

the placebo group participants. If they were told that they were getting a placebo, there would be no expectation for getting better and hence no (or a much reduced) placebo effect. But why should the experimenters not know the group assignments of the participants? This is to control for the effects of experimenter expectation (Rosenthal, 1966, 1994). If the experimenter knew which condition the participants were in, she might unintentionally treat them differently and thereby affect their behavior. In addition, the experimenter might interpret and record the behavior of the participants differently if she needed to make judgments about their behavior (their anxiety level in the example study). The key for the participant assignments to groups is kept by a third party and then given to the experimenter once the study has been conducted.

Now let's think about experiments that are more complex than our sample experiment with its one independent variable (aerobic exercise) and two control groups. In most experiments, the researcher examines multiple values of the independent variable. With respect to the aerobic exercise variable, the experimenter might examine the effects of different amounts or different types of aerobic exercise. Such manipulations would provide more detailed information about the effects of aerobic exercise on anxiety. An experimenter might also manipulate more than one independent variable. For example, an experimenter might manipulate diet as well as aerobic exercise. Different diets (high-protein vs. high-carbohydrate diets) might affect a person's level of anxiety in different ways. The two independent variables (diet and aerobic exercise) might also interact to determine one's level of anxiety. An experimenter could also increase the number of dependent variables. For example, both level of anxiety and level of depression could be measured in our sample experiment, even if only aerobic exercise were manipulated. If aerobic exercise reduces anxiety, then it might also reduce depression. As an experimenter increases the number of values of an independent variable, the number of independent variables, or the number of dependent variables, the possible gain in knowledge about the relationship between the variables also increases. Thus, most experiments are more complex in design than our simple example with an experimental group and two control groups. In addition, many experimental studies, including replications, must be conducted to address any experimental question. Researchers have a statistical technique called **meta-analysis** that combines the results for a large number of studies on one experimental question into one analysis to arrive at an overall conclusion. Because a meta-analysis involves the results of numerous experimental studies, its conclusion is considered much stronger evidence than the results of an individual study in answering an experimental question.

The various research methods that have been discussed are summarized in Table 1.3 (page 26). Their purposes and data-gathering procedures are described. Make sure you understand each of these research methods before going on to the next section, where we'll discuss how to understand research results. If you feel that you don't understand a particular method, go back and reread the information about it until you do.

meta-analysis A statistical technique that combines the results of a large number of studies on one experimental question into one analysis to arrive at an overall conclusion.

Table 1.3 Summary of Research Methods

Research Method	Goal of Method	How Data Are Collected
Laboratory observation	Description	Unobtrusively observe behavior in a laboratory setting
Naturalistic observation	Description	Unobtrusively observe behavior in its natural setting
Participant observation	Description	Observer becomes part of the group whose behavior is being observed
Case study	Description	Study an individual in depth over an extended period of time
Survey	Description	Representative sample of a group completes questionnaires or interviews to determine behavior, beliefs, and attitudes of the group
Correlational study	Prediction	Measure two variables to determine whether they are related
Experiment	Explanation	Manipulate one or more independent variables in a controlled setting to determine their impact on one or more measured dependent variables

Section Summary

Research methods fall into three categories—descriptive, correlational, and experimental. There are three descriptive methods—observation, case studies, and surveys. Observational studies can be conducted in the laboratory or in a natural setting (naturalistic observation). Sometimes participant observation is used. In participant observation, the observer becomes a part of the group being observed. The main goal of all is an in-depth study of one individual. Hypotheses generated from case studies in a clinical setting have often led to important experimental findings. Surveys attempt to describe the behavior, attitudes, or beliefs of particular populations (groups of people). It is essential in conducting surveys to ensure that a representative sample of the population is obtained for the study. Random sampling in which each person in the population has an equal opportunity to be in the sample is used for this purpose.

Descriptive methods only allow description, but correlational studies allow the researcher to make predictions about the relationships between variables. In a correlational study, two variables are measured and these measurements are compared to see if they are related. A statistic, the correlation coefficient, tells us both the type of the relationship (positive or negative) and the strength of the relationship. The sign of the coefficient (+ or -) tells us the type, and the absolute value of the coefficient (0 to 1.0) tells us the strength. Zero and values near zero indicate no relationship. As the absolute value approaches 1.0, the strength increases. Correlational data may also be

depicted in scatterplots. A positive correlation is indicated by data points that extend from the bottom left of the plot to the top right. Scattered data points going from the top left to the bottom right indicate a negative correlation. The strength is reflected in the amount of scatter—the more the scatter, the lower the strength. When two variables are not perfectly correlated, regression toward the mean will be observed. This means that extreme values on one of the variables tend to be matched on average with less extreme values on the other variable. Thus, when trying to explain or predict changes in scores, levels of performance, and other events over time, remember regression toward the mean is likely operating so extreme or extraordinary events will tend to be followed by more normal ones. A correlation of 1.0 gives us perfect predictability about the two variables involved, but it does not allow us to make cause-effect statements about the variables. This is because “third variables” may be responsible for the observed relationship, making the correlation a spurious one.

