.
 - Choose S_x and s^2 for sample data; choose σ_x and σ^2 for an entire population.
 - Interpret the mean and standard deviation in the original units, and label grouped results as approximate.
 - When comparing groups measured in the same units, compare both center and spread rather than relying on only one statistic.
- TI-84 grouped-data workflow:
 - Enter the values or class midpoints in L1 and their frequencies or weights in L2.
 - Choose STAT → CALC → 1:1-Var Stats, set List:L1 and FreqList:L2, then calculate.
 - Confirm that the reported n equals the sum of the frequencies, then read $\bar{x}$, S_x , or σ_x as required.

Example 1. Use the frequency table to find the mean quiz score and identify the total number of scores.

Score	6	8	10
Frequency	2	3	5

Answer: There are $n = 2 + 3 + 5 = 10$ scores. The mean is $\bar{x} = \frac{6(2)+8(3)+10(5)}{10} = 8.6$ points.

Example 2. The table gives parking fines from a sample of 30 citations. Find each class midpoint, then use 1-Var Stats to approximate the sample mean and sample standard deviation.

Fine	\$0-\$49	\$50-\$99	\$100-\$149	\$150-\$249
Frequency	6	12	7	5
Midpoint	24.5	74.5	124.5	199.5

Answer: The midpoints are 24.5, 74.5, 124.5, and 199.5 dollars. Enter them in L1 and the frequencies in L2. The output gives $\bar{x} = 97$ dollars and S_x approximately 57.37 dollars. These are estimates because each citation is represented by its class midpoint.

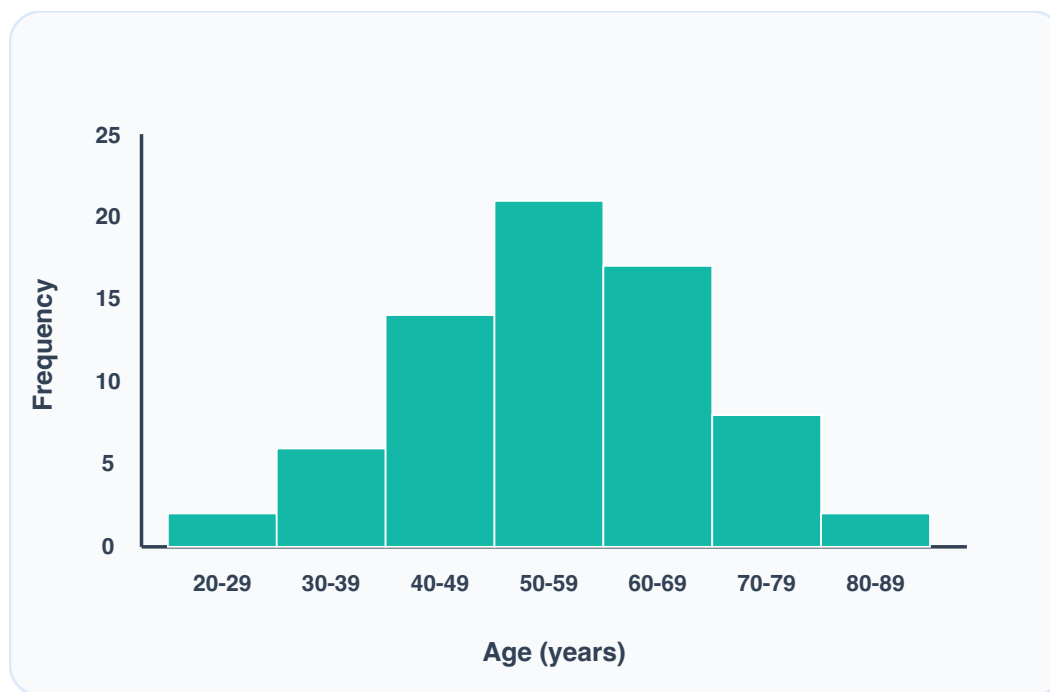
Example 3. The grouped distribution contains every employee at a company. A TI-84 output reports $\bar{x} = 42.625$, $S_x = 13.650$, $\sigma_x = 13.564$, and $n = 80$. Report the approximate population mean, population standard deviation, and population variance.

Answer: The approximate population mean is 42.625 years. Because all employees are included, use σ_x : $\sigma \approx 13.564$ years and $\sigma^2 \approx 184.0$ square years.

Example 4. The grouped distribution contains every employee at a company. Use the class midpoints and frequencies with TI-84 weighted 1-Var Stats to find the approximate population mean and population standard deviation. Then explain why the Empirical Rule is reasonable and estimate the interval containing about 95% of the ages.

Age	20-29	30-39	40-49	50-59	60-69	70-79	80-89
Frequency	2	6	14	21	17	8	2

Grouped Employee Ages



Answer: Enter the class midpoints 24.5, 34.5, 44.5, 54.5, 64.5, 74.5, and 84.5 in L1 and the frequencies in L2, then run 1-Var Stats L1,L2. The output gives an approximate population mean of $\mu \approx 55.5$ years and, using σ_x , an approximate population standard deviation of $\sigma \approx 13.2$ years. The histogram has one peak and is approximately bell-shaped, even though opposite bars are not exact matches, so the Empirical Rule is reasonable. About 95% of the ages should lie from $55.5 - 2(13.2) = 29.1$ to $55.5 + 2(13.2) = 81.9$ years.

Example 5. A student earns an A in a 3-credit course, a B in a 4-credit course, a C in a 2-credit course, and a B in a 3-credit course. Using A = 4, B = 3, C = 2, D = 1, and F = 0, find the GPA.

Answer: Use credit hours as weights: $\frac{4(3)+3(4)+2(2)+3(3)}{3+4+2+3} = \frac{37}{12} \approx 3.08$. The GPA is 3.08.

Example 6. A course grade uses homework 15%, quizzes 20%, tests 25%, and the final exam 40%. A student earns 95, 86, 78, and 82 in those categories, respectively. Find the course average.

Answer: The weighted mean is $95(0.15) + 86(0.20) + 78(0.25) + 82(0.40) = 83.75$. The course average is 83.75%.

Example 7. A snack mix contains 2 pounds of nuts at \$4.50 per pound, 3 pounds of cereal at \$3.20 per pound, and 5 pounds of pretzels at \$2.10 per pound. Find the cost per pound of the mixture.

Answer: The quantities are the weights. The total cost is \$29.10 and the total weight is 10 pounds, so the weighted cost is $\$29.10 / 10 = \2.91 per pound.

Example 8. Two grouped populations of customer ages have the approximate summaries shown. Compare their centers and spreads.

Population	Approximate mean	Approximate population SD
A	40.36 years	20.87 years
B	41.02 years	19.27 years

Answer: Population B has the slightly higher mean age, 41.02 years compared with 40.36 years. Population A has the larger standard deviation, 20.87 years compared with 19.27 years, so its ages are more spread out.

Section 3.4 Measures of Position and Outliers

Use z-scores, percentiles, quartiles, and outlier fences.

QUICK REFERENCE

- A z-score gives a value's distance from the mean in standard-deviation units.
 - Use $z = \frac{x - \mu}{\sigma}$ for a population or $z = \frac{x - \bar{x}}{s}$ for a sample.
 - A positive z-score is above the mean, a negative z-score is below the mean, and $z = 0$ is at the mean. A z-score has no units.
 - To recover a raw value, use $x = \mu + z\sigma$ or $x = \bar{x} + zs$.
- Z-scores allow fair relative comparisons between different distributions.
 - When larger values are better, the larger z-score is the stronger relative result.
 - When smaller values are better, such as race times, the more negative z-score is the stronger relative result.
 - After standardizing every value in a population, the z-scores have mean 0 and population standard deviation 1.
- The kth percentile is a value with about k% of the observations at or below it.
 - The remaining observations are above it. For example, the 80th percentile has about 80% at or below and 20% above.
 - On an ogive, the vertical cumulative relative frequency can be read as a percentile. Read upward from a value to find its percentile or across from a percentile to estimate its value.

- Quartiles divide ordered data into four parts.
 - Q_1 is the 25th percentile, Q_2 is the median or 50th percentile, and Q_3 is the 75th percentile.
 - About 25% of observations are at or below Q_1 , about 50% are at or below Q_2 , and about 75% are at or below Q_3 .
 - The interquartile range is $IQR = Q_3 - Q_1$. It measures the spread of the middle 50% and is resistant to extreme values.
- Use the 1.5(IQR) rule to identify potential outliers.
 - Lower fence: $Q_1 - 1.5(IQR)$. Upper fence: $Q_3 + 1.5(IQR)$.
 - A value below the lower fence or above the upper fence is an outlier. A value equal to a fence is not outside it.
 - An outlier may be a valid unusual observation or may come from measurement, entry, or sampling error. Investigate it rather than automatically deleting it.
- Choose resistant summaries when the distribution is skewed or has outliers.
 - Median and IQR are resistant; mean, range, and standard deviation are not.
 - For roughly symmetric data without strong outliers, mean and standard deviation are commonly reported. For skewed data or data with outliers, median and IQR usually give a more stable summary.
- TI-84 quartile workflow:
 - Enter the data in L1 and run 1-Var Stats with List:L1.
 - Scroll down to read minX, Q_1 , Med, Q_3 , and maxX.
 - The TI-84 reports quartiles but does not label outliers; calculate the fences and compare the original observations with them.

Example 1. Adult men in a city have mean height 69.6 inches with standard deviation 3.2 inches. Adult women have mean height 64.1 inches with standard deviation 3.9 inches. Who is relatively taller: a 75-inch man or a 71-inch woman?

Answer: The man's z-score is $(75 - 69.6)/3.2 \approx 1.69$. The woman's is $(71 - 64.1)/3.9 \approx 1.77$. The woman is relatively taller because her height is more standard deviations above her group's mean.

Example 2. Runner A finishes in 52 minutes in a race with mean 60 minutes and standard deviation 4 minutes. Runner B finishes in 44 minutes in a race with mean 50 minutes and standard deviation 2.5 minutes. Which is the more convincing victory relative to that race?

Answer: Runner A has $z = (52 - 60)/4 = -2$. Runner B has $z = (44 - 50)/2.5 = -2.4$. Runner B's more negative z-score represents the stronger relative finish because a lower time is better.

Example 3. A school admits applicants who score at least 2 standard deviations above a test mean of 200 with standard deviation 20. What is the minimum qualifying score?

Answer: Use $x = \mu + z\sigma$: $x = 200 + 2(20) = 240$. The minimum qualifying score is 240.

Example 4. A bolt should be destroyed if its length is more than 4 standard deviations from the mean. Bolt lengths have mean 9.0 cm and standard deviation 0.05 cm. Which lengths should be destroyed?

Answer: The cutoffs are $9.0 - 4(0.05) = 8.8$ cm and $9.0 + 4(0.05) = 9.2$ cm. Destroy bolts shorter than 8.8 cm or longer than 9.2 cm.

Example 5. Use the ogive. Interpret the percentile rank of a score of 150, and identify the score at the 25th percentile.

Percentile Ranks of Test Scores



Answer: A score of 150 is at the 60th percentile, so about 60% of scores are at or below 150 and about 40% are above it. The 25th percentile corresponds to a score of 130.

Example 6. For the ordered data 12, 14, 15, 17, 19, 20, 21, 24, 27, and 35, find and interpret the quartiles and IQR.

Answer: $Q_1 = 15$, $Q_2 = 19.5$, and $Q_3 = 24$. About 25%, 50%, and 75% of the values are at or below those respective quartiles. The middle 50% span $IQR = 24 - 15 = 9$ units.

Example 7. A data set has $Q_1 = 15$ and $Q_3 = 24$. Find the fences and determine whether 1.5, 35, and 49 are outliers.

Answer: The IQR is 9. The lower fence is $15 - 1.5(9) = 1.5$, and the upper fence is $24 + 1.5(9) = 37.5$. The value 49 is an outlier. The values 1.5 and 35 are not outside the fences.

Example 8. The table gives TI-84 quartile output for delivery times in two regions. Which region has the more dispersed middle 50%?

Region	Q1	Median	Q3
A	18	25	38
B	20	26	31

Answer: Region A has $IQR = 38 - 18 = 20$ minutes. Region B has $IQR = 31 - 20 = 11$ minutes. Region A's middle 50% is more dispersed.

Example 9. The data 20, 22, 24, 25, 26, 28, 30, 32, 34, and 100 contain an extreme value. If 100 is changed to 200, what happens to the standard deviation and IQR? What property does this illustrate?

Answer: The standard deviation increases because the extreme value moves farther from the mean. The IQR does not change because the middle half is unchanged. This illustrates that IQR is resistant while standard deviation is not.

Section 3.5 The Five-Number Summary and Boxplots

Summarize data with quartiles, IQR, outliers, and boxplots.

QUICK REFERENCE

- The five-number summary lists minimum, Q_1 , median, Q_3 , and maximum in ascending order.
 - Use ordered data or the lower portion of TI-84 1-Var Stats output to find the five values.
 - The five-number summary describes position but does not show every observation, the mean, or the standard deviation.
- A modified boxplot displays quartiles, non-outlier extremes, and outliers on a number line.
 - The box runs from Q_1 to Q_3 , so its length represents the IQR. The line inside the box marks the median.
 - Whiskers extend to the smallest and largest observations that are not outliers; they do not extend to the fences unless an observation occurs there.
 - Plot observations outside the 1.5(IQR) fences separately.
- To construct a modified boxplot:
 - Find Q_1 , the median, Q_3 , the IQR, and both outlier fences.
 - Draw a scaled number line, make the box from Q_1 to Q_3 , and mark the median inside it.
 - Draw whiskers to the most extreme non-outliers and plot each outlier separately.
- A boxplot gives clues about distribution shape.
 - Roughly symmetric: the median is near the center of the box and the whiskers have similar lengths.
 - Skewed right: the median is usually left of the box's center and the right side or right whisker is longer.
 - Skewed left: the median is usually right of the box's center and the left side or left whisker is longer.Treat these as visual clues rather than absolute rules.
- Use side-by-side boxplots to compare groups on the same scale.
 - Compare medians for center, IQRs and overall ranges for spread, and separately plotted points for outliers.
 - A group with a smaller IQR or range is more consistent. A longer box and whiskers usually suggest greater spread, but a boxplot does not give an exact standard deviation.
 - A boxplot does not show clusters, gaps, or sample size, so use the available evidence without claiming more than the display supports.

- TI-84 modified-boxplot workflow:
 - Enter the data in a list. Open 2nd → Y= (STAT PLOT), turn one plot on, and choose the modified boxplot icon with separate outlier marks.
 - Set Xlist to the data list and Freq to 1, then use ZOOM → 9:ZoomStat.
 - Press TRACE and move left or right to read the quartiles, whisker endpoints, and outliers.

Anatomy of a Modified Boxplot

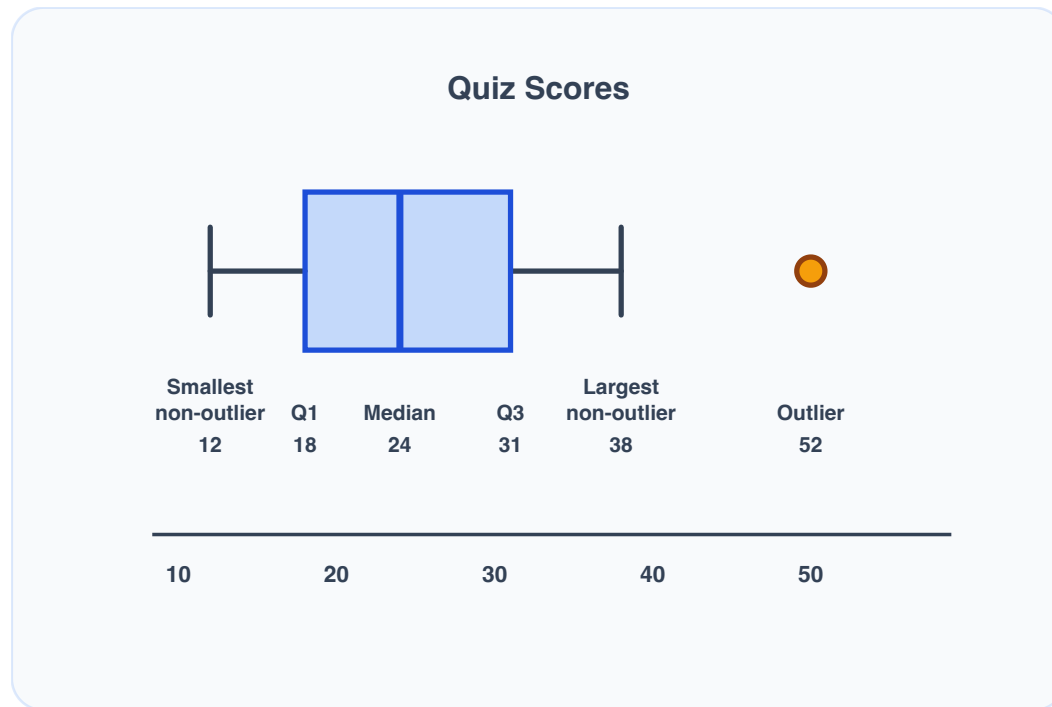


Example 1. Find the five-number summary for 2, 4, 5, 8, 10, 12, 15.

Answer: Minimum 2, $Q_1 = 4$, median 8, $Q_3 = 12$, maximum 15.

Example 2. Use the modified boxplot to identify and interpret the smallest non-outlier, Q_1 , median, Q_3 , largest non-outlier, and outlier. Find the IQR and state what the box represents.

Modified Boxplot

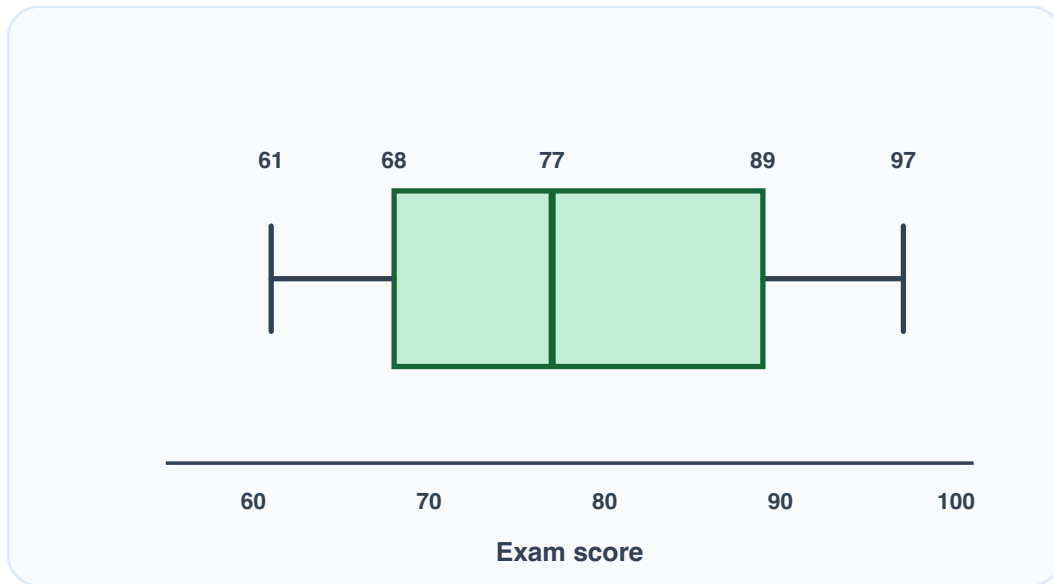


Answer: The smallest non-outlier is 12, and the largest non-outlier is 38, so the whiskers end at those scores. $Q_1 = 18$, so about 25% of scores are at or below 18. The median is 24, so about 50% are at or below 24 and about 50% are at or above 24. $Q_3 = 31$, so about 75% are at or below 31. The IQR is $31 - 18 = 13$, and the box represents the middle 50% of scores. The separate point at 52 is an outlier.

Example 3. An exam has five-number summary 61, 68, 77, 89, 97 and no outliers. Construct its boxplot on a scaled number line.

Answer: The completed boxplot has a box from 68 to 89, a median line at 77, and whiskers extending to 61 and 97.

Exam Score Boxplot



Example 4. Use the side-by-side boxplots to compare typical delivery time, consistency, shape, and outliers for Routes A and B.

Delivery Times by Route



Answer: Route A has the lower median delivery time, 30 minutes compared with 42. Route A is also more consistent because its IQR is 12 minutes compared with 20. Route A is roughly symmetric. Route B is skewed right and has an outlier at 88 minutes.

Example 5. Boxplot I spans approximately 5 to 18 with an IQR of 6. Boxplot II spans approximately 1 to 29 with an IQR of 16. Which data set likely has the larger standard deviation? Can the exact standard deviations be found from the boxplots?

Answer: Boxplot II likely has the larger standard deviation because its values are more spread out. The exact standard deviations cannot be calculated from boxplots alone.

Example 6. A student enters tablet weights in L1. Give the TI-84 steps to display a modified boxplot and inspect its values.

Answer: Open STAT PLOT, turn Plot1 on, select the modified boxplot, set Xlist:L1 and Freq:1, then choose ZoomStat. Use TRACE and the arrow keys to inspect quartiles, whisker endpoints, and any outliers.

Module 4: Correlation and Regression

Students describe relationships between two quantitative variables and use linear regression carefully.

Section 4.1 Scatter Diagrams and Correlation

Describe scatterplots and interpret the correlation coefficient.

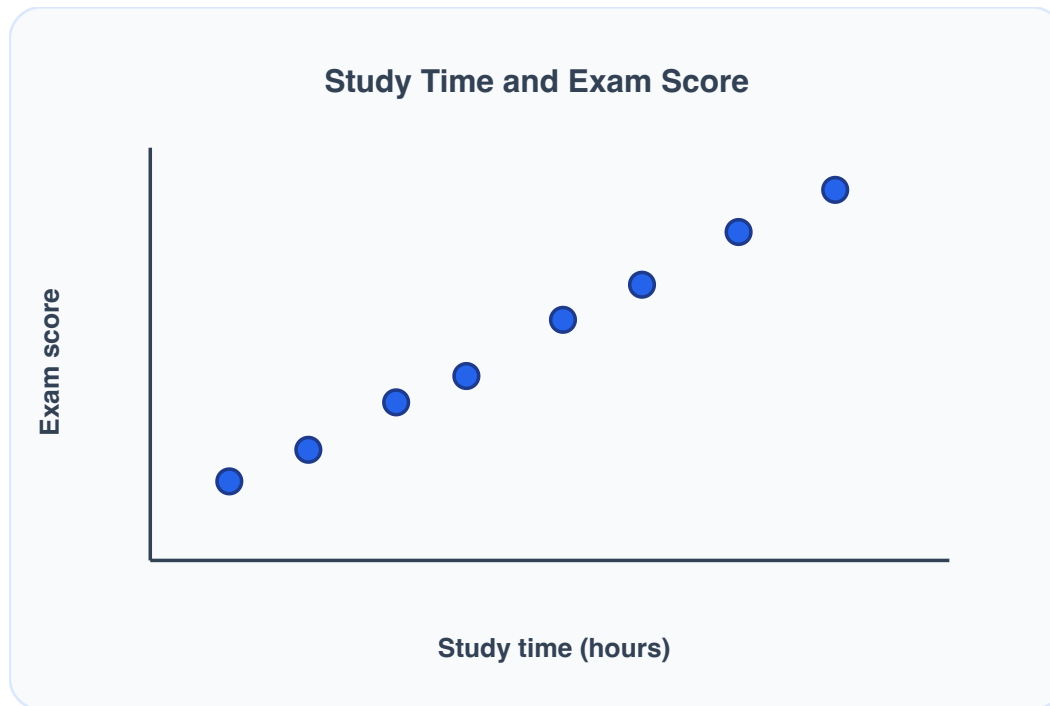
QUICK REFERENCE

- A scatterplot displays paired quantitative data measured on the same individuals or objects.
 - Plot the explanatory or predictor variable on the horizontal axis and the response variable on the vertical axis.
 - The explanatory variable may help explain or predict the response, but the roles are not always determined by the data alone. Use the study's question and context.
- Describe a scatterplot by direction, form, strength, and unusual features.
 - Direction may be positive, negative, or neither. Form may be linear, curved, or show no clear pattern.
 - Strength describes how closely the points follow the form. Also note clusters, gaps, and points that do not follow the overall pattern.
 - Always inspect the scatterplot. A numerical summary alone can hide curvature or an influential unusual point.
- The linear correlation coefficient measures the direction and strength of a linear relationship.
 - Use r for a sample and ρ for a population. The long formula is $r = \frac{1}{n-1} \sum \left(\frac{x-\bar{x}}{s_x} \right) \left(\frac{y-\bar{y}}{s_y} \right)$, but technology is used for computation.
 - The value is between -1 and 1. The sign gives direction; the closer $|r|$ is to 1, the stronger the linear relationship.
 - A value near 0 means little evidence of a linear relationship; a strong curved relationship can still have r near 0.
- Correlation is unitless and is not resistant.
 - Changing inches to centimeters or applying another positive linear unit conversion does not change r .
 - An unusual point can change the direction or strength of r , so investigate it with the scatterplot.
- Use the homework critical-value table to decide whether the sample shows a linear relation.
 - Find $|r|$ and the critical value for the sample size n .
 - If $|r|$ is greater than the critical value, conclude that a linear relation exists; use the sign of r to call it positive or negative.
 - If $|r|$ is not greater than the critical value, conclude that the data do not provide evidence of a linear relation.
- Correlation does not by itself establish causation.
 - With observational data, say the variables are associated, not that one causes the other.
 - A lurking variable may be related to both variables and may help explain the observed association.

- TI-84 scatterplot and correlation workflow:
 - Enter explanatory values in L1 and matching response values in L2. Keep each pair in the same row.
 - Use STAT PLOT to turn on a scatterplot with Xlist:L1 and Ylist:L2, then use ZoomStat.
 - Run LinReg(ax+b) to read r . In this form, a is the slope and b is the intercept. If r is hidden, run DiagnosticOn from the calculator's catalog first.

Example 1. Use the scatterplot. Describe the direction, form, and strength of the association.

Scatterplot: Study Time and Score



Answer: The association is positive, roughly linear, and strong.

Example 2. Which correlation is stronger: $r = -0.91$ or $r = 0.45$?

Answer: $r = -0.91$, because it is closer to -1 or 1 .

Example 3. Can $r = 1.24$ be a correlation coefficient?

Answer: No. r must be between -1 and 1 .

Example 4. A study finds a strong correlation between ice cream sales and drowning incidents. What causal wording should be avoided?

Answer: Do not say ice cream sales cause drowning incidents. The data show an association, and warm weather or season is a plausible lurking variable related to both.

Example 5. A researcher wants to use commute time to predict a well-being score. Identify the explanatory and response variables. Then use the paired table and TI-84 output to describe the relationship.

Commute time (min)	5	20	35	60	85
Well-being score	69.1	67.4	65.6	63.0	60.5

Answer: Commute time is the explanatory variable and well-being score is the response. The output $r = -0.999$ and the scatterplot show a very strong negative linear relationship: longer commutes are associated with lower well-being scores.

Example 6. For $n = 8$ paired observations, a TI-84 gives $r = 0.682$. The homework critical value for $n = 8$ is 0.707. Is there evidence of a linear relation?

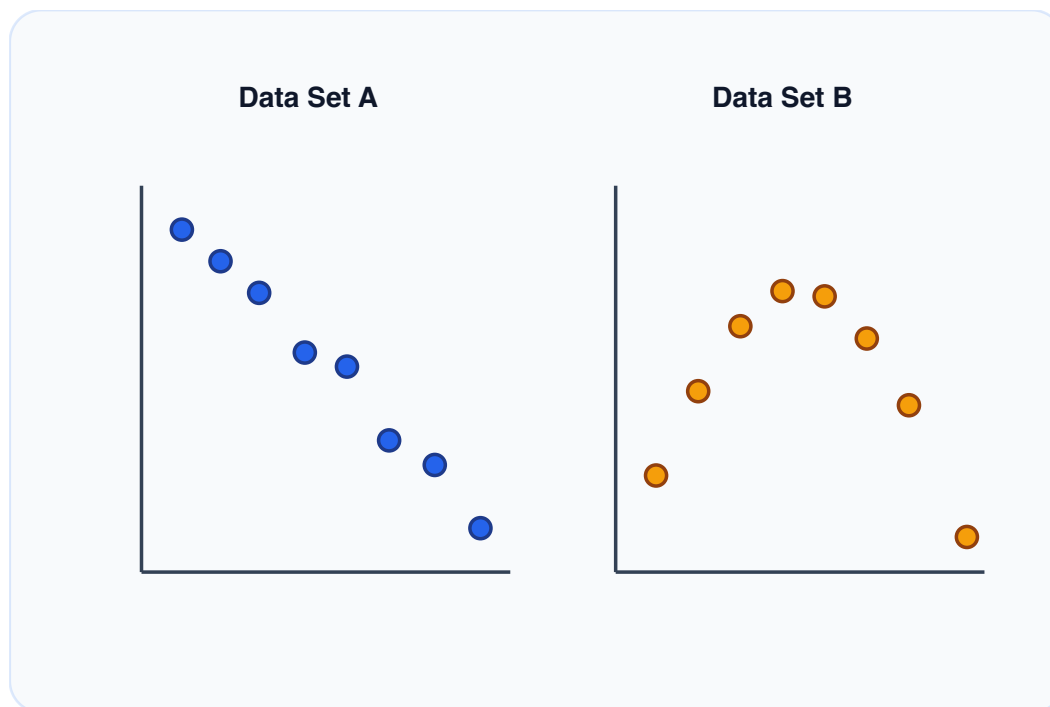
Answer: No. Although the sample correlation is positive, $|0.682| = 0.682$ is not greater than 0.707, so the data do not provide evidence of a linear relation by the critical-value rule.

Example 7. Heights and head circumferences measured in inches have $r = 0.955$. What happens to r if both variables are converted to centimeters?

Answer: The correlation remains $r = 0.955$. Multiplying both variables by the positive conversion factor 2.54 changes the units but not the correlation.

Example 8. Both data sets have nearly the same correlation, $r \approx -0.82$. Why is the correlation coefficient alone not an adequate summary?

Similar Correlations, Different Patterns



Answer: Data Set A follows a reasonably linear downward pattern, while Data Set B contains a curved cluster and an unusual point. The same correlation can hide very different structures, so the scatterplot and r must be used together.

Example 9. Predict the likely correlation direction for each pair. (a) Number of dinner guests and total dinner bill. (b) Day of the week and outdoor temperature. (c) Outdoor temperature and number of people wearing coats.

Answer: (a) Positive correlation. (b) No clear correlation. (c) Negative correlation.

Section 4.2 Least-Squares Regression

Use regression equations for prediction and interpret slope, intercept, residuals, and extrapolation.

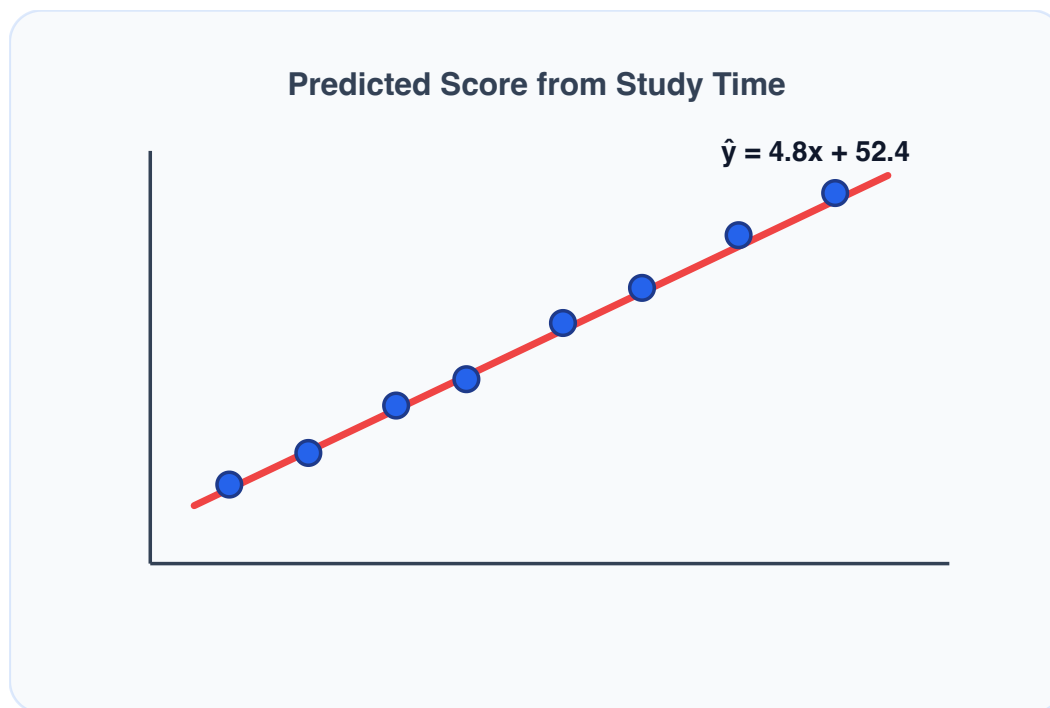
QUICK REFERENCE

- The least-squares regression line predicts a quantitative response from a quantitative explanatory variable.
 - Write the line as $\hat{y} = b_0 + b_1x$, where $\hat{y}$ is the predicted response.
 - The line minimizes the sum of the squared vertical residuals. For formula context, $b_1 = r \frac{s_y}{s_x}$ and $b_0 = \bar{y} - b_1\bar{x}$, but use technology to fit the model.
 - Fit and use a linear model only after the scatterplot and correlation support a linear relationship.
- Interpret the slope and intercept in context.
 - The slope is the predicted change in the response for each one-unit increase in the explanatory variable, on average. Include both variables and their units.
 - The intercept is the predicted response when $x = 0$. Interpret it only when zero is meaningful and reasonably within the scope of the data.
 - A negative slope predicts a decrease in y as x increases; a positive slope predicts an increase.
- Use the regression equation for prediction and comparison.
 - Substitute an explanatory value for x to find the predicted response $\hat{y}$. The same prediction can estimate a mean response at that x or an individual response, though individual values vary around the line.
 - A residual is observed - predicted: $e = y - \hat{y}$. A positive residual means the observation is above the line; a negative residual means it is below the line.
 - Keep more digits in the regression equation during calculations, then round the final prediction as requested.
- Limit predictions to the model's scope.
 - Interpolation predicts within the observed range of x ; extrapolation predicts outside it and may be unreliable because the pattern may change.
 - Do not automatically apply a model to a different population or type of subject, even when the new x -value lies in the numerical range.
- The coefficient of determination measures explained variation.
 - For simple linear regression, $R^2 = r^2$. Convert it to a percent for interpretation.
 - Say: 'R-squared percent of the variation in the response is explained by the linear model with the explanatory variable.' Do not say that percent of the response itself is explained.

- Use a residual plot to check whether a linear model is appropriate.
 - A useful linear model has residuals scattered around zero without a clear pattern and with reasonably similar vertical spread across x .
 - A curved pattern suggests a nonlinear relationship. A fan shape suggests changing residual variation. A point far above or below the others may be an outlier.
- TI-84 regression workflow:
 - Enter explanatory values in L1 and paired responses in L2. Turn on DiagnosticOn when r and r^2 are needed.
 - Run LinReg(ax+b) using L1,L2. This course uses that form consistently, so a is the slope and b is the intercept.
 - The TI-84 also lists LinReg(a+bx), which reverses the coefficient roles: a is the intercept and b is the slope. Check the displayed equation form before interpreting output from another source.
 - To graph the fitted line, store the regression equation to Y1, turn on the scatterplot, and use ZoomStat.

Example 1. Use the regression graph. Interpret the slope in context.

Least-Squares Regression Line



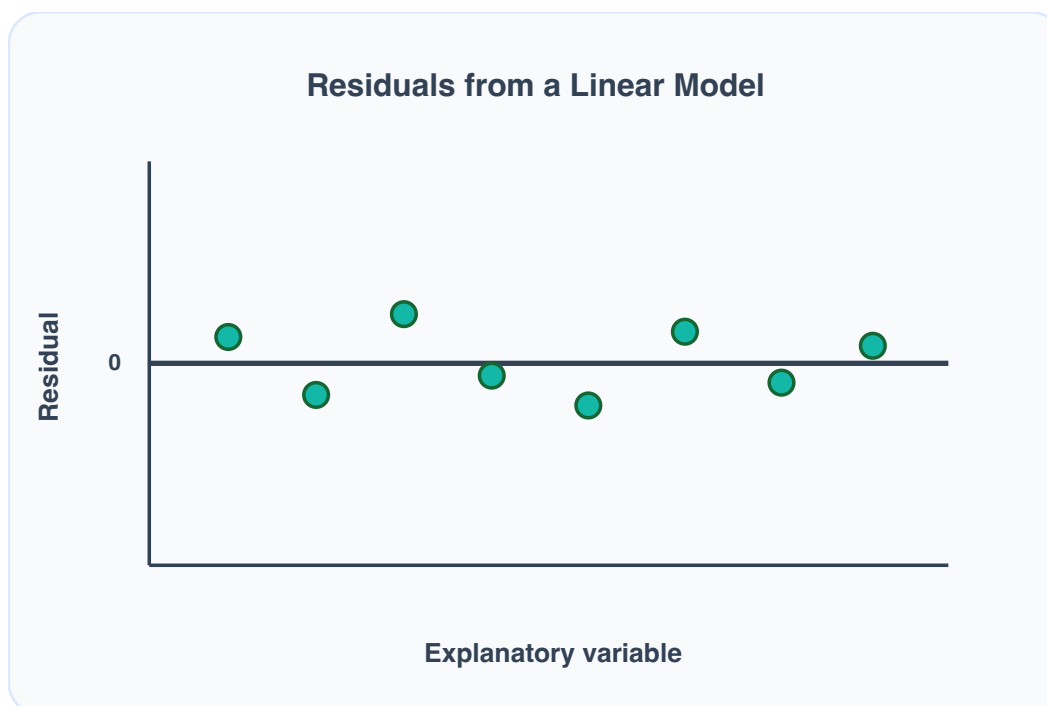
Answer: For each additional hour of study time, the predicted exam score increases by about 4.8 points.

Example 2. The TI-84 LinReg(ax+b) output gives $a = 4.8$, $b = 52.4$, $r = 0.94$, and $r^2 = 0.884$ for predicting exam score from study time in hours. Write the regression equation and interpret each output value.

Answer: The regression equation is $\hat{y} = 4.8x + 52.4$. The slope $a = 4.8$ means the predicted exam score increases by about 4.8 points for each additional hour of study time. The intercept $b = 52.4$ is the predicted exam score for a student who studies 0 hours. The correlation $r = 0.94$ indicates a strong positive linear relationship between study time and exam score. The value $r^2 = 0.884$ means about 88.4% of the variation in exam score is explained by the linear model with study time.

Example 3. Use the residual plot. Does the residual plot show an obvious curved pattern?

Residual Plot



Answer: No. The residuals are scattered around 0, so this plot does not show an obvious curved pattern.

Example 4. Data were collected for x -values from 2 to 15. Is predicting at $x = 40$ interpolation or extrapolation?

Answer: Extrapolation. The value 40 is outside the observed range, so the prediction may be unreliable.

Example 5. A commute-time model is $\hat{y} = 69.051 - 0.095x$, where x is commute time in minutes and y is a well-being score. Interpret the slope, predict the score for a 25-minute commute, and find the residual if the observed score is 66.0.

Answer: For each additional minute of commute time, the predicted well-being score decreases by 0.095 point, on average. At 25 minutes, $\hat{y} = 69.051 - 0.095(25) = 66.676$, or about 66.7. The residual is $66.0 - 66.676 = -0.676$, so the observed score is about 0.7 point below the prediction.

Example 6. For regions with 15% to 60% of adults holding a bachelor's degree, the model is $\hat{y} = 17884 + 622.2x$, where x is the percentage of adults holding a bachelor's degree and $\hat{y}$ is the predicted median income in dollars. Interpret the slope. Should the intercept be interpreted?

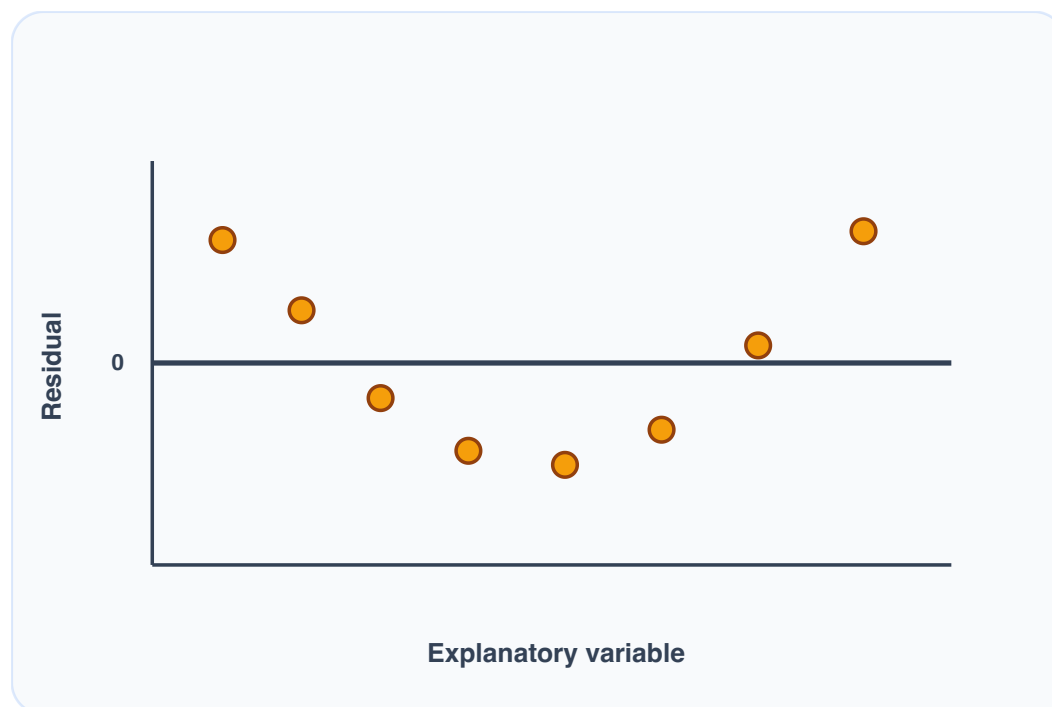
Answer: For each 1-percentage-point increase in adults holding a bachelor's degree, predicted median income increases by about \$622.20, on average. The intercept would predict income at 0%, which is outside the observed range, so it should not be interpreted.

Example 7. A regression model was built from gas-powered cars. A hybrid car's weight falls within the observed weight range. Is the model automatically appropriate for predicting the hybrid's city mileage?

Answer: No. A hybrid is a different type of car, so being inside the numerical weight range is not enough. The model should not be applied to a different population without evidence that the same relationship holds.

Example 8. The scatterplot appears strongly associated and has $r = -0.97$, but the residual plot has a clear curved pattern. Is the linear model appropriate?

Residual Plot with Curvature



Answer: No. The curved residual pattern shows that the linear model systematically overpredicts in some regions and underpredicts in others. A high absolute correlation does not override the residual evidence of nonlinearity.

Module 5: Probability

Students compute probabilities using complements, addition rules, multiplication rules, conditional probability, and counting techniques.

Section 5.1 Probability Rules

Use basic probability language and simple probability rules.

QUICK REFERENCE

- Probability measures the likelihood of an outcome or event.
 - A probability experiment is a repeatable process with uncertain results. An outcome is one possible result, the sample space is the collection of all possible outcomes, and an event contains one or more outcomes.
 - Probabilities are between 0 and 1 inclusive. Probability 0 represents an impossible event, and probability 1 represents a certain event.
 - An event is often called unusual when $P(E) \leq 0.05$, meaning its probability is 0.05 or less, unless the problem sets another cutoff.
- A probability model lists each outcome and its probability.
 - Every probability must be between 0 and 1 inclusive, and all probabilities must sum to 1.
 - When outcomes are equally likely, $P(E) = \frac{\text{number of outcomes in } E}{\text{number of outcomes in the sample space}}$.
- Probability can come from different sources.
 - Theoretical probability uses a known model with equally likely outcomes, such as a fair die.
 - Empirical probability uses observed relative frequency: number of times the event occurred / number of trials.
 - Subjective probability is an informed judgment based on experience or available information rather than equally likely outcomes or repeated data.
- Long-run relative frequency tends to settle near the true probability as the number of trials grows.
 - Short runs can vary substantially, so an empirical result does not need to match a theoretical probability exactly.
 - When a theoretical probability is known, compare it with the empirical probability. A difference may be random variation and does not by itself prove that the model is wrong; more trials provide stronger evidence.
 - A probability can be used to predict an approximate long-run count: expected count = number of trials x probability.
- The complement of event A , written A^c , contains all outcomes where A does not occur.
 - Complement rule: $P(A^c) = 1 - P(A)$.
 - The event and its complement are mutually exclusive and together cover the entire sample space.

Example 1. If $P(A) = 0.37$, find $P(A^c)$.

Answer: $1 - 0.37 = 0.63$.

Example 2. Can a probability equal 1.08?

Answer: No. Probabilities cannot be greater than 1.

Example 3. A fair die is rolled. What is the sample space?

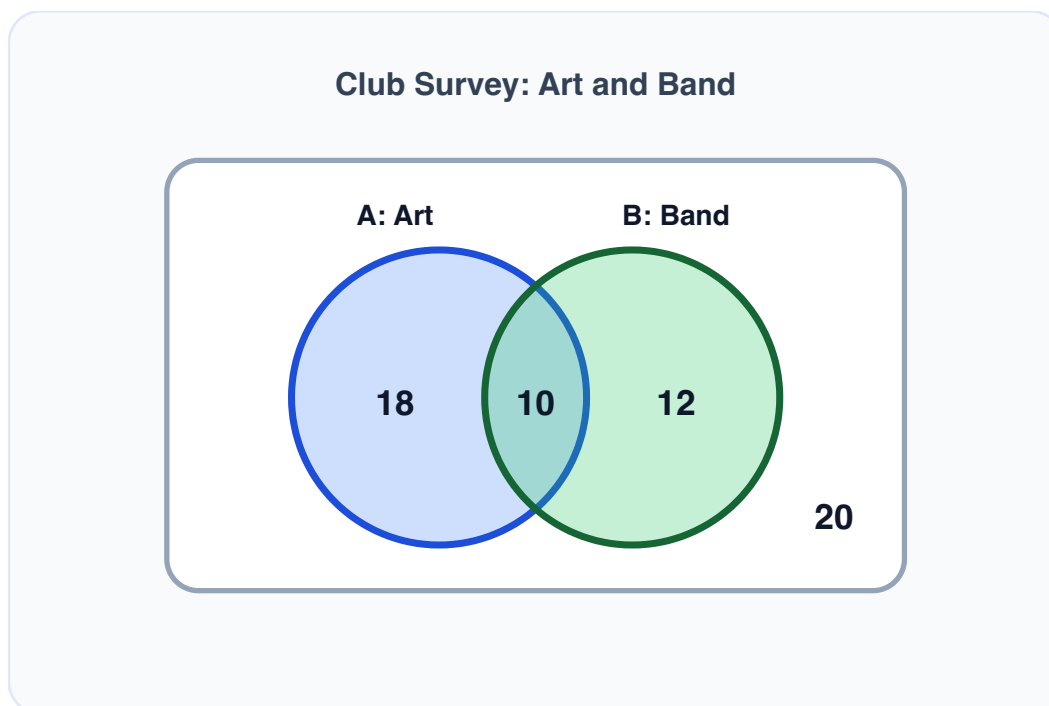
Answer: $\{1, 2, 3, 4, 5, 6\}$.

Example 4. A fair die is rolled. Find the probability of rolling an even number.

Answer: $\frac{3}{6} = \frac{1}{2}$.

Example 5. Use the Venn diagram. How many students are in neither Art nor Band?

Two-Event Venn Diagram



Answer: 20 students are outside both circles, so 20 students are in neither group.

Example 6. A coin is tossed 100 times and lands tails 36 times. (a) Find the empirical probability of tails. (b) If the coin is fair, state the theoretical probability of tails. (c) Compare the results. Does this experiment prove that the coin is unfair?

Answer: (a) The empirical probability is $36 / 100 = 0.36$. (b) For a fair coin, the theoretical probability is $1/2 = 0.50$. (c) The empirical result is below the theoretical probability. This difference could be random variation and does not by itself prove that the coin is unfair. More trials would provide stronger evidence; for a fair coin, the empirical proportion should tend to move closer to 0.50 over many trials.

Example 7. Does the table define a probability model? If not, explain why.

Color	Red	Blue	Green	Yellow
Probability	0.25	0.40	0.30	0.15

Answer: No. Although every probability is between 0 and 1, they sum to 1.10 rather than 1.

Example 8. Classify each probability. (a) A fair die has probability $1/6$ of landing on 4. (b) In 400 rolls, a die landed on 4 a total of 94 times. (c) A meteorologist judges a 30% chance of rain using current conditions and experience.

Answer: (a) Theoretical. (b) Empirical. (c) Subjective.

Example 9. A type of theft has probability 0.010. Is it unusual using the 0.05 guideline? In 2,000 similar records, about how many such thefts would be expected?

Answer: Yes. The unusual-event guideline is $P(E) \leq 0.05$, and $0.010 < 0.05$. The expected long-run count is $2000(0.010) = 20$ thefts.

Example 10. A person draws one playing card and records its suit, then spins a three-space spinner labeled 1, 2, and 3. Use C, D, H, and S for the card suits. List the sample space.

Answer: The 12 outcomes are $S = \{C1, C2, C3, D1, D2, D3, H1, H2, H3, S1, S2, S3\}$. Each outcome records the card suit first and the spinner result second.

Section 5.2 The Addition Rule and Complements

Compute probabilities involving "or" statements.

QUICK REFERENCE

- The union A or B contains outcomes in A , in B , or in both.
 - In probability, or is inclusive unless the wording explicitly says exactly one.
 - The intersection A and B contains outcomes shared by both events.
- Mutually exclusive or disjoint events cannot happen together.
 - For disjoint events, $P(A \text{ and } B) = 0$, so $P(A \text{ or } B) = P(A) + P(B)$.
 - Do not assume events are disjoint merely because their labels differ; decide whether one outcome can belong to both.
- Use the general addition rule whenever events may overlap.
 - $P(A \text{ or } B) = P(A) + P(B) - P(A \text{ and } B)$.
 - Subtract the overlap once because it was included in both individual probabilities. The same idea appears visually in a Venn diagram.
- Use the complement rule for not, none, or neither.
 - $P(A^c) = 1 - P(A)$.
 - For neither A nor B , find the complement of A or B : $1 - P(A \text{ or } B)$.
- Use the structure of a standard 52-card deck when card events appear.
 - The deck has four suits with 13 cards in each suit. Hearts and diamonds are red; clubs and spades are black.
 - Each rank appears four times, once in each suit. There are four aces, four twos, and so on, including four jacks, four queens, and four kings.
 - Events involving different ranks are disjoint for one draw, but a rank and a suit overlap in one card, such as the six of diamonds.

- Read a two-way table by identifying the correct cells and the grand total.
 - For an and event, use the single intersection cell. For an or event, combine the relevant row and column without double-counting their shared cell.
 - For an unconditional probability, divide the desired count by the table's grand total. Conditional probabilities use a restricted denominator and are covered in Section 5.4.

Example 1. If $P(A) = 0.30$, $P(B) = 0.25$, and A and B are mutually exclusive, find $P(A \text{ or } B)$.

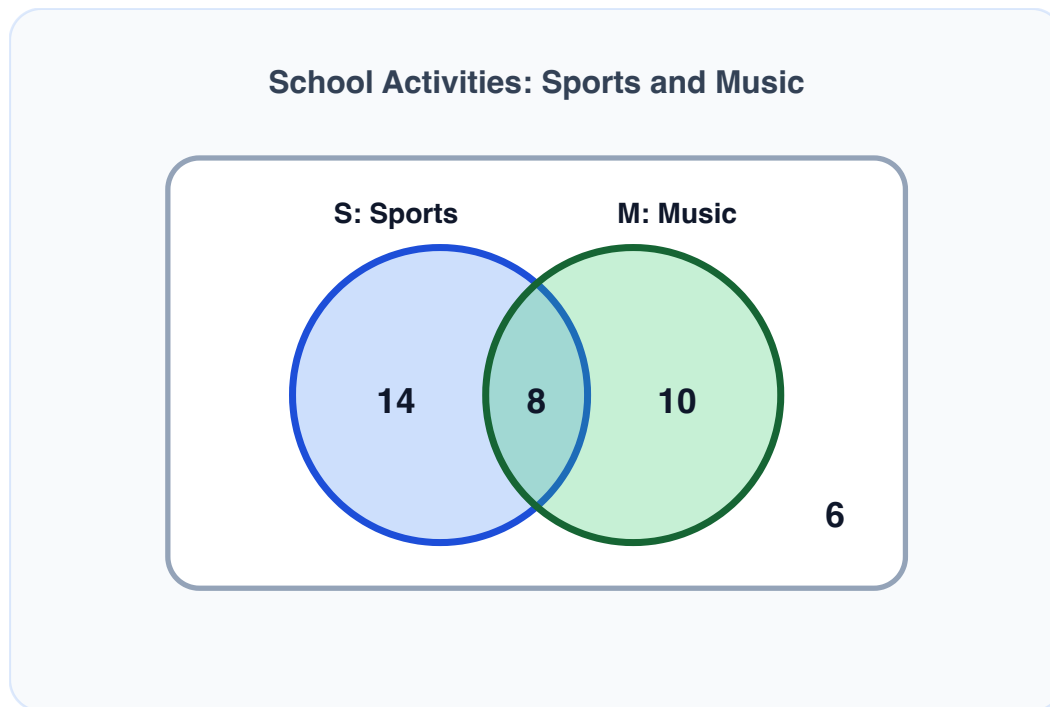
Answer: $0.30 + 0.25 = 0.55$.

Example 2. If $P(A) = 0.40$, $P(B) = 0.35$, and $P(A \text{ and } B) = 0.10$, find $P(A \text{ or } B)$.

Answer: $0.40 + 0.35 - 0.10 = 0.65$.

Example 3. Use the Venn diagram. (a) How many students were surveyed? (b) How many participate only in Sports? (c) Find $P(S \text{ and } M)$. (d) Find $P(S \text{ or } M)$. (e) Find the probability of neither activity. (f) Are Sports and Music mutually exclusive?

Reading Every Region of a Venn Diagram



Answer: (a) The total is $14 + 8 + 10 + 6 = 38$ students. (b) Sports only contains 14 students. (c) $P(S \text{ and } M) = 8/38 = 4/19$. (d) $P(S \text{ or } M) = (14 + 8 + 10)/38 = 32/38 = 16/19$. (e) $P(\text{neither}) = 6/38 = 3/19$. (f) No. The events are not mutually exclusive because 8 students participate in both activities.

Example 4. If the probability of at least one success is needed, what complement is often easier?

Answer: Find the probability of no successes, then subtract from 1.

Example 5. Use the two-way table to find the probability that a randomly selected adult is age 45–54 or uses social media.

Social media	18–34	35–44	45–54	55+	Total
Uses	117	86	83	47	333
Does not use	30	33	56	65	184
Total	147	119	139	112	517

Answer: The age 45–54 total is 139, the social-media total is 333, and their overlap is 83. Thus $P(45\text{--}54 \text{ or social media}) = (139 + 333 - 83)/517 = 389/517 \approx 0.752$.

Example 6. One card is drawn from a standard 52-card deck. Find each probability. (a) An ace or a ten. (b) An ace, a ten, or an eight. (c) A six or a diamond.

Answer: (a) The ranks are disjoint, so $P(\text{ace or ten}) = (4 + 4)/52 = 8/52 = 2/13$. (b) The three ranks are disjoint, so $P(\text{ace, ten, or eight}) = (4 + 4 + 4)/52 = 12/52 = 3/13$. (c) The six of diamonds is in both events, so $P(\text{six or diamond}) = (4 + 13 - 1)/52 = 16/52 = 4/13$.

Example 7. If $P(A \text{ or } B) = 0.64$, find the probability of neither A nor B .

Answer: Neither is the complement of the union, so the probability is $1 - 0.64 = 0.36$.

Example 8. A fair die is rolled once. Let $A = \{2, 4, 6\}$ be the event of rolling an even number and $B = \{4, 5, 6\}$ be the event of rolling a number greater than 3. List the outcomes in A or B , A and B , and A^c .

Answer: The union is $A \text{ or } B = \{2, 4, 5, 6\}$. The intersection is $A \text{ and } B = \{4, 6\}$. The complement is $A^c = \{1, 3, 5\}$.

Section 5.3 Independence and the Multiplication Rule

Use multiplication rules for independent and dependent events.

QUICK REFERENCE

- Independent events do not change each other's probabilities; dependent events do.
 - Use the multiplication check: the events are independent when $P(A \text{ and } B) = P(A)P(B)$.
 - Independence is an assumption that must fit the context. Repeated real-world outcomes are not automatically independent.
- Mutually exclusive and independent do not mean the same thing.
 - To check mutual exclusivity, ask: Can both events happen in the same trial? If not, they are mutually exclusive.
 - To check independence, ask: Does knowing that one event occurred change the probability of the other? If not, they are independent; if it does, they are dependent.
 - Mutually exclusive events with positive probabilities are dependent. Once one occurs, the probability of the other becomes 0.
 - Mutually exclusive events can also be independent only in the special case where at least one event has probability 0.

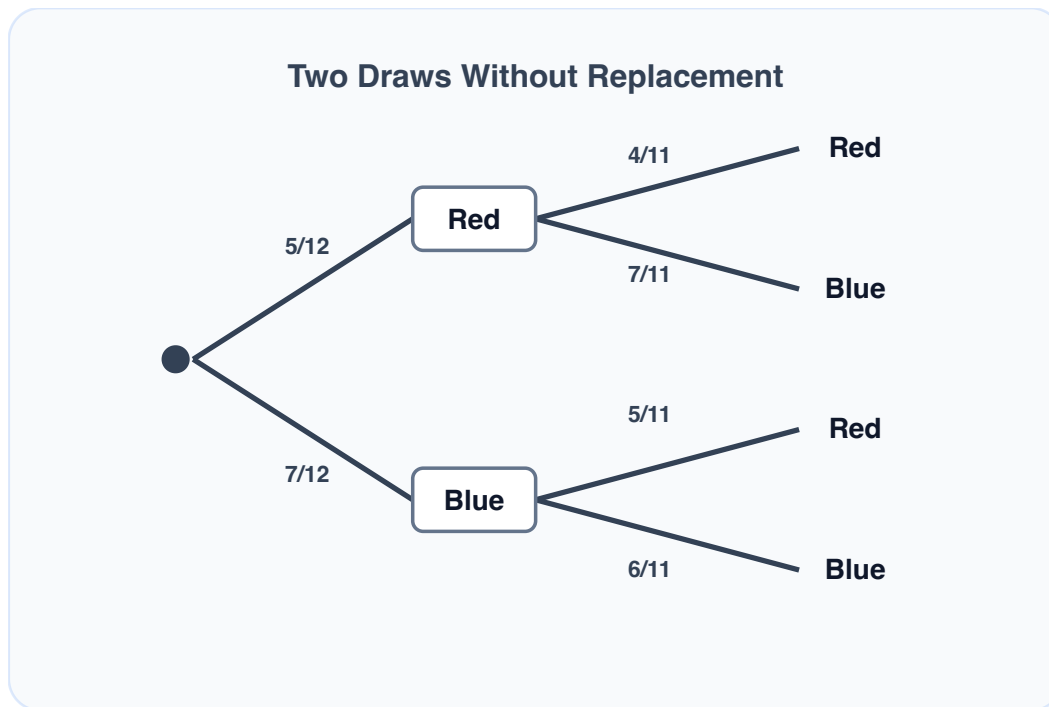
- For independent events, multiply probabilities to find an and probability.
 - $P(A \text{ and } B) = P(A)P(B)$. For repeated independent trials with the same probability p , the probability of success every time is p^n .
 - Interpret a small probability as a long-run frequency when requested: expected repetitions = number of repetitions $\times$ probability.
- With replacement usually keeps selections independent; without replacement usually makes them dependent.
 - With replacement, restore the selected item before the next draw, so both the numerator and denominator return to their original values.
 - Without replacement, update the remaining favorable items and total items after each draw. A tree diagram helps display the changing branches.
- Use a complement for at least one.
 - $P(\text{at least one}) = 1 - P(\text{none})$.
 - For independent trials, multiply the failure probabilities to find none, then subtract from 1.

Example 1. A coin is flipped twice. Find the probability of two heads.

Answer: $\frac{1}{2} \cdot \frac{1}{2} = \frac{1}{4}$.

Example 2. Use the tree diagram. A bag has 5 red and 7 blue marbles. Two marbles are drawn without replacement. Find the probability both are red.

Without-Replacement Tree Diagram



Answer: $\frac{5}{12} \cdot \frac{4}{11} = \frac{5}{33}$.

Example 3. Use the comparison table. Why does drawing without replacement usually create dependent events?

Setting	First red	Second red after first red	Are draws independent?
With replacement	5/12	5/12	Yes
Without replacement	5/12	4/11	No

Answer: Without replacement, the first draw changes the number of items left, so the second probability changes.

Example 4. Are two rolls of a fair die independent?

Answer: Yes. The first roll does not affect the second roll.

Example 5. A fair coin is flipped five times. Find and interpret the probability of five tails.

Answer: The flips are independent, so $P(\text{five tails}) = (1/2)^5 = 1/32 = 0.03125$. In 10,000 repetitions of five flips, about $10,000(0.03125) = 313$ repetitions would be expected to produce five tails.

Example 6. A screening test correctly returns positive for an affected person with probability 0.80. For five independent affected people, find the probability that all five tests are positive and the probability that at least one is negative.

Answer: All positive: $0.80^5 = 0.32768$. At least one negative: $1 - 0.80^5 = 0.67232$.

Example 7. Are the events 'a randomly selected student is a senior' and 'the student is a sophomore' mutually exclusive? Are they independent?

Answer: They are mutually exclusive because one student cannot be both a senior and a sophomore.

They are not independent. Once we know the student is a senior, the probability that the same student is a sophomore becomes 0, so knowing one event occurred changes the probability of the other.

Example 8. Use the multiplication check to determine whether the events are independent. In both cases, $P(A) = 0.40$ and $P(B) = 0.25$. (a) $P(A \text{ and } B) = 0.10$. (b) $P(A \text{ and } B) = 0.16$.

Answer: The product is $P(A)P(B) = 0.40(0.25) = 0.10$. (a) The events are independent because the intersection probability equals the product. (b) The events are dependent because 0.16 does not equal 0.10.

Example 9. One card is drawn from a standard 52-card deck. For each pair of events, determine whether the events are mutually exclusive or not mutually exclusive and whether they are independent or dependent. Support each classification using outcomes or probabilities. (a) Event A: Draw a heart. Event B: Draw a spade. (b) Event A: Draw a heart. Event B: Draw a face card. (c) Event A: Draw a heart. Event B: Draw a red card. (d) Event A: Draw a joker. Event B: Draw a heart.

Answer: (a) Mutually exclusive and dependent. One card cannot be both a heart and a spade, so $P(A \text{ and } B) = 0$, but $P(A)P(B) = (1/4)(1/4) = 1/16$. The events cannot occur together, and the multiplication check for independence fails.

(b) Not mutually exclusive and independent. The jack, queen, and king of hearts satisfy both events, so the events can occur together. Also, $P(A \text{ and } B) = 3/52$, which equals $P(A)P(B) = (13/52)(12/52) = 3/52$, so the multiplication check confirms independence.

(c) Not mutually exclusive and dependent. Every heart is red, so the events can occur together. However, $P(A \text{ and } B) = 13/52 = 1/4$, while $P(A)P(B) = (1/4)(1/2) = 1/8$, so the multiplication check shows dependence.

(d) Mutually exclusive and independent as a special zero-probability case. A standard 52-card deck has no jokers, so $P(A) = 0$, $P(A \text{ and } B) = 0$, and $P(A)P(B) = 0$. The events cannot occur together and satisfy the multiplication check only because one event has probability 0.

Example 10. One standard die is rolled. For each pair of events, determine whether the events are mutually exclusive or not mutually exclusive and whether they are independent or dependent. Support each classification using outcomes or probabilities. (a) Event A: Roll a 1. Event B: Roll an even number. (b) Event A: Roll an even number. Event B: Roll a number greater than 2. (c) Event A: Roll an even number. Event B: Roll a number less than 4. (d) Event A: Roll a 7. Event B: Roll an even number.

Answer: (a) Mutually exclusive and dependent. A roll cannot be both 1 and even, so $P(A \text{ and } B) = 0$, but $P(A)P(B) = (1/6)(3/6) = 1/12$. The events cannot occur together, and the multiplication check for independence fails.

(b) Not mutually exclusive and independent. Rolls 4 and 6 satisfy both events, so the events can occur together. Also, $P(A \text{ and } B) = 2/6 = 1/3$, which equals $P(A)P(B) = (3/6)(4/6) = 1/3$, so the multiplication check confirms independence.

(c) Not mutually exclusive and dependent. Roll 2 satisfies both events, so the events can occur together. However, $P(A \text{ and } B) = 1/6$, while $P(A)P(B) = (3/6)(3/6) = 1/4$, so the multiplication check shows dependence.

(d) Mutually exclusive and independent as a special zero-probability case. Rolling a 7 is impossible on a standard die, so $P(A) = 0$, $P(A \text{ and } B) = 0$, and $P(A)P(B) = 0$. The events cannot occur together and satisfy the multiplication check only because one event has probability 0.

Section 5.4 Conditional Probability and the General Multiplication Rule

Compute and interpret conditional probabilities.

QUICK REFERENCE

- Conditional probability restricts the sample space to known information.
 - $P(A | B)$ means the probability of A given that B has occurred.
 - Use $P(A | B) = \frac{P(A \text{ and } B)}{P(B)}$ or, with counts, number in both A and B / number in B .
 - The event after the vertical bar supplies the denominator.
- The order in conditional notation matters.
 - $P(A | B)$ and $P(B | A)$ usually have different restricted sample spaces and different denominators.
 - Translate the condition into words before selecting numbers from a table.
- In a two-way table, condition on a row or column total.
 - Locate the row or column named after given. Its total is the denominator.
 - The numerator is the intersection cell that also satisfies the event before the vertical bar.
- The general multiplication rule works for independent or dependent events.
 - $P(A \text{ and } B) = P(A)P(B | A) = P(B)P(A | B)$.
 - For independent events, the conditional probability equals the original probability, reducing the rule to $P(A)P(B)$.
 - With replacement or replay allowed, the original probabilities return before the next selection. Without replacement, multiply along the desired path while updating the favorable count and total after every selection.
- Use addition when an event can occur through more than one disjoint path.
 - For exactly one of two types, calculate first type then second plus second type then first.
 - For one red and one yellow without replacement, add $P(R \text{ then } Y) + P(Y \text{ then } R)$.
- Check independence by comparing conditional and unconditional probabilities.
 - For independent events, $P(A | B) = P(A)$ whenever $P(B) > 0$. This is equivalent to $P(A \text{ and } B) = P(A)P(B)$.
 - Compare the overall rate for an event with the rate inside the given group. These are $P(A)$ and $P(A | B)$.
 - If the rates are the same, allowing for normal rounding differences, the events are independent. If knowing the group changes the rate meaningfully, the events are dependent.

Example 1. If $P(A \text{ and } B) = 0.18$ and $P(B) = 0.30$, find $P(A | B)$.

Answer: $\frac{0.18}{0.30} = 0.60$.

Example 2. Use the two-way table. Find the probability a student takes biology given they take statistics.

Student group	Biology	Not biology	Total
Statistics	18	12	30
Not statistics	14	36	50
Total	32	48	80

Answer: Given Statistics, use the Statistics row: $\frac{18}{30} = 0.60$.

Example 3. Does $P(A | B)$ always equal $P(B | A)$?

Answer: No.

Example 4. If $P(B) = 0.40$ and $P(A | B) = 0.25$, find $P(A \text{ and } B)$.

Answer: $0.40(0.25) = 0.10$.

Example 5. Using the enrollment table from Example 2, find $P(\text{Statistics} | \text{Biology})$. Explain why it differs from $P(\text{Biology} | \text{Statistics})$.

Answer: Among the 32 biology students, 18 take statistics, so $P(\text{Statistics} | \text{Biology}) = 18/32 = 0.5625$. The reverse conditional is $18/30 = 0.60$ because it conditions on the 30 statistics students instead.

Example 6. A bag has 11 red, 10 yellow, and 9 purple bulbs. Two bulbs are drawn without replacement. Find the probability that the first is red and the second is yellow.

Answer: Use the general multiplication rule: $P(R \text{ then } Y) = 11/30 \cdot 10/29 = 110/870 \approx 0.1264$.

Example 7. Using the same bulb bag, find the probability that exactly one bulb is red and the other is yellow.

Answer: Add the two possible orders: $11/30 \cdot 10/29 + 10/30 \cdot 11/29 = 220/870 \approx 0.2529$.

Example 8. In a survey of adults, $P(\text{uses social media}) = 0.644$ and $P(\text{uses social media} | \text{age 35-44}) = 0.723$. Use the conditional-probability test to determine whether social-media use and being age 35-44 are independent.

Answer: Let A be using social media and B be being age 35-44. Independence requires $P(A | B) = P(A)$. Here, $P(A | B) = 0.723 \neq 0.644 = P(A)$, so the events are dependent. In words, 64.4% of all surveyed adults use social media, compared with 72.3% of adults ages 35-44. Knowing the age group changes the probability of social-media use.

Example 9. A Spotify playlist has 15 tracks, and you like 2 of them. The playlist first plays every track once in random order. For the first two tracks, find the probability that (a) you like both, (b) you like neither, and (c) you like exactly one. Is liking both unusual using the 0.05 guideline? Then repeat (a)-(c) if replay is allowed before all 15 tracks have played.

Answer: Without replay, the selections are dependent. (a) $2/15 \cdot 1/14 \approx 0.010$, which is unusual because $0.010 < 0.05$. (b) $13/15 \cdot 12/14 \approx 0.743$. (c) Exactly one can occur in either order: $2/15 \cdot 13/14 + 13/15 \cdot 2/14 \approx 0.248$. With replay allowed, the probabilities reset and the selections are independent. (a) $(2/15)^2 \approx 0.018$. (b) $(13/15)^2 \approx 0.751$. (c) $2(2/15)(13/15) \approx 0.231$.

Example 10. A five-card flush consists of five cards from the same suit. (a) Find the probability of five diamonds. (b) Find the probability of a flush in any of the four suits. (c) A royal flush is the ten, jack, queen, king, and ace of one suit. Find the probability of a royal flush.

Answer: (a) Without replacement, $P(\text{five diamonds}) = 13/52 \cdot 12/51 \cdot 11/50 \cdot 10/49 \cdot 9/48 \approx 0.000495$. (b) The four suit events are mutually exclusive, so $P(\text{flush}) = 4(0.000495198) \approx 0.001981$, or about 0.002. (c) Choose the suit in 4 ways. For that suit, the five required cards can appear in any order, so the successive favorable counts are 5, 4, 3, 2, and 1. Thus, $P(\text{royal flush}) = 4 \cdot 5/52 \cdot 4/51 \cdot 3/50 \cdot 2/49 \cdot 1/48 \approx 0.00000154$.

Section 5.5 Counting Techniques

Use the fundamental counting principle, permutations, and combinations.

QUICK REFERENCE

- The Fundamental Counting Principle multiplies the number of choices across successive stages.
 - If one stage has a choices and the next has b , there are ab ordered outcomes.
 - When repetition is allowed, the same choice count may be reused at each stage, such as 5^7 seven-digit keypad codes using five keys.
- A factorial counts arrangements of distinct objects.
 - $n! = n(n-1)(n-2) \cdots 1$, and $0! = 1$.
 - Arranging all n distinct objects uses $n!$.
- Use a permutation when order or assigned roles matter.
 - ${}_nP_r = \frac{n!}{(n-r)!}$ counts ordered selections of r distinct objects from n .
 - Examples include race finishes, officer positions, and ordered playlists without repetition.
- Use a combination when only the selected group matters.
 - ${}_nC_r = \frac{n!}{r!(n-r)!}$ counts unordered selections of r objects from n .
 - Examples include committees, card hands, lottery selections, and simple random samples.
- For arrangements with identical repeated objects, divide out duplicate orders.
 - If n objects include repeated groups of sizes $n_1, n_2, \dots$, the number of distinguishable arrangements is $\frac{n!}{n_1!n_2!\cdots}$.
 - Use this for repeated letters, repeated tree types, or birth-order categories with fixed counts.
- Counting can supply the numerator and denominator of a probability.
 - For equally likely outcomes, $P(E) = \frac{\text{favorable outcomes}}{\text{total outcomes}}$. Count both with methods that match the selection rules.
 - When selecting without replacement from two groups, count a specific composition by multiplying combinations. For example, choose the desired number of liked tracks and the remaining number of unliked tracks, then divide by the number of all possible selections.
 - On a TI-84, enter n , press MATH $\rightarrow$ PRB, select nPr or nCr, enter r , and press ENTER. Factorial is in the same menu.

Example 1. A code uses 3 letters followed by 2 digits. If repetition is allowed, how many codes are possible?

Answer: $26^3 \cdot 10^2 = 1,757,600$.

Example 2. How many ways can 4 students be arranged in a line from a group of 10?

Answer: ${}_{10}P_4 = 5040$.

Example 3. How many committees of 4 can be chosen from 10 students?

Answer: ${}_{10}C_4 = 210$.

Example 4. Should order matter when choosing a committee?

Answer: No, so use a combination.

Example 5. Evaluate $8!$, ${}_{13}P_0$, and ${}_8C_8$.

Answer: $8! = 40,320$, ${}_{13}P_0 = 1$, and ${}_8C_8 = 1$. There is one way to arrange zero objects and one way to choose all available objects.

Example 6. A race has 39 cars. In how many ways can first, second, and third place be awarded? Which TI-84 command applies?

Answer: Order matters, so use ${}_{39}nPr$: ${}_{39}P_3 = 39(38)(37) = 54,834$ finishes.

Example 7. How many distinguishable arrangements can be made from the letters in BALLOON?

Answer: There are 7 letters with two Ls and two Os, so the count is $7!/(2!2!) = 1260$.

Example 8. A shipment has 20 components, 3 defective and 17 good. Four components are selected without replacement. Find the probability that the shipment is rejected if at least one selected component is defective.

Answer: Use the complement of selecting four good components: $1 - \frac{{}_{17}C_4}{{}_{20}C_4} = 1 - 2380/4845 \approx 0.5088$.

Example 9. A Spotify playlist has 17 tracks, and you like 6 of them. Each track plays once in random order. Among the first 5 tracks, find the probability that you like (a) exactly 2, (b) exactly 3, and (c) all 5.

Answer: The first 5 tracks form an unordered selection from 17 tracks, so the denominator is ${}_{17}C_5 = 6188$. (a) Choose 2 of the 6 liked tracks and 3 of the 11 unliked tracks: $\frac{{}_6C_2 \cdot {}_{11}C_3}{{}_{17}C_5} = 2475/6188 \approx 0.4000$.

(b) $\frac{{}_6C_3 \cdot {}_{11}C_2}{{}_{17}C_5} = 1100/6188 \approx 0.1778$. (c) $\frac{{}_6C_5 \cdot {}_{11}C_0}{{}_{17}C_5} = 6/6188 \approx 0.0010$.

Example 10. A player is dealt a five-card hand from a standard 52-card deck. Find the probability of receiving exactly four cards of one rank.

Answer: There are ${}_{52}C_5 = 2,598,960$ possible five-card hands. Choose the four-of-a-kind rank in 13 ways. All four suits of that rank must be included, and the fifth card can be any 1 of the 48 cards from the other ranks. The number of favorable hands is $13 \cdot {}_4C_4 \cdot {}_{48}C_1 = 13(1)(48) = 624$, so $P(\text{exactly four of a kind}) = 624/2,598,960 \approx 0.000240$.

Module 6: Discrete Random Variables and Binomial Distributions

Students work with discrete probability distributions, expected value, and binomial probabilities.

Section 6.1 Discrete Random Variables

Determine whether a distribution is valid and compute expected values.

QUICK REFERENCE

- A random variable assigns a numerical value to each outcome of a probability experiment.
 - A discrete random variable has a finite or countable set of possible values, often counts such as 0, 1, 2, and so on.
 - A continuous random variable can take any value in an interval, often measurements such as time, weight, area, or distance.
 - State realistic possible values when classifying a variable; counts use whole numbers, while measurements may include decimals.
- A discrete probability distribution lists each possible value and its probability.
 - Every probability must satisfy $0 \leq P(x) \leq 1$, and $\sum P(x) = 1$.
 - To find a missing probability, subtract the sum of the known probabilities from 1.
 - A probability histogram uses separated bars because the random variable takes discrete values.
- Expected value is the probability-weighted long-run mean.
 - $\mu_X = E(X) = \sum xP(x)$. Multiply each value by its probability and add the products.
 - Interpret expected value as an average over many repetitions, not as a prediction for one trial. It does not need to be a possible individual outcome.
 - For games, insurance, or investments, a positive expected value is an average gain and a negative expected value is an average loss from the stated point of view.
 - Distinguish a payout or prize from net gain. Subtract the cost to play when the listed values are payouts, but do not subtract it again when the problem already gives net gains and losses.
- The standard deviation describes the typical distance of outcomes from the expected value.
 - For formula context, $\sigma_X = \sqrt{\sum (x - \mu_X)^2 P(x)}$.
 - Report the standard deviation in the same units as the random variable. A larger value means more variable outcomes.
- TI-84 weighted-statistics workflow:
 - Enter possible values in L1 and probabilities in L2.
 - Run 1-Var Stats with List:L1 and FreqList:L2.
 - Read $\bar{x}$ as the expected value and σ_x as the distribution's standard deviation. The displayed n may equal 1 when probabilities are used as weights; it is not the number of trials.

Example 1. Is this a valid probability distribution?

x	0	1	2	3
$P(x)$	0.20	0.35	0.25	0.20

Answer: Yes. All probabilities are between 0 and 1, and they sum to 1.

Example 2. Find the expected value for $x = 0, 1, 2$ with probabilities 0.2, 0.5, 0.3.

Answer: $\mu = 0(0.2) + 1(0.5) + 2(0.3) = 1.1$.

Example 3. A probability table sums to 1.08. Is it valid?

Answer: No. The probabilities must sum to 1.

Example 4. Classify each random variable and state its possible values. (a) Number of website visits in a day. (b) Time needed to download a file.

Answer: (a) Discrete; possible values are 0, 1, 2, and so on. (b) Continuous; possible values are nonnegative times, including decimals.

Example 5. Find the missing probability and verify the distribution.

x	0	1	2	3
P(x)	0.18	0.42	0.27	?

Answer: The missing probability is $1 - (0.18 + 0.42 + 0.27) = 0.13$. Every probability is between 0 and 1 and the completed probabilities sum to 1, so the distribution is valid.

Example 6. A \$10 game pays a net gain of \$340 with probability $1/38$ and a net loss of \$10 otherwise. Find and interpret the expected value.

Answer: $E(X) = 340(1/38) + (-10)(37/38) \approx -0.79$. In the long run, a player loses about \$0.79 per play on average; this does not mean every play loses exactly \$0.79.

Example 7. The table gives the cash-prize distribution for a \$1 lottery ticket. Find and interpret the expected cash prize. Then find and interpret the expected profit from one ticket.

Cash prize, x	P(x)
\$13,000,000	0.00000000699
\$200,000	0.00000015
\$10,000	0.000001655
\$100	0.000172835
\$7	0.004331105
\$4	0.006950843
\$3	0.01357866
\$0	0.97496474501

Answer: The expected cash prize is approximately \$0.25. Since the ticket costs \$1, the expected profit is $\$0.25 - \$1 = -\$0.75$. Over many tickets, a player receives about \$0.25 in prizes and loses about \$0.75 per ticket on average; this does not describe the result of every ticket.

Example 8. An insurance company sells a \$220,000 one-year term life insurance policy for a \$190 premium. The probability that the insured person survives the year is 0.999544. Determine the company's value if the insured person survives and if the insured person dies. Then find and interpret the expected value of the policy to the insurance company.

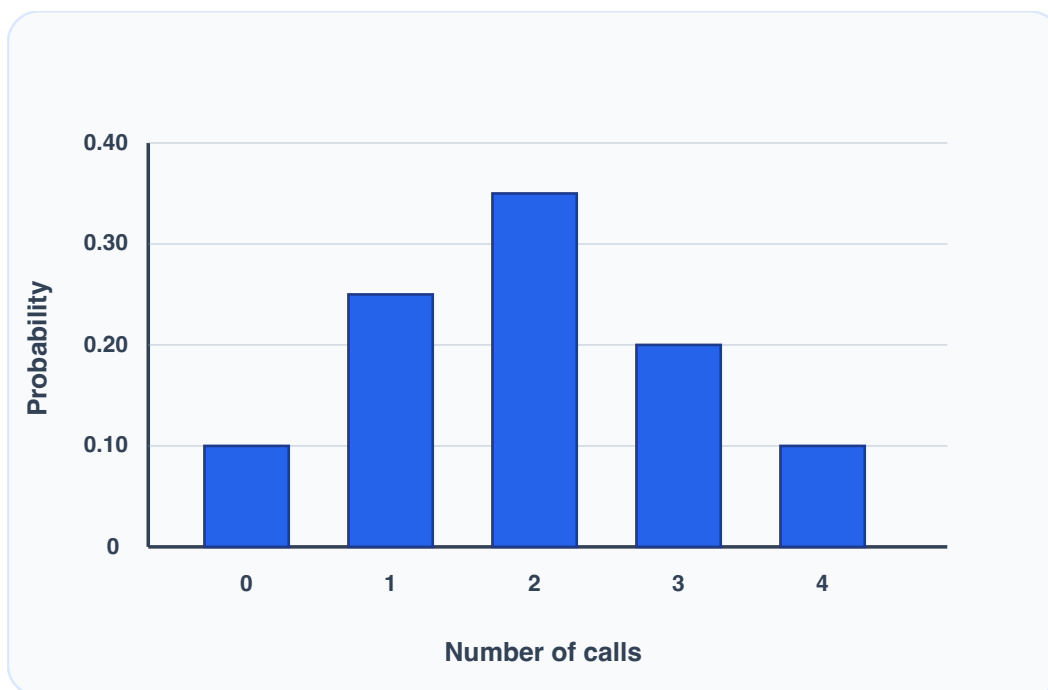
Answer: If the insured person survives, the company gains \$190. If the insured person dies, the company's value is $\$190 - \$220,000 = -\$219,810$, with probability $1 - 0.999544 = 0.000456$. Therefore, $E(X) = 190(0.999544) + (-219,810)(0.000456) = 89.68$. Across many similar one-year policies, the insurer averages about \$89.68 profit per policy; it does not make exactly that amount on every policy.

Example 9. The table gives the probability distribution for the number of urgent service calls received in one hour. Verify the distribution, construct its probability histogram, then use TI-84 weighted 1-Var Stats to find and interpret the mean and standard deviation.

Calls, x	0	1	2	3	4
P(x)	0.10	0.25	0.35	0.20	0.10

Answer: Every probability is between 0 and 1, and the probabilities total 1.00, so the distribution is valid. Enter 0, 1, 2, 3, 4 in L1 and 0.10, 0.25, 0.35, 0.20, 0.10 in L2. Weighted 1-Var Stats gives $\bar{x} = 1.95$ and σ_x approximately 1.12. Over many hours, the company averages about 1.95 urgent calls per hour, and hourly counts typically vary from that mean by about 1.12 calls.

Urgent Service Calls per Hour



Lists	L1: 0, 1, 2, 3, 4	L2: 0.10, 0.25, 0.35, 0.20, 0.10
Output	$\bar{x} = 1.95$	$\sigma_x \approx 1.12$

Section 6.2 The Binomial Probability Distribution

Recognize binomial settings and compute binomial probabilities with technology.

QUICK REFERENCE

- A binomial experiment must satisfy four conditions.
 - There is a fixed number of trials, n .
 - Each trial has two outcomes labeled success and failure, the trials are independent, and the success probability p remains constant.
 - Sampling without replacement can be treated as approximately independent when the sample is no more than 5% of a large population.
- Define the binomial variables before calculating.
 - n is the number of trials, p is the success probability, $q = 1 - p$ is the failure probability, and X counts successes.
 - State what success means in context; it is simply the outcome being counted, not necessarily a desirable result.
- The binomial probability formula shows the probability structure.
 - $P(X = x) = \binom{n}{x} p^x q^{n-x}$. The combination counts where the successes can occur.
 - Use the formula for context and the TI-84 for routine probability calculations.
- Translate probability wording before choosing a calculator command.
 - `binompdf` finds the probability of one exact count. `binomcdf` accumulates probabilities from 0 through its final argument.
 - Exactly k : $P(X = k)$, so use `binompdf(n,p,k)`.
 - At most or no more than k : $P(X \leq k)$, so use `binomcdf(n,p,k)`.
 - Fewer than or less than k : $P(X < k) = P(X \leq k - 1)$, so use `binomcdf(n,p,k-1)`.
 - More than or greater than k : $P(X > k) = 1 - P(X \leq k)$, so use `1-binomcdf(n,p,k)`.
 - At least k : $P(X \geq k) = 1 - P(X \leq k - 1)$, so use `1-binomcdf(n,p,k-1)`.
 - Between a and b , inclusive: use `binomcdf(n,p,b)-binomcdf(n,p,a-1)`.
 - Strictly between a and b : use `binomcdf(n,p,b-1)-binomcdf(n,p,a)`.
- A binomial distribution has mean $\mu = np$ and standard deviation $\sigma = \sqrt{npq}$.
 - The mean is the expected number of successes over many repetitions of the whole experiment; it need not be a whole number.
 - A count more than about 2 standard deviations from the mean may be considered unusual. Use probabilities when the question gives a specific unusual-event cutoff.
- The binomial shape depends on n and p .
 - When $p = 0.5$, the distribution is symmetric. Small p tends to skew right, and large p tends to skew left.
 - As n grows and both np and nq grow, the distribution becomes more bell-shaped.

Example 1. A basketball player makes 70% of free throws and shoots 8 free throws. Identify n and p .

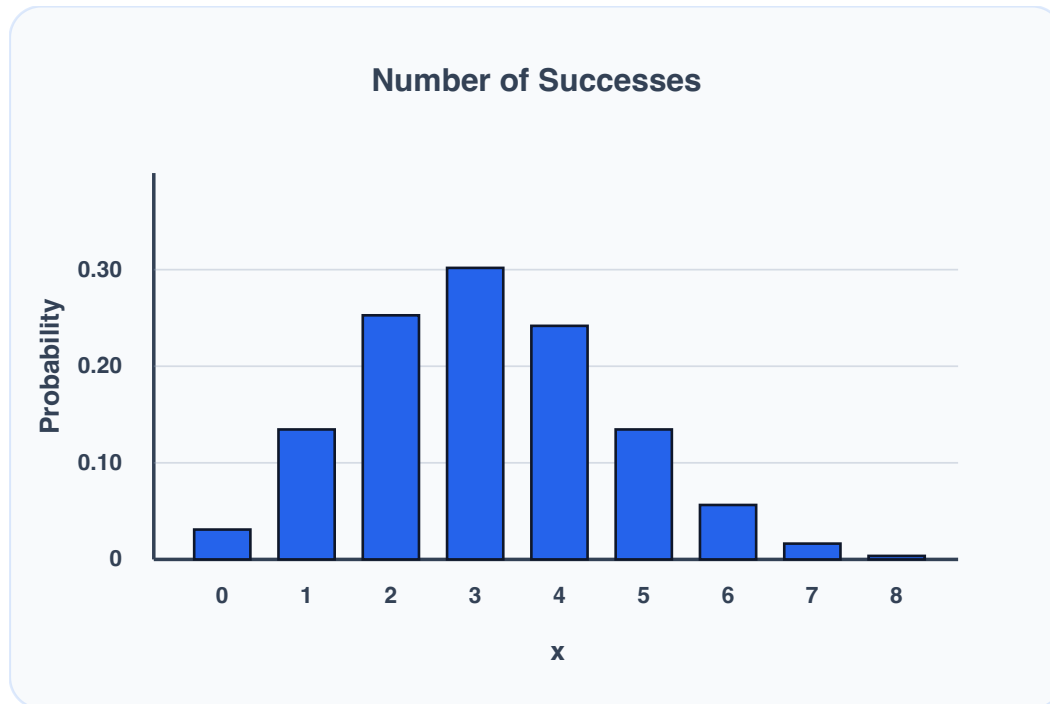
Answer: $n = 8, p = 0.70$.

Example 2. Use binomial notation to compute the probability of exactly 3 successes in 10 trials when $p = 0.40$.

Answer: Use `binompdf(10,0.40,3)`; approximately 0.215.

Example 3. Use the binomial distribution graph. For $X \sim \text{Binomial}(8, 0.35)$, near which value should the highest bars occur?

Binomial Distribution: $n = 8, p = 0.35$



Answer: Near the mean $np = 8(0.35) = 2.8$, so the highest bars should be near $x = 2$ or $x = 3$.

Example 4. For $X \sim \text{Binomial}(12, 0.25)$, what TI-84 command gives $P(X \leq 4)$, meaning "at most 4"?

Answer: Use `binomcdf(12,0.25,4)`, giving approximately 0.8424.

Example 5. For $X \sim \text{Binomial}(12, 0.25)$, how can $P(X \geq 5)$, meaning "at least 5," be found using the complement?

Answer: $P(X \geq 5) = 1 - P(X \leq 4)$, so use `1-binomcdf(12,0.25,4)`, giving approximately 0.1576.

Example 6. For $X \sim \text{Binomial}(12, 0.25)$, what commands find $P(3 < X < 7)$?

Answer: This means $P(4 \leq X \leq 6)$. Use `binomcdf(12,0.25,6)-binomcdf(12,0.25,3)`, giving approximately 0.3370.

Example 7. For $X \sim \text{Binomial}(12, 0.25)$, find the mean and standard deviation.

Answer: $\mu = np = 12(0.25) = 3$. Since $q = 0.75$, $\sigma = \sqrt{12(0.25)(0.75)} = 1.5$.

Example 8. Decide whether each setting is binomial. (a) A goalie attempts 20 penalty saves with a constant 0.42 save probability. (b) A goalie keeps attempting saves until the first goal is allowed. (c) Thirty patients receive a drug, but their success probabilities differ by dosage.

Answer: (a) Binomial with $n = 20$ and $p = 0.42$. (b) Not binomial because the number of trials is not fixed. (c) Not binomial because the success probability is not constant.

Example 9. For $X \sim \text{Binomial}(10, 0.30)$, describe the shape. What change would make a binomial distribution more bell-shaped?

Answer: Because $p = 0.30$, smaller success counts are more likely and the tail extends toward the larger counts, so the distribution is skewed right. For a fixed p , increasing n generally makes the distribution less skewed and more nearly bell-shaped as the expected numbers of successes and failures increase.

Example 10. A poll has $n = 250$ and $p = 0.45$. Find the expected count and standard deviation. Would 90 successes be unusual by the 2-standard-deviation guideline?

Answer: $\mu = 250(0.45) = 112.5$ and $\sigma = \sqrt{250(0.45)(0.55)} \approx 7.87$. The count 90 is about $(90 - 112.5)/7.87 \approx -2.86$ standard deviations from the mean, so it is unusual by that guideline.

Example 11. For $X \sim \text{Binomial}(12, 0.25)$, find each probability. (a) Fewer than 3 successes. (b) More than 4 successes.

Answer: (a) Fewer than 3 means $X < 3$, or $X \leq 2$. Use $\text{binomcdf}(12, 0.25, 2)$, giving approximately 0.3907. (b) More than 4 means $X > 4$. Use $1 - \text{binomcdf}(12, 0.25, 4)$, giving approximately 0.1576.

Example 12. An airline reports that each flight arrives on time with probability 0.80. For 25 independent flights, find the probability that between 14 and 16 flights, inclusive, arrive on time. Interpret the result in 100 repeated samples of 25 flights.

Answer: Let X count the on-time flights. Use $\text{binomcdf}(25, 0.80, 16) - \text{binomcdf}(25, 0.80, 13)$, giving $P(14 \leq X \leq 16) \approx 0.0452$. In 100 repeated samples of 25 flights, about 5 samples would have between 14 and 16 on-time flights, inclusive.

Module 7: Normal Distributions

Students use normal distribution properties, normal probabilities, inverse normal values, and normality checks.

Section 7.1 Properties of the Normal Distribution

Understand normal curves, z-scores, uniform distributions, and area interpretations.

QUICK REFERENCE

- A continuous probability density curve represents probabilities by area.
 - The total area under the curve is 1, and the area over an interval equals the probability that the variable falls in that interval.
 - For a continuous variable, a single exact value has area and probability 0, so $P(X < a) = P(X \leq a)$.

- A continuous uniform distribution has constant height over its interval.
 - Every equal-width subinterval has the same probability. Probability = interval length / total length.
 - For a uniform distribution from a to b , the rectangle's height is $1/(b - a)$.
- A normal distribution is bell-shaped, symmetric, and determined by μ and σ .
 - The mean, median, and mode are equal at the center. Half the area lies on each side of the mean, and the curve approaches but never touches the horizontal axis.
 - Changing μ shifts the center. A larger σ makes the curve wider and lower; a smaller σ makes it narrower and taller.
 - The inflection points occur at $\mu - \sigma$ and $\mu + \sigma$, where the curve changes bending direction.
- A z-score standardizes a normal value.
 - $z = \frac{x - \mu}{\sigma}$ tells how many standard deviations x is above or below the mean.
 - The standard normal distribution has $\mu = 0$ and $\sigma = 1$. Standardizing changes the scale, not the corresponding area.
- Normal distributions satisfy the bell-shaped condition for the Empirical Rule reviewed in Section 3.2.
 - Refer to the Section 3.2 visual and examples for the 68-95-99.7 estimates.
 - Use exact normal calculator methods in Section 7.2 when cutoffs are not whole numbers of standard deviations from the mean.

Example 1. A normal distribution has mean 50 and standard deviation 6. Find the z-score for $x = 62$.

Answer: $z = \frac{62 - 50}{6} = 2.$

Example 2. A normal distribution has mean 100 and standard deviation 15. What interval is within 1 standard deviation?

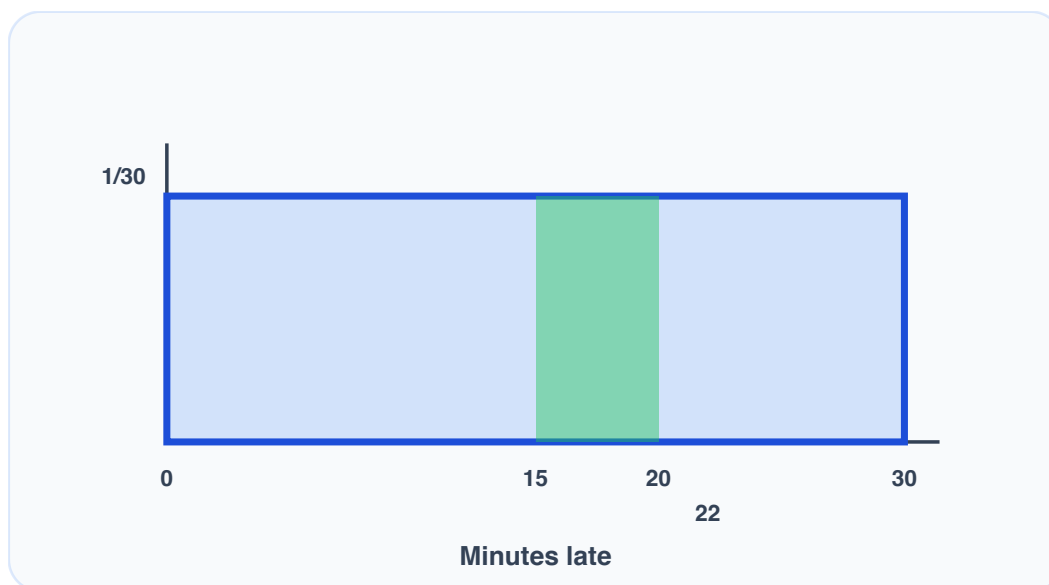
Answer: 85 to 115.

Example 3. What is the total area under a normal curve?

Answer: 1, or 100%.

Example 4. A friend is uniformly distributed from 0 to 30 minutes late. Find the probability the friend is between 15 and 20 minutes late and the probability the friend is at least 22 minutes late.

Uniform Arrival-Time Model



Answer: Use interval length / total length. $P(15 < X < 20) = 5/30 \approx 0.167$. $P(X \geq 22) = 8/30 \approx 0.267$.

Example 5. A normal curve is centered at 530, and adjacent inflection points are labeled 528 and 532. Identify μ and σ .

Answer: The center is $\mu = 530$. Each inflection point is 2 units from the mean, so $\sigma = 2$.

Example 6. Normal Curve A has $\mu = 9$, $\sigma = 4$. Normal Curve B has $\mu = 16$, $\sigma = 1$. Which curve is farther right, and which is narrower and taller?

Answer: Curve B is centered farther right because it has the larger mean. It is also narrower and taller because it has the smaller standard deviation.

Example 7. For pregnancy lengths, the area to the left of 250 days is 0.1587. Give two interpretations.

Answer: A randomly selected pregnancy has probability 0.1587 of lasting 250 days or less. In many pregnancies, about 15.87% would be expected to last 250 days or less.

Section 7.2 Applications of the Normal Distribution

Use TI-84 normalcdf and invNorm to compute probabilities and cutoff values.

QUICK REFERENCE

- Start by sketching or naming the region under the normal curve.
 - Less than shades left, greater than shades right, and between shades the region between two cutoffs.
 - Probability is area, so the answer must be between 0 and 1. Use symmetry and the cutoff's location to check whether the result is reasonable.

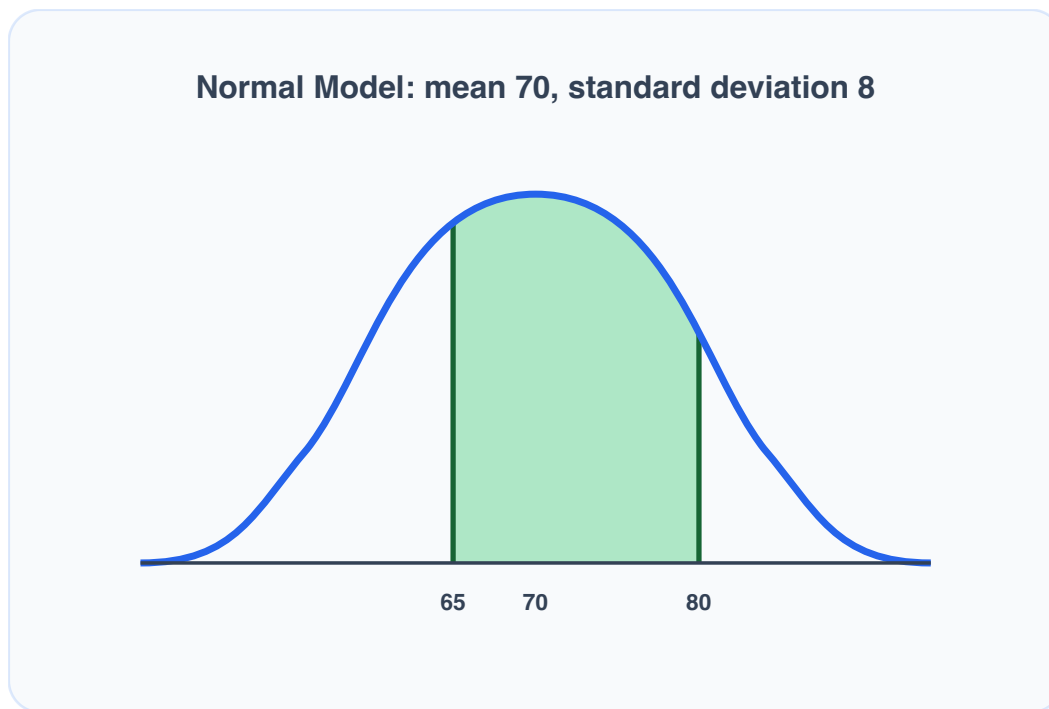
- Use `normalcdf(lower,upper,mean,standard deviation)` to find normal probabilities.
 - For less than a , use a very large negative lower bound and upper bound a . For greater than a , use lower bound a and a very large positive upper bound.
 - For between a and b , enter the two cutoffs in increasing order. For two tails, calculate each tail and add or use a complement when easier.
 - For standard-normal z-scores, use mean 0 and standard deviation 1.
- Use `invNorm(left-tail area,mean,standard deviation)` to find a cutoff from an area.
 - A percentile is already a left-tail area. For a right-tail area r , enter $1 - r$.
 - For the middle $C\%$ of a symmetric normal distribution, split the remaining area equally between the two tails, then find both cutoff percentiles.
 - Use mean 0 and standard deviation 1 for z-cutoffs, or enter the original μ and σ to obtain raw-value cutoffs directly.
- The notation z_α identifies a standard-normal right-tail cutoff.
 - z_α is the z-score with area α to its right, so the area to its left is $1 - \alpha$.
 - On a TI-84, use `invNorm(1-alpha,0,1)`. For $z_{\alpha/2}$, use left-tail area $1 - \alpha/2$; the middle area between $-z_{\alpha/2}$ and $z_{\alpha/2}$ is $1 - \alpha$.
- Interpret the calculator result in context.
 - A probability describes one randomly selected observation; the same decimal can be interpreted as an approximate long-run proportion or percentage.
 - A percentile cutoff is a value with the stated percentage of observations at or below it.

Example 1. Let X be normal with $\mu = 70$ and $\sigma = 8$. Use a TI-84 to find $P(X < 82)$.

Answer: Use `normalcdf(-1E99,82,70,8)`. The probability is approximately 0.9332.

Example 2. Use the shaded normal curve and a TI-84 to find $P(65 < X < 80)$ when $\mu = 70$ and $\sigma = 8$.

Shaded Normal Probability



Answer: Use `normalcdf(65,80,70,8)`. The probability is approximately 0.6284.

Example 3. Let X be normal with $\mu = 70$ and $\sigma = 8$. Use a TI-84 to find the 90th percentile.

Answer: Use `invNorm(0.90,70,8)`. The 90th-percentile cutoff is approximately 80.25, meaning about 90% of values are at or below 80.25.

Example 4. If the calculator gives $P(X > 90) = 0.0228$, interpret the answer.

Answer: About 2.28% of values are greater than 90.

Example 5. Find the standard-normal area to the right of $z = -0.17$. Give the TI-84 command and result.

Answer: Use `normalcdf(-0.17,1E99,0,1)`. The area is approximately 0.5675.

Example 6. Let X be normal with $\mu = 50$ and $\sigma = 8$. Find $P(34 < X < 62)$.

Answer: Use `normalcdf(34,62,50,8)`. The probability is approximately 0.9104.

Example 7. Find the combined standard-normal area below $z = -1.96$ and above $z = 1.96$.

Answer: The two symmetric tails total about $2(0.025) = 0.0500$. On the TI-84, add `normalcdf(-1E99,-1.96,0,1)` and `normalcdf(1.96,1E99,0,1)`.

Example 8. Incubation times are normal with mean 22 days and standard deviation 1 day. Find the cutoffs for the middle 95%.

Answer: The tails contain 2.5% each. Use `invNorm(0.025,22,1)` and `invNorm(0.975,22,1)`. The cutoffs are approximately 20.04 and 23.96 days.

Example 9. Find $z_{0.08}$ and explain what its subscript means.

Answer: The subscript means the area to the right is 0.08, so the left-tail area is 0.92. Use `invNorm(0.92,0,1)`, giving $z_{0.08} \approx 1.405$.

Section 7.3 Assessing Normality

Use graphs and normal probability plots to decide whether normal methods are reasonable.

QUICK REFERENCE

- Assess normality with graphs rather than expecting perfect normal data.
 - A roughly bell-shaped, symmetric histogram with one main peak supports a normal model.
 - Strong skewness, multiple peaks, large gaps, or clear outliers are warnings against a normal model. A boxplot can help reveal skewness and outliers but does not show the full shape.
- A normal probability plot compares ordered observations with expected normal z-scores.
 - Points that follow an approximately straight line support a normal model.
 - Systematic curvature suggests skewness or another nonnormal shape. A point far from the line may indicate an outlier.
 - To construct the plot, order the observations, calculate $f_i = (i - 0.375)/(n + 0.25)$, find each expected score with $\text{invNorm}(f_i, 0, 1)$, and plot observed values against expected z-scores.
- Use the normal-probability-plot critical-value table for the correlation check.
 - Find the correlation between the ordered observations and their expected z-scores, then find the critical value for n .
 - For more than 30 observations, use the $n = 30$ critical value, following the homework convention.
 - If the correlation is greater than the critical value, the plot is sufficiently linear to support normality. If it is not greater, the data do not pass that normality check.
- Use all available evidence and context.
 - A small sample can make a histogram jagged, while a very large sample can reveal minor departures from normality. Focus on departures that matter for the intended method.
 - Normality checks decide whether normal-based methods are reasonable; they do not prove that the population is exactly normal.
- Technology can create the required displays.
 - StatCrunch is useful for normal probability plots, histograms, and larger data sets.
 - On a TI-84, enter the ordered observations in L1 and expected z-scores in L2. A normal probability plot can be selected in STAT PLOT; LinReg(ax+b) L1,L2 gives the correlation when diagnostics are on. Compare that correlation with the quick-reference critical value.

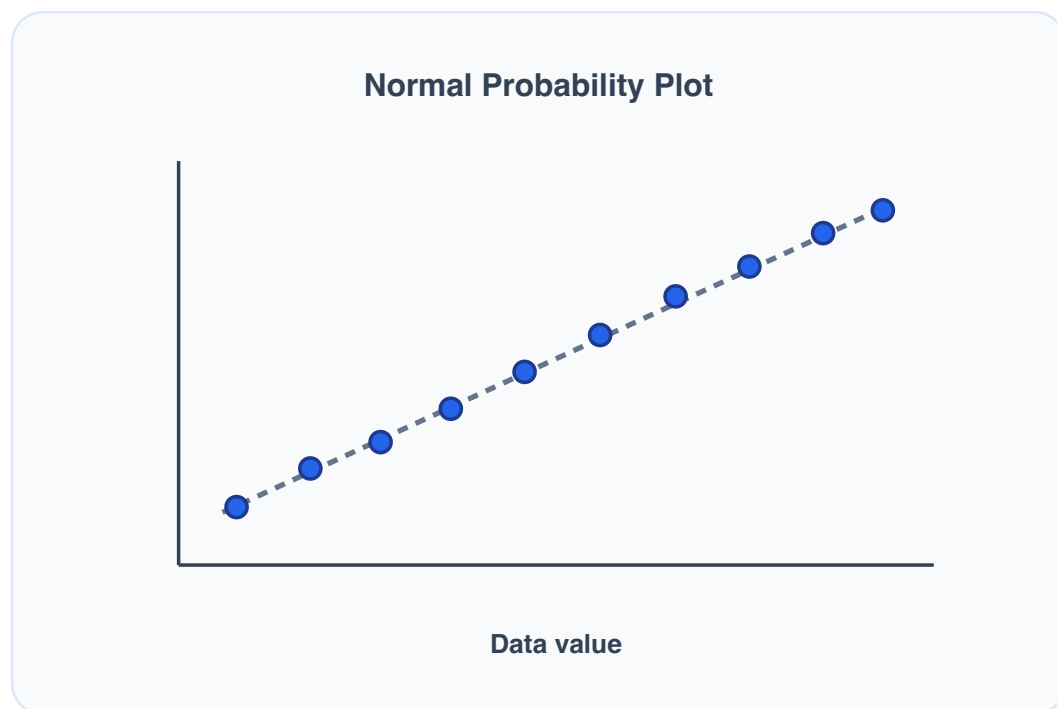
Critical Values for Normal Probability Plots

Sample size, n	Critical value	Sample size, n	Critical value
5	0.880	16	0.941
6	0.888	17	0.944
7	0.898	18	0.946
8	0.906	19	0.949
9	0.912	20	0.951

Sample size, n	Critical value	Sample size, n	Critical value
10	0.918	21	0.952
11	0.923	22	0.954
12	0.928	23	0.956
13	0.932	24	0.957
14	0.935	25	0.959
15	0.939	30	0.960

Example 1. Use the normal probability plot. What does this suggest about normality?

Normal Probability Plot



Answer: The data appear approximately normal because the points are close to a straight line.

Example 2. A histogram is strongly skewed right with several outliers. Are normal methods automatically reasonable?

Answer: No. Strong skewness and outliers are concerns.

Example 3. What graph is specifically designed to assess whether data are approximately normal?

Answer: Normal probability plot.

Example 4. Does a data set have to be perfectly normal for normal methods to be useful?

Answer: No. The data need to be approximately normal enough for the method being used.

Example 5. A normal probability plot for $n = 16$ observations has correlation 0.962. Use the quick-reference table to determine whether the sample passes the correlation check for normality.

Answer: Yes. Because 0.962 is greater than 0.941, the points are sufficiently linear by the supplied critical-value rule, supporting use of a normal model.

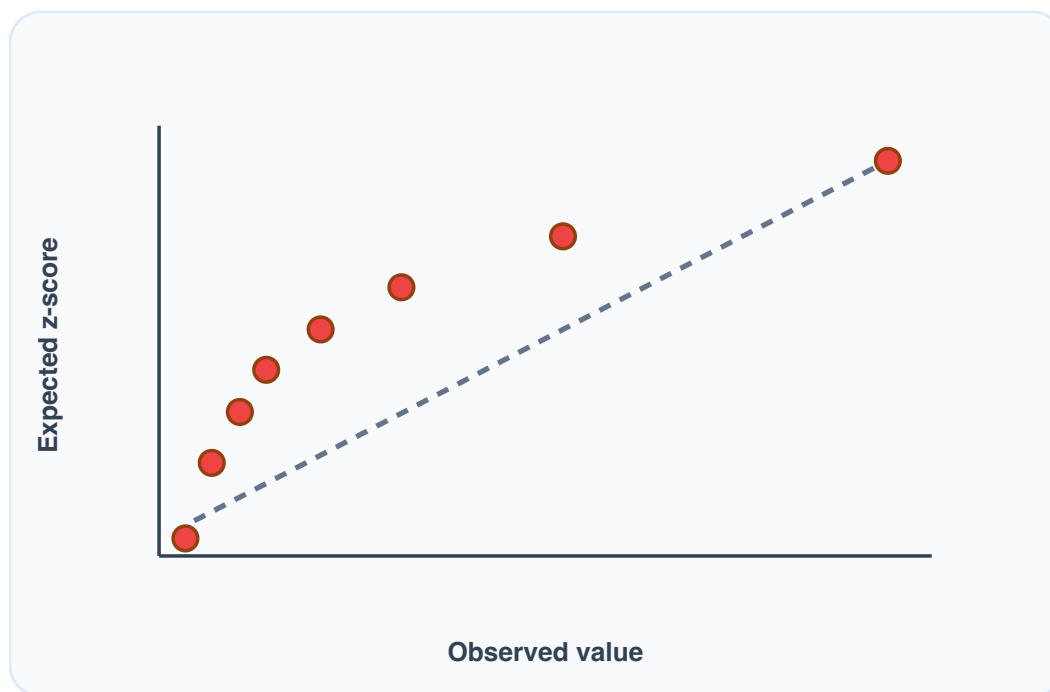
Example 6. A histogram is roughly symmetric and bell-shaped, but its normal probability plot has one point far from the line. What should be reported?

Answer: Most of the distribution supports normality, but the isolated point is a possible outlier that should be investigated before relying on a normal model.

Example 7. The observed values and their expected z-scores are shown for $n = 8$. Enter the columns in L1 and L2, use LinReg(ax+b) L1,L2, and compare the correlation with the quick-reference critical value. Does the sample pass the normality check?

Observed value	2	3	4	5	7	10	16	28
Expected z-score	-1.41	-0.84	-0.47	-0.15	0.15	0.47	0.84	1.41

Curved Normal Probability Plot



Answer: The calculator gives $r \approx 0.900$. Because 0.900 does not exceed the critical value 0.906, the sample does not pass the correlation check. The systematic curvature in the plot also warns that the data are right-skewed rather than approximately normal.

Module 8: Sampling Distributions

Students study sampling distributions for sample means and sample proportions.

Section 8.1 Distribution of the Sample Mean

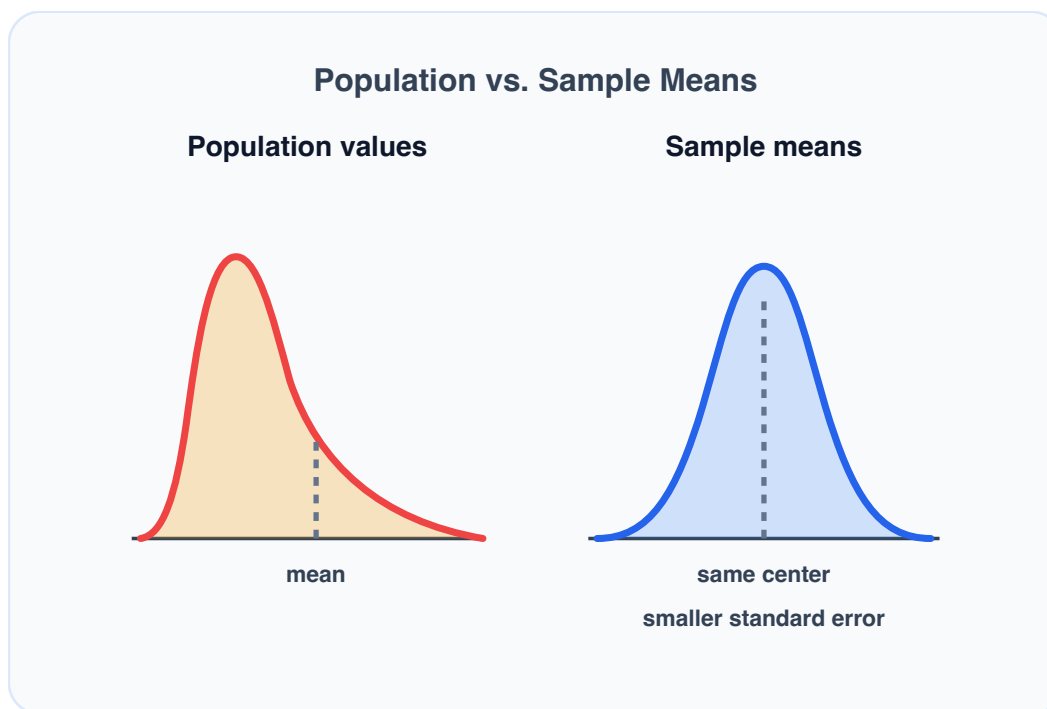
Use the sampling distribution of $\bar{x}$ and the Central Limit Theorem.

QUICK REFERENCE

- The sampling distribution of $\bar{x}$ is the distribution of sample means from all possible random samples of the same size n .
 - The population distribution describes individual values; one sample has its own data distribution; the sampling distribution describes how $\bar{x}$ varies from sample to sample.
 - A simulation approximates the sampling distribution by repeatedly taking samples, computing each mean, and graphing those means.
- The sampling distribution is centered at the population mean and has less spread than individual values.
 - $\mu_{\bar{x}} = \mu$. The sample mean is an unbiased estimator because its long-run mean equals the population mean.
 - $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$. This standard deviation of sample means is called the standard error.
 - As n increases, standard error decreases, so larger samples produce means that cluster more tightly around μ .
- Check independence and shape before using a normal model for $\bar{x}$.
 - Use a simple random sample or a design that produces representative, approximately independent observations. When sampling without replacement, check $n \leq 0.05N$.
 - If the population is normal, the sampling distribution of $\bar{x}$ is normal for any sample size.
 - If the population is not normal, the Central Limit Theorem says the sampling distribution becomes approximately normal as n grows. $n \geq 30$ is a common introductory guideline, but strong skewness or outliers may require a larger sample.
- Use the sampling-distribution parameters, not the population spread, in probability calculations.
 - For probabilities involving one individual value X , use mean μ and standard deviation σ .
 - For probabilities involving a sample mean $\bar{x}$, use mean μ and standard error $\sigma/\sqrt{n}$.
 - On the TI-84, use `normalcdf(lower,upper,mean,standard deviation)`, entering μ for the mean and $\sigma/\sqrt{n}$ for the standard deviation. Interpret the result as a proportion of samples or a probability for a randomly selected sample.

Example 1. Use the comparison figure. How does the sampling distribution of $\bar{x}$ compare with the population distribution?

Population Distribution versus Sampling Distribution



Answer: The population graph shows individual values and is skewed right. The sampling distribution shows sample means and is more nearly normal. Both are centered at μ , but the sample means are less spread out because their standard deviation is $\sigma/\sqrt{n}$.

Example 2. A population has $\mu = 80$, $\sigma = 12$, and samples of size 36. Find $\mu_{\bar{x}}$ and $\sigma_{\bar{x}}$.

Answer: $\mu_{\bar{x}} = 80$, $\sigma_{\bar{x}} = \frac{12}{\sqrt{36}} = 2$.

Example 3. For the same setting, what TI-84 command finds $P(\bar{x} < 83)$?

Answer: `normalcdf(-1E99,83,80,2)` gives $P(\bar{x} < 83) \approx 0.9332$.

Example 4. What happens to the standard error as sample size increases?

Answer: It decreases.

Example 5. Why is the Central Limit Theorem useful?

Answer: It lets sample means be modeled approximately normally in many large-sample situations.

Example 6. A population is strongly skewed right. Compare the expected sampling-distribution shapes for samples of size $n = 4$ and $n = 35$.

Answer: For $n = 4$, the sampling distribution may remain noticeably right-skewed. For $n = 35$, the Central Limit Theorem makes the distribution of sample means much closer to normal. Both are centered at the population mean, and the $n = 35$ distribution has smaller standard error.

Example 7. Pregnancy lengths are normal with $\mu = 266$ days and $\sigma = 16$ days. Compare the TI-84 setups for the probability that one pregnancy is under 260 days and the probability that the mean of $n = 16$ pregnancies is under 260 days.

Answer: One pregnancy uses `normalcdf(-1E99,260,266,16)`. A sample mean uses standard error $16/\sqrt{16} = 4$: `normalcdf(-1E99,260,266,4)`. The second probability is smaller because sample means vary less.

Example 8. A population has $\mu = 200$ and $\sigma = 10$. Without calculating exact values, order these probabilities from least to greatest: an individual value from 195 to 205, a sample mean with $n = 10$ in that interval, and a sample mean with $n = 20$ in that interval.

Answer: All three distributions are centered at 200. Their standard deviations are 10, $10/\sqrt{10} \approx 3.16$, and $10/\sqrt{20} \approx 2.24$, respectively. A smaller standard deviation places more values inside the same interval from 195 to 205. Therefore, $P(195 < X < 205) < P(195 < \bar{X}_{10} < 205) < P(195 < \bar{X}_{20} < 205)$.

Example 9. Which procedure correctly simulates the sampling distribution of the sample mean for samples of size $n = 20$? (a) Take one random sample of 20 observations and graph those observations. (b) Repeatedly take random samples of 20 observations, calculate each sample mean, and graph the collection of means. (c) Repeatedly take samples of changing sizes and graph every individual observation.

Answer: Procedure (b). A sampling distribution is built from the statistic produced by many samples of the same size. Procedures (a) and (c) graph individual observations rather than a collection of sample means from a fixed sample size.

Section 8.2 Distribution of the Sample Proportion

Use the sampling distribution of $\hat{p}$ and check normal approximation conditions.

QUICK REFERENCE

- The sample proportion is $\hat{p} = x/n$, where x is the number of sampled individuals with the characteristic.
 - The population proportion p is a fixed parameter; $\hat{p}$ varies from sample to sample and estimates p .
 - The sampling distribution of $\hat{p}$ describes the proportions from all possible random samples of the same size.
- The sampling distribution is centered at p and becomes less variable as n increases.
 - $\mu_{\hat{p}} = p$, so the sample proportion is an unbiased estimator of the population proportion.
 - $\sigma_{\hat{p}} = \sqrt{\frac{p(1-p)}{n}}$. This is the standard error when the population proportion is used in a sampling-distribution problem.

- Check the sampling conditions before using a normal model for $\hat{p}$.
 - Use a simple random sample or another representative random process. For sampling without replacement, check $n \leq 0.05N$.
 - For this course's homework, check $np(1 - p) \geq 10$. If this value is below 10, the sampling distribution may be too skewed for the normal approximation.
 - When solving for a required sample size, isolate n , round the result up, and subtract the current sample size if the question asks how many additional individuals are needed.
- Use the proportion mean and standard error in normal probability calculations.
 - On the TI-84, use `normalcdf(lower,upper,p,standard deviation)`, entering $\sqrt{p(1 - p)/n}$ for the standard deviation.
 - Interpret the result as the probability that a random sample produces a proportion in the stated range, or as the long-run proportion of such samples.

Example 1. Let $p = 0.40$ and $n = 100$. Find the mean and standard deviation of $\hat{p}$.

Answer: Mean = 0.40. Standard deviation = $\sqrt{\frac{0.40(0.60)}{100}} \approx 0.049$.

Example 2. Check the normal approximation conditions for $p = 0.40$, $n = 100$.

Answer: $np(1 - p) = 100(0.40)(0.60) = 24$, so the value is at least 10.

Example 3. What TI-84 command finds $P(\hat{p} < 0.45)$ when $p = 0.40$, $n = 100$, and standard error is 0.049?

Answer: `normalcdf(-1E99,0.45,0.40,0.049)`.

Example 4. If $np(1 - p) = 7$, what concern should be noted?

Answer: The normal approximation condition is not met.

Example 5. In a sample of 1400 adults, 154 report being afraid to fly. Find the sample proportion. Is it required to equal a previously reported population proportion of 0.10?

Answer: $\hat{p} = 154/1400 = 0.11$. No. Sample proportions vary naturally from sample to sample and are not expected to match p exactly.

Example 6. A population has $N = 20,000$, $p = 0.10$, and a simple random sample has $n = 400$. Check the conditions and describe the sampling distribution of $\hat{p}$.

Answer: The 5% condition holds because 400 is at most 1000. Also, $np(1 - p) = 400(0.10)(0.90) = 36$, which is at least 10. Thus $\hat{p}$ is approximately normal with mean 0.10 and standard error $\sqrt{0.10(0.90)/400} = 0.015$.

Example 7. Suppose $p = 0.82$ and $n = 100$. Find $P(\hat{p} > 0.88)$ and interpret it.

Answer: The standard error is $\sqrt{0.82(0.18)/100} \approx 0.0384$. Use `normalcdf(0.88,1E99,0.82,0.0384)`, giving approximately 0.059. About 5.9% of random samples of 100 would have more than 88% successes.

Example 8. A historical proportion is $p = 0.10$. In a new random sample of $n = 1200$, 132 people have the characteristic. Explain how the sampling distribution helps decide whether the sample proportion is surprising.

Answer: The sample proportion is $\hat{p} = 132/1200 = 0.11$, and the normal condition holds because $np(1 - p) = 1200(0.10)(0.90) = 108$. Model $\hat{p}$ with mean 0.10 and standard error $\sqrt{0.10(0.90)/1200} \approx 0.00866$. Use `normalcdf(0.11,1E99,0.10,0.00866)`, giving approximately 0.1241. If the population proportion is still 0.10, about 12.4% of random samples of 1200 would produce a sample proportion of 0.11 or greater. This is not unusual using the 5% guideline, so the sample result is not necessarily evidence that the population proportion increased.

Example 9. A researcher already plans to sample 50 adults. How many additional adults are needed for the sampling distribution of $\hat{p}$ to meet $np(1 - p) \geq 10$ if (a) $p = 0.20$ and (b) $p = 0.25$?

Answer: (a) Solve $n(0.20)(0.80) \geq 10$, giving $n \geq 62.5$. Round up to 63 total adults, so 13 more are needed. (b) Solve $n(0.25)(0.75) \geq 10$, giving $n \geq 53.33$. Round up to 54 total adults, so 4 more are needed.

Module 9: Confidence Intervals

Students estimate population proportions and means using confidence intervals.

Section 9.1 Estimating a Population Proportion

Use one-proportion confidence intervals and interpret them correctly.

QUICK REFERENCE

- A point estimate is a sample statistic used to estimate a population parameter.
 - For a population proportion p , the point estimate is $\hat{p} = x/n$. The variable must have two outcomes after defining success and failure.
 - A confidence interval has the form point estimate $\pm$ margin of error. Its midpoint is the point estimate, and $E = (\text{upper} - \text{lower})/2$.
- The confidence level describes the long-run success rate of the interval method.
 - The population proportion is fixed but unknown. Different random samples produce different sample proportions and therefore different intervals.
 - If many random samples were taken and an interval were built from each, about $C\%$ of those intervals would capture the fixed population proportion.
 - Once one interval is computed, the fixed proportion either is or is not inside it. Say, 'We are $C\%$ confident that the interval captures the population proportion,' not that there is a $C\%$ probability the parameter is inside.

- Check conditions for a one-proportion z-interval.
 - Use a random or representative sample and, when sampling without replacement, check $n \leq 0.05N$.
 - For this course's homework, check $n\hat{p}(1 - \hat{p}) \geq 10$ using the observed sample proportion.
- For formula context, $\hat{p} \pm z^* \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}$.
 - Higher confidence uses a larger critical value and produces a wider interval when other information is fixed.
 - A larger sample size reduces standard error and produces a narrower interval when the confidence level and sample proportion are fixed.
- The critical value $z^* = z_{\alpha/2}$ captures the middle $(1 - \alpha)100\%$ of the standard normal distribution.
 - On the TI-84, use `invNorm(1-alpha/2,0,1)`. Common values are 1.645 for 90%, 1.960 for 95%, and 2.576 for 99% confidence.
 - Each tail has area $\alpha/2$, so the positive critical value has left-tail area $1 - \alpha/2$.
- Choose a sample size large enough to achieve the requested margin of error E .
 - With a prior estimate $\tilde{p}$, use $n = \tilde{p}(1 - \tilde{p})(z^*/E)^2$. Without a prior estimate, use $n = 0.25(z^*/E)^2$.
 - Write a percent margin of error as a decimal and always round the calculated sample size up to the next whole number.
- TI-84 1-PropZInt workflow:
 - Open STAT → TESTS → 1-PropZInt. Enter the number of successes x , sample size n , and confidence level, then calculate.
 - Read the interval and $\hat{p}$, then write a contextual interpretation naming the population and the characteristic.

Example 1. In a sample of 250 adults, 142 support a proposal. Find $\hat{p}$.

Answer: $\hat{p} = \frac{142}{250} = 0.568$.

Example 2. A sample has $x = 142$ supporters out of $n = 250$. Use TI-84 1-PropZInt for a 95% confidence interval and interpret the output.

Answer: Use 1-PropZInt with $x = 142$, $n = 250$, and C-Level=.95. We are 95% confident the population proportion is between about 0.507 and 0.629.

Input	$x = 142$	$n = 250$	C-Level = 0.95
Output	$(0.507, 0.629), \hat{p} = 0.568$		

Example 3. A 95% confidence interval is $(0.51, 0.63)$. Interpret it in context.

Answer: We are 95% confident the population proportion is between 0.51 and 0.63.

Example 4. Is it correct to say there is a 95% probability the fixed population proportion is in this computed interval?

Answer: No. The population proportion is fixed, so after the interval is computed, the proportion either is or is not inside it. The randomness came from selecting the sample and constructing the interval. A correct statement is, 'We are 95% confident that the interval captures the population proportion.'

Example 5. Start to finish: In a random sample of 420 voters, 231 favor a proposal. Build and interpret a 95% confidence interval for the population proportion.

Answer: Conditions: random sample; $\hat{p} = 231/420 = 0.55$, so $n\hat{p}(1 - \hat{p}) = 420(0.55)(0.45) = 103.95$, which is at least 10. Use 1-PropZInt with $x = 231$, $n = 420$, and C-Level = .95. Output is about (0.502, 0.598). We are 95% confident the population proportion of voters who favor the proposal is between 0.502 and 0.598.

Conditions	Random sample; $n\hat{p}(1 - \hat{p}) = 420(0.55)(0.45) = 103.95 \geq 10$
TI-84 setup	1-PropZInt: $x = 231$, $n = 420$, C-Level = 0.95
Output	(0.502, 0.598), $\hat{p} = 0.55$
Conclusion	We are 95% confident the population proportion is between 0.502 and 0.598.

Example 6. A confidence interval is (0.141, 0.679). Find the point estimate and margin of error.

Answer: The midpoint is $\hat{p} = (0.141 + 0.679)/2 = 0.410$. The margin of error is $E = (0.679 - 0.141)/2 = 0.269$.

Example 7. With the same sample size and sample proportion, order 82%, 88%, 94%, and 99% confidence intervals from narrowest to widest.

Answer: 82%, 88%, 94%, 99%. Higher confidence requires a wider interval.

Example 8. If 100 different 91% confidence intervals are constructed under the model conditions, about how many should capture the population proportion? Must exactly that many capture it?

Answer: About 91 intervals should capture the parameter in the long run, but the actual count from one set of 100 intervals does not have to be exactly 91.

Example 9. Find the critical value z^* for a 98% confidence interval.

Answer: Here $\alpha = 0.02$, so each tail has area 0.01. Use invNorm(0.99,0,1) to get $z^* = z_{0.01} \approx 2.326$.

Example 10. Find the minimum sample size for a 95% confidence interval with margin of error 0.03 (a) using a prior estimate of 0.40 and (b) without a prior estimate.

Answer: Use $z^* = 1.96$. (a) $n = 0.40(0.60)(1.96/0.03)^2 \approx 1024.43$, so sample at least 1025 people. (b) $n = 0.25(1.96/0.03)^2 \approx 1067.11$, so sample at least 1068 people. Always round a required sample size up.

Section 9.2 Estimating a Population Mean

Use one-sample t-intervals and interpret them correctly.

QUICK REFERENCE

- The sample mean $\bar{x}$ is the point estimate for a population mean μ .
 - A mean confidence interval has the form $\bar{x} \pm E$. Its midpoint is $\bar{x}$, and its margin of error is half the interval width.
 - Use a one-sample t-interval when the population standard deviation σ is unknown and is estimated with the sample standard deviation s .

- The t-distribution accounts for the additional uncertainty from estimating σ .
 - Degrees of freedom tell how many sample values can vary freely after using the sample to estimate its mean. For one sample, $df = n - 1$: once the mean and $n - 1$ values are fixed, the last value is forced to make the deviations balance around the mean.
 - Degrees of freedom select the appropriate t-distribution and critical value. Smaller df means more uncertainty, heavier tails, and a larger critical value.
 - The t-distribution is symmetric and bell-shaped. As degrees of freedom increase, it approaches the standard normal distribution.
- The critical value $t^* = t_{\alpha/2}$ depends on the confidence level and degrees of freedom.
 - For a $(1 - \alpha)100\%$ interval, use the t-table or $\text{invT}(1-\alpha/2, df)$ with $df = n - 1$.
 - The interval extends t^* standard errors below and above the sample mean.
- Check conditions for a one-sample t-interval.
 - Use a random or representative sample and check $n \leq 0.05N$ when sampling without replacement.
 - For a small sample, the population should be approximately normal with no strong skewness or outliers. Larger samples make t-procedures more robust, but severe outliers still matter.
- For formula context, $\bar{x} \pm t^* \frac{s}{\sqrt{n}}$.
 - Higher confidence increases t^* and the margin of error. Increasing n decreases the standard error and margin of error.
 - Interpret the interval as a plausible range for the population mean in the original units, not as a range containing a stated percentage of individual observations.
- Use confidence wording because the population mean is fixed while the interval varies from sample to sample.
 - In repeated sampling, a $C\%$ confidence-interval method captures the fixed population mean in about $C\%$ of the intervals it produces.
 - After one interval is calculated, the mean either is or is not inside it. Say, 'We are $C\%$ confident that the interval captures the population mean,' not that there is a $C\%$ probability the fixed mean is inside.
- For planning, estimate the sample size with $n = (z^* \sigma / E)^2$.
 - Use a known or reasonable planning value for σ ; a standard deviation from an earlier study may be used when the population standard deviation is unavailable.
 - Use the z critical value because the t critical value depends on the sample size being found, and always round the result up.
- TI-84 TInterval workflow:
 - Choose Data when raw observations are in a list, or Stats when $\bar{x}$, s , and n are given. Enter the confidence level and calculate.
 - Read the interval and degrees of freedom, then state: 'We are $C\%$ confident that the population mean [context] is between [lower] and [upper] [units].'

Example 1. A sample has $n = 18$. What are the degrees of freedom for a one-sample t-interval?

Answer: $df = n - 1 = 18 - 1 = 17$. After the sample mean is fixed, 17 values can vary freely; the final value must balance the others around that mean. These 17 degrees of freedom determine which t-distribution and critical value to use.

Example 2. Use TI-84 TInterval with summary statistics $n = 18$, $\bar{x} = 14.6$, $s = 3.8$, and confidence level 90%. Interpret the interval.

Answer: Use TInterval with Stats input. The interval is approximately (13.0, 16.2). We are 90% confident the population mean is between 13.0 and 16.2.

Input	$\bar{x} = 14.6$	$Sx=3.8$	$n = 18$	C-Level = 0.90
Output	(13.0, 16.2)			

Example 3. A 90% confidence interval for mean wait time is (12.4, 16.8) minutes. Interpret it.

Answer: We are 90% confident the population mean wait time is between 12.4 and 16.8 minutes.

Example 4. Why should outliers be checked for a small-sample t-interval?

Answer: Outliers can pull the sample mean and sample standard deviation away from the typical values. In a small sample, one unusual observation can have a large effect, so the resulting t-interval may be shifted or unnecessarily wide.

Example 5. Start to finish: A random sample of 24 wait times has $\bar{x} = 18.4$ minutes and $s = 4.2$ minutes. A boxplot shows no strong skew or outliers. Build and interpret a 95% confidence interval for the population mean wait time.

Answer: Conditions: random sample; small sample, but no strong skew or outliers. Use TInterval with Stats input: $\bar{x} = 18.4$, $Sx=4.2$, $n = 24$, C-Level = 0.95. Output is about (16.6, 20.2). We are 95% confident the population mean wait time is between 16.6 and 20.2 minutes.

Conditions	Random sample; no strong skew or outliers for the small sample.
TI-84 setup	TInterval Stats: $\bar{x} = 18.4$, $Sx=4.2$, $n = 24$, C-Level = 0.95
Output	(16.6, 20.2)
Conclusion	We are 95% confident the population mean wait time is between 16.6 and 20.2 minutes.

Example 6. A mean confidence interval is (21, 25). Find the point estimate and margin of error.

Answer: The point estimate is the midpoint, $\bar{x} = 23$. The margin of error is half the width, $E = 2$.

Example 7. Two 96% t-intervals use the same sample mean and standard deviation, but one has $n = 13$ and the other has $n = 26$. Which is narrower?

Answer: The interval with $n = 26$ is narrower because its standard error is smaller.

Example 8. A 95% confidence interval for mean reading time is 12.4 to 16.8 minutes. Is it correct to say 95% of individual reading times lie in this interval?

Answer: No. The interval estimates the population mean reading time; it is not an interval containing 95% of individual observations.

Example 9. Find the critical value t^* for a 90% confidence interval based on a sample of 13 observations.

Answer: The degrees of freedom are $df = 13 - 1 = 12$, and $\alpha/2 = 0.05$. Use a t-table or invT(0.95,12) to get $t^* \approx 1.782$.

Example 10. An earlier study found a standard deviation of 12 minutes. Estimate the minimum sample size needed for a 95% confidence interval for the population mean with margin of error 2 minutes.

Answer: Use $z^* = 1.96$ and the earlier standard deviation as the planning estimate: $n = (1.96(12)/2)^2 = 138.30$. Round up, so at least 139 observations are needed.

Module 10: Hypothesis Testing

Students set up and run hypothesis tests for population proportions and means using correct statistical wording.

Section 10.1 The Language of Hypothesis Testing

Use hypothesis-test vocabulary and write correct conclusions.

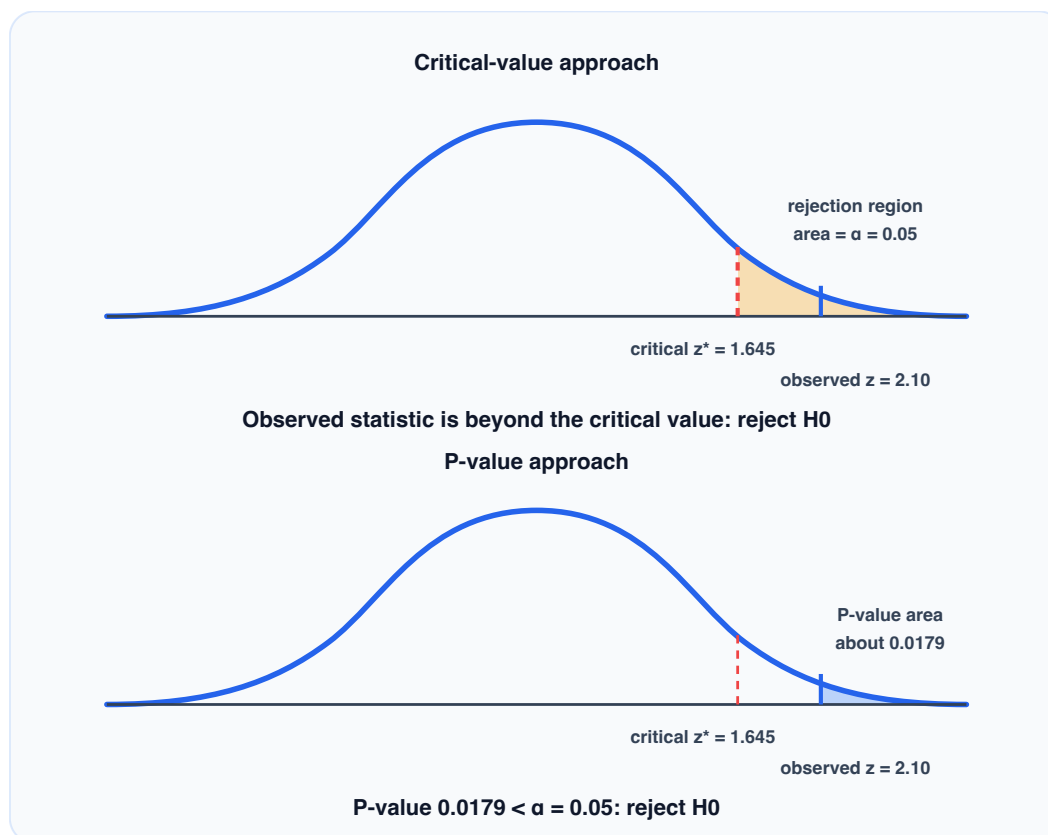
QUICK REFERENCE

- A hypothesis test evaluates a claim about a population parameter using sample evidence.
 - Write hypotheses with parameters such as p , μ , or σ , not sample statistics such as $\hat{p}$, $\bar{x}$, or s .
 - The null hypothesis H_0 contains equality and represents the starting status or claimed value. The alternative H_a states the change or difference the evidence is meant to support.
- The alternative hypothesis determines the test direction.
 - $<$ gives a left-tailed test, $>$ gives a right-tailed test, and $\neq$ gives a two-tailed test. Left- and right-tailed tests are also called one-tailed tests.
 - Words such as decreased or less suggest left-tailed; increased, more, or greater suggest right-tailed; changed or different suggest two-tailed.
- The P-value measures evidence against H_0 .
 - It is the probability, assuming H_0 is true, of obtaining a sample result at least as extreme in the direction of H_a as the observed result.
 - Smaller P-values give stronger evidence against H_0 . A P-value is not the probability that H_0 is true.
- The P-value and critical-value approaches use the same evidence and make the same decision.
 - P-value approach: reject H_0 when P-value $\leq \alpha$; otherwise, fail to reject H_0 .
 - Critical-value approach: use α and the direction of H_a to mark the rejection region. Reject H_0 when the test statistic falls beyond the critical value in that region.
 - When the test statistic is beyond the critical value, its P-value area is no larger than α , so both approaches reject. Otherwise, both approaches fail to reject.
 - Fail to reject does not mean accept or prove H_0 ; the evidence may simply be insufficient.
 - A smaller α makes rejection harder and lowers the chance of a Type I error.

- A testing decision can be wrong.
 - Type I error: reject a true H_0 . Its probability is α . Describe the false conclusion in the problem's context.
 - Type II error: fail to reject a false H_0 . Describe the real change or difference the test failed to detect.
- Write the conclusion in terms of the original claim.
 - When rejecting H_0 , say there is sufficient evidence for the alternative claim.
 - When failing to reject H_0 , say there is not sufficient evidence for the alternative claim. If the original claim is the null, say the data do not provide sufficient evidence against that claim; do not say the claim was proven true.

Actual situation	Fail to reject H_0	Reject H_0
H_0 is true	Correct decision - no error. The test does not claim a change when none exists.	Type I error. The test claims a change or difference when none actually exists.
H_0 is false H_a is true	Type II error. The test misses a real change or difference.	Correct decision - no error. The test detects a real change or difference.

Critical-Value and P-Value Approaches



Example 1. A claim says the population proportion is greater than 0.40. Write the alternative hypothesis.

Answer: $H_a : p > 0.40$.

Example 2. If $H_a : \mu < 25$, what type of test is this?

Answer: Left-tailed test.

Example 3. If $P\text{-value} = 0.032$ and $\alpha = 0.05$, what is the decision?

Answer: Reject H_0 , because $0.032 < 0.05$.

Example 4. If the $P\text{-value}$ is large, should we say we accept H_0 ?

Answer: No. Say fail to reject H_0 .

Example 5. In a medical screening study, H_0 says the test has the advertised accuracy. What kind of error is rejecting H_0 when the advertised accuracy is actually true?

Answer: Type I error, because a true null hypothesis was rejected.

Example 6. A manager tests $H_0 : \mu = 7.9$ against $H_a : \mu < 7.9$. Describe a Type II error in context.

Answer: A Type II error would be failing to conclude that the population mean pressure has decreased when the true mean pressure is actually below 7.9 psi.

Example 7. Three years ago, the population standard deviation of monthly cell phone bills was \$48.79. A researcher tests $H_0 : \sigma = 48.79$ against $H_a : \sigma < 48.79$ and rejects H_0 . State the conclusion in context.

Answer: There is sufficient evidence to conclude that the population standard deviation of monthly cell phone bills is less than its level three years ago of \$48.79.

Example 8. If a Type I error would have severe consequences, which common significance level is the most conservative: 0.10, 0.05, or 0.01?

Answer: Use $\alpha = 0.01$. It gives the smallest probability of rejecting a true null hypothesis.

Section 10.2 Hypothesis Tests for a Population Proportion

Use one-proportion z-tests with TI-84 and state conclusions correctly.

QUICK REFERENCE

- Use a one-proportion z-test for a claim about one population proportion p .
 - Write $H_0 : p = p_0$. Choose $H_a : p < p_0$, $p > p_0$, or $p \neq p_0$ from the claim's wording.
 - Use the population parameter p in the hypotheses, not the observed sample proportion $\hat{p}$.
- Check the one-proportion test conditions using the null value.
 - Use a random or representative sample and check $n \leq 0.05N$ when sampling without replacement.
 - For this course's homework, check $np_0(1 - p_0) \geq 10$. Test conditions use the null proportion p_0 , unlike interval conditions, which use the observed sample proportion.
- The one-proportion test statistic compares the observed proportion with the null proportion.
 - For formula context, $z = \frac{\hat{p} - p_0}{\sqrt{p_0(1 - p_0)/n}}$.
 - The $P\text{-value}$ is the null-model tail area at least as extreme as the observed z , using the direction in H_a .

- Use either the P-value approach or the critical-value approach for a one-proportion z-test.
 - P-value approach: reject H_0 when the calculator's P-value is no larger than α .
 - Critical-value approach: for a left-tailed test use $\text{invNorm}(\alpha, 0, 1)$; for a right-tailed test use $\text{invNorm}(1-\alpha, 0, 1)$; for a two-tailed test use the positive and negative values from $\text{invNorm}(1-\alpha/2, 0, 1)$.
 - Reject when the observed z falls beyond the critical value in the rejection region. Both approaches must give the same decision.
- TI-84 1-PropZTest workflow:
 - Open STAT → TESTS → 1-PropZTest. Enter p_0 , successes x , sample size n , and the alternative symbol.
 - Read z , the P-value, and $\hat{p}$. Compare the P-value with α ; do not decide from whether $\hat{p}$ is merely above or below p_0 .
- Use the Draw option in 1-PropZTest to see the P-value graph.
 - After entering the 1-PropZTest information and selecting the alternative symbol, highlight Draw instead of Calculate and press ENTER.
 - The calculator draws the standard normal curve, shades the P-value area, and displays the test statistic and P-value. A less-than alternative shades the left tail, a greater-than alternative shades the right tail, and a not-equal alternative shades both tails.
 - The Draw screen illustrates the P-value approach. It does not mark the critical z-value, so find that separately with invNorm when the critical-value approach is required.
- State the conclusion with evidence wording and context.
 - Reject H_0 : 'There is sufficient evidence to conclude that [alternative claim about the population proportion].'
 - Fail to reject H_0 : 'There is not sufficient evidence to conclude that [alternative claim].'. Do not say the null is true or that no difference exists.
- Statistical significance does not guarantee practical importance.
 - A very large sample can make a small departure from p_0 statistically significant. Report the observed size of the difference and its context when practical importance matters.

Example 1. A claim says fewer than 30% of students work full time. Write H_0 and H_a .

Answer: $H_0 : p = 0.30$; $H_a : p < 0.30$.

Example 2. For $H_0 : p = 0.30$, $n = 120$, check the large-count conditions.

Answer: $np_0(1 - p_0) = 120(0.30)(0.70) = 25.2$, so the value is at least 10.

Example 3. A sample of 120 students has 27 who work full time. Test the claim $p < 0.30$ at $\alpha = 0.05$ using TI-84 1-PropZTest.

Answer: Use $p_0 = 0.30$, $x = 27$, $n = 120$, and $p < p_0$. The output gives $z \approx -1.79$ and P -value ≈ 0.037 . P -value approach: $0.037 < 0.05$, so reject H_0 . Critical-value approach: $\text{invNorm}(0.05, 0, 1)$ gives $z^* \approx -1.645$, and $-1.79 < -1.645$, so the test statistic is in the left rejection region. Both approaches reject H_0 . There is sufficient evidence to support the claim that the population proportion is less than 0.30.

Input	$p_0 = 0.30$	$x = 27$	$n = 120$	$\text{prop} < p_0$
Output	$z \approx -1.79$	$p \approx 0.037$	$\hat{p} = 0.225$	

Example 4. A test gives P -value = 0.081 with $\alpha = 0.05$. What is the decision?

Answer: Fail to reject H_0 .

Example 5. State the conclusion for a failed rejection when the claim was $p < 0.30$.

Answer: There is not sufficient evidence to support the claim that the population proportion is less than 0.30.

Example 6. Start to finish: A college claims fewer than 40% of students buy a printed textbook. In a random sample of 180 students, 62 bought a printed textbook. Test the claim at $\alpha = 0.05$.

Answer: Hypotheses: $H_0 : p = 0.40$, $H_a : p < 0.40$. Conditions: $np_0(1 - p_0) = 180(0.40)(0.60) = 43.2$, which is at least 10. Use 1-PropZTest with $p_0 = 0.40$, $x = 62$, $n = 180$, and $p < p_0$. Since P -value $\approx 0.064 > 0.05$, fail to reject H_0 . There is not sufficient evidence to support the claim that fewer than 40% of students buy a printed textbook.

Hypotheses	$H_0 : p = 0.40$; $H_a : p < 0.40$
Conditions	$np_0(1 - p_0) = 180(0.40)(0.60) = 43.2 \geq 10$
TI-84 setup	1-PropZTest: $p_0 = 0.40$, $x = 62$, $n = 180$, $\text{prop} < p_0$
Output	$z \approx -1.52$, P -value ≈ 0.064
Conclusion	Fail to reject H_0 ; not sufficient evidence to support the claim.

Example 7. Order these P -values from weakest to strongest evidence against H_0 : 0.048, 0.013, 0.008, 0.036.

Answer: 0.048, 0.036, 0.013, 0.008. Smaller P -values provide stronger evidence against the null hypothesis.

Example 8. A test of $H_0 : p = 0.60$ gives P -value 0.18. Interpret this P -value without stating a decision.

Answer: If the population proportion were truly 0.60, results at least as extreme as the observed sample result in the direction of the alternative would occur in about 18% of random samples.

Example 9. A very large study finds $\hat{p} = 0.503$ significantly different from $p_0 = 0.500$. What additional question should be asked?

Answer: Ask whether a difference of 0.003, or 0.3 percentage point, is practically meaningful in the study's context. Statistical significance alone does not answer that question.

Example 10. Return to the left-tailed one-proportion z-test in Example 3. Use the Draw option in TI-84 1-PropZTest. Describe the graph, report the displayed test statistic and P-value, and state the decision at $\alpha = 0.05$. How does this agree with the critical-value approach?

Answer: Open STAT → TESTS → 1-PropZTest. Enter $p_0 = 0.30$, $x = 27$, and $n = 120$; select $p < p_0$; then highlight Draw and press ENTER. The calculator draws a standard normal curve with the area to the left of $z \approx -1.79$ shaded and displays a P-value of approximately 0.037. Because $0.037 < 0.05$, reject H_0 . The Draw screen does not mark the critical value, but the decision agrees with the critical-value approach because $z = -1.79$ is left of the critical value -1.645 .

Section 10.3 Hypothesis Tests for a Population Mean

Use one-sample t-tests with TI-84 and state conclusions correctly.

QUICK REFERENCE

- Use a one-sample t-test for a claim about one population mean when σ is unknown.
 - Write $H_0 : \mu = \mu_0$. Choose a left-, right-, or two-tailed alternative from the claim.
 - Use the population mean μ in the hypotheses, not the sample mean $\bar{x}$. Degrees of freedom are $df = n - 1$.
- Check the one-sample t-test conditions.
 - Use a random or representative sample and check $n \leq 0.05N$ when sampling without replacement.
 - For a small sample, verify approximate normality and no strong outliers with the data, histogram, boxplot, or normal probability plot. Larger samples are more robust, but extreme outliers can still distort the test.
- The t statistic measures the sample mean's distance from the null mean in standard-error units.
 - For formula context, $t = \frac{\bar{x} - \mu_0}{s/\sqrt{n}}$.
 - The P-value is the t-distribution tail area at least as extreme as the observed statistic, using $df = n - 1$ and the direction in H_a .
- Tail direction controls the critical value and the P-value area.
 - For a right-tailed test use t_α ; for a left-tailed test use $-t_\alpha$; for a two-tailed test use $\pm t_{\alpha/2}$, all with $df = n - 1$.
 - Because invT uses left-tail area, enter invT(1-alpha,df) for a right critical value, invT(alpha,df) for a left critical value, or invT(1-alpha/2,df) for the positive two-tailed critical value.
 - Shade the P-value to the left for $H_a : \mu < \mu_0$, to the right for $H_a : \mu > \mu_0$, and in both tails for $H_a : \mu \neq \mu_0$.
 - The P-value approach rejects when P-value $\leq \alpha$. The critical-value approach rejects when t falls beyond the critical value. Both approaches must agree.
- TI-84 T-Test workflow:
 - Choose Data when raw observations are stored in a list, or Stats when $\bar{x}$, s , and n are given. Enter μ_0 and select the correct alternative.
 - Read t , the P-value, $\bar{x}$, and n . Compare the P-value with α .

- Use the Draw option in T-Test to see the P-value graph.
 - After entering the T-Test information and selecting the alternative symbol, highlight Draw instead of Calculate and press ENTER.
 - The calculator draws the t distribution, shades the P-value area, and displays the test statistic and P-value. A less-than alternative shades the left tail, a greater-than alternative shades the right tail, and a not-equal alternative shades both tails.
 - The Draw screen illustrates the P-value approach. It does not mark the critical t-value, so find that separately with invT when the critical-value approach is required.
- Finish with a decision and contextual conclusion.
 - Reject H_0 : there is sufficient evidence for the alternative claim about the population mean.
 - Fail to reject H_0 : there is not sufficient evidence for the alternative claim. Do not say the population mean equals μ_0 or that H_0 was accepted.
- Separate statistical significance from practical significance.
 - A small P-value says the sample result would be unusual under H_0 ; it does not say the difference $\bar{x} - \mu_0$ is large enough to matter in practice.
 - Report the estimated difference and its units when judging real-world importance.

Example 1. A claim says the population mean wait time is different from 15 minutes. Write H_0 and H_a .

Answer: $H_0 : \mu = 15$; $H_a : \mu \neq 15$.

Example 2. If $H_a : \mu > 42$, what type of test is this?

Answer: Right-tailed test.

Example 3. A sample has $n = 16$, $\bar{x} = 43.8$, and $s = 3.2$. Test the claim $\mu > 42$ at $\alpha = 0.05$ using TI-84 T-Test.

Answer: Use Stats input with $\mu_0 = 42$, $\bar{x} = 43.8$, $Sx=3.2$, $n = 16$, and $\mu > \mu_0$. The output gives $t = 2.25$, $P\text{-value} \approx 0.020$, and $df = 15$. P-value approach: $0.020 < 0.05$, so reject H_0 . Critical-value approach: invT(0.95,15) gives $t^* \approx 1.753$, and $2.25 > 1.753$, so the test statistic is in the right rejection region. Both approaches reject H_0 . There is sufficient evidence to support the claim that the population mean is greater than 42.

Input	$\mu_0 = 42$	$\bar{x} = 43.8$	$Sx=3.2$	$n = 16$
Output	$t = 2.25$	$p \approx 0.020$	$df = 15$	

Example 4. A t-test gives $P\text{-value} = 0.014$ with $\alpha = 0.05$. What is the decision?

Answer: Reject H_0 .

Example 5. State the conclusion for rejecting H_0 when the claim was $\mu > 42$.

Answer: There is sufficient evidence to support the claim that the population mean is greater than 42.

Example 6. Start to finish: A tutoring center claims the mean improvement is more than 5 points. A random sample of 20 students has $\bar{x} = 6.3$, $s = 2.4$, and no strong skew or outliers. Test the claim at $\alpha = 0.05$.

Answer: Hypotheses: $H_0 : \mu = 5$, $H_a : \mu > 5$. Conditions: random sample; small sample with no strong skew or outliers. Use T-Test Stats with $\mu_0 = 5$, $\bar{x} = 6.3$, $Sx=2.4$, $n = 20$, and $\mu > \mu_0$. Since $P\text{-value} \approx 0.013 < 0.05$, reject H_0 . There is sufficient evidence to support the claim that the population mean improvement is more than 5 points.

Hypotheses	$H_0 : \mu = 5; H_a : \mu > 5$
Conditions	Random sample; no strong skew or outliers for the small sample.
TI-84 setup	T-Test Stats: $\mu_0 = 5$, $\bar{x} = 6.3$, $Sx=2.4$, $n = 20$, $\mu > \mu_0$
Output	$t \approx 2.42$, $P\text{-value} \approx 0.013$, $df = 19$
Conclusion	Reject H_0 ; sufficient evidence supports the claim.

Example 7. A left-tailed t-test of $H_0 : \mu = 20$ gives $t = -1.90$ and P-value 0.036. Interpret the P-value and decide at $\alpha = 0.05$.

Answer: If the population mean were 20, a sample mean producing a t statistic of -1.90 or lower would occur about 3.6% of the time. Because $0.036 < 0.05$, reject H_0 ; there is sufficient evidence that the population mean is below 20.

Example 8. A two-tailed mean test at $\alpha = 0.01$ has a matching 99% confidence interval that contains μ_0 . What decision should the test make?

Answer: Fail to reject H_0 . For a two-tailed test, containing the null mean in the matching confidence interval agrees with a P-value greater than 0.01.

Example 9. A large-sample study finds a statistically significant mean increase of 0.02 unit. What should be considered before calling the result important?

Answer: Consider whether an increase of 0.02 unit is large enough to matter in the real setting, along with its units, costs, benefits, and consequences. Statistical significance alone does not establish practical significance.

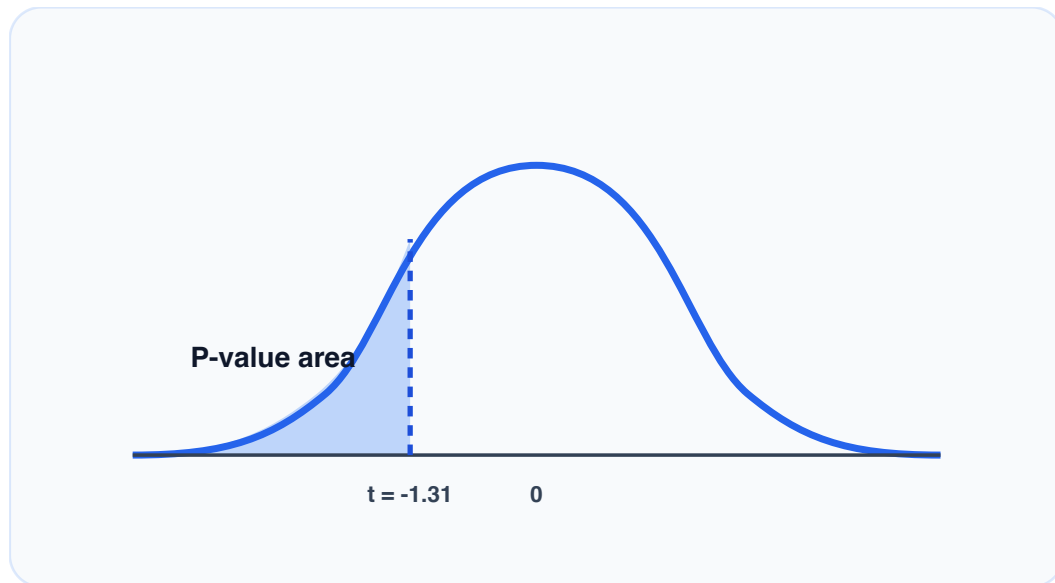
Example 10. Find the critical value or values for each t-test: (a) right-tailed with $\alpha = 0.05$ and $df = 15$; (b) left-tailed with $\alpha = 0.10$ and $n = 10$; (c) two-tailed with $\alpha = 0.01$ and $n = 16$.

Answer: (a) $\text{invT}(0.95,15)$ gives $t \approx 1.753$. (b) $df = 9$, and $\text{invT}(0.10,9)$ gives $t \approx -1.383$. (c) $df = 15$; use $\text{invT}(0.995,15)$ and both signs to get critical values $t \approx \pm 2.947$.

Example 11. Start to finish: A random sample of 17 completion times has mean 19.6 minutes and standard deviation 4.4 minutes, with no strong skew or outliers. Test $H_0 : \mu = 21$ against $H_a : \mu < 21$ at $\alpha = 0.01$.

Answer: Conditions: random sample and no strong skew or outliers for the small sample. On the TI-84, use T-Test with Stats input: $\mu_0 = 21$, $\bar{x} = 19.6$, $Sx=4.4$, $n = 17$, and $\mu < \mu_0$. The output gives $t \approx -1.31$, $df = 16$, and a P-value of approximately 0.104. Because $0.104 > 0.01$, fail to reject H_0 . There is not sufficient evidence to conclude that the population mean completion time is less than 21 minutes.

Left-Tailed P-Value



Example 12. Return to the right-tailed one-sample t-test in Example 3. Use the Draw option in TI-84 T-Test. Describe the graph, report the displayed test statistic and P-value, and state the decision at $\alpha = 0.05$. How does this agree with the critical-value approach?

Answer: Open STAT → TESTS → T-Test and choose Stats. Enter $\mu_0 = 42$, $\bar{x} = 43.8$, $Sx=3.2$, and $n = 16$; select $\mu > \mu_0$; then highlight Draw and press ENTER. The calculator draws a t distribution with the area to the right of $t = 2.25$ shaded and displays a P-value of approximately 0.020. Because $0.020 < 0.05$, reject H_0 . The Draw screen does not mark the critical value, but the decision agrees with the critical-value approach because $t = 2.25$ is right of the critical value 1.753.