To draw cause-effect conclusions, the researcher must conduct well-controlled experiments. In a simple experiment, the researcher manipulates the independent variable (the hypothesized cause) and measures its effect upon the dependent variable (the variable hypothesized to be affected). These variables are operationally defined so that other researchers understand exactly how they were manipulated or measured. In more complex experiments, more than one independent variable is manipulated or more than one dependent variable is measured. The experiment is conducted in a controlled environment in which possible third variables are held constant; the individual characteristics of participants are controlled through random assignment of participants to groups or conditions. Other controls employed in experiments include using a control group that is not exposed to the experimental manipulation, a placebo group that receives a placebo to control for the placebo effect, and the double-blind procedure that controls for the effects of experimenter and participant expectation. The researcher uses inferential statistics to interpret the results of an experiment. These statistics determine the probability that the results are due to chance. For the results to be statistically significant, this probability has to be .05 or less, but this significance standard may be too lenient. In addition, statistically significant results may or may not have practical significance or value in our everyday world. Because researchers have to conduct many studies, including replications, to address an experimental question, meta-analysis, a statistical technique that combines the results of a large number of experiments on one experimental question into one analysis, can be used to arrive at an overall conclusion.

ConceptCheck | 2

Explain why the results of a case study cannot be generalized to a population.

Explain the differences between random sampling and random assignment.

Explain how the scatterplots for correlation coefficients of $+ .90$ and $- .90$ would differ.

Waldman, Nicholson, Adilov, and Williams (2008) found that autism prevalence rates among school-aged children were positively correlated with annual precipitation levels in various states. Autism rates were higher in counties with higher precipitation levels. Try to identify some possible third variables that might be responsible for this correlation.

Explain why a double-blind procedure is necessary in an experiment in which there is a placebo group.

How to Understand Research Results

Once we have completed an experiment, we need to understand our results and to describe them concisely so that others can understand them. To do this, we need to use statistics. There are two types of statistics—descriptive and inferential. We described inferential statistics when we discussed how to interpret the results of experimental studies. In this section, we will be discussing **descriptive statistics**—the statistics used to describe the data of a research study in a concise fashion. The correlation coefficient that we discussed earlier is a descriptive statistic that allows us to describe the results of a correlational study precisely. For experimental findings, we need two types of descriptive statistics to summarize our data—measures of central tendency and measures of variability. In addition, a researcher often constructs a frequency distribution for the data. A **frequency distribution** depicts, in a table or a graph, the number of participants receiving each score for a variable. The bell curve, or normal distribution, is the most famous frequency distribution. We begin with the two types of descriptive statistics necessary to describe a data set: measures of central tendency and measures of variability.

Descriptive Statistics

In an experiment, the data set consists of the measured scores on the dependent variable for the sample of participants. A listing of this set of scores, or any set of numbers, is referred to as a distribution of scores, or a distribution of numbers. To describe such distributions in a concise summary manner, we use two types of descriptive statistics: measures of central tendency and measures of variability.

Measures of central tendency. The measures of central tendency define a “typical” score for a distribution of scores. There are three measures of central tendency (three ways to define the “typical” score)—mean, median, and mode. The first, the mean, is one that you are already familiar with from our earlier discussion of the regression toward the mean phenomenon. Remember, the mean is the numerical average for a distribution of scores. To compute the mean, you

descriptive statistics Statistics that describe the results of a research study in a concise fashion.

frequency distribution A depiction, in a table or figure, of the number of participants (frequency) receiving each score for a variable.

median The score positioned in the middle of a distribution of scores when all of the scores are arranged from lowest to highest.

mode The most frequently occurring score in a distribution of scores.

merely add up all of the scores and divide by the number of scores. A second measure of central tendency is the **median**—the score positioned in the middle of the distribution of scores when all of the scores are listed from the lowest to the highest. If there is an odd number of scores, the median is the middle score. If there is an even number of scores, the median is the halfway point between the two center scores. The final measure of central tendency, the **mode**, is the most frequently occurring score in a distribution of scores. Sometimes there are two or more scores that occur most frequently. In these cases, the distribution has multiple modes. Now let's consider a small set of scores to see how these measures are computed.

Let's imagine a class with five students who just took an exam. That gives us a distribution of five test scores: 70, 80, 80, 85, and 85. First, let's compute the mean or average score. The sum of all five scores is 400. Now divide 400 by 5, and you get the mean, 80. What's the median? It's the middle score when the scores are arranged in ascending order. Because there is an odd number of scores (5), it's the third score—80. If there had been an even number of scores, the median would be the halfway point between the center two scores. For example, if there had been only four scores in our sample distribution (70, 80, 85, and 85), the median would be the halfway point between 80 and 85, 82.5. Now, what's the mode or most frequently occurring score? For the distribution of five scores, there are two numbers that occur twice, so there are two modes—80 and 85. This kind of distribution is referred to as a bimodal distribution (a distribution with two modes). Remember that a distribution can have one or more than one mode.

Of the three measures of central tendency, the mean is the one that is most commonly used. This is mainly because it is used to analyze the data in many inferential statistical tests. The mean can be distorted, however, by a small set of unusually high or low scores. In this case, the median, which is not distorted by such scores, should be used. To understand how atypical scores can distort the mean, let's consider changing one score in our sample distribution of five scores. Change 70 to 20. Now, the mean is 70 ($350/5$). The median remains 80, however; it hasn't changed. This is because the median is only a positional score. The mean is distorted because it has to average in the value of any unusual scores.

Measures of variability. In addition to knowing the typical score for a distribution, you need to determine the variability between the scores. For example, two distributions might have the same mean, but one distribution might have little variability between scores and the other, considerable variability between scores. So how is such variability measured? There are two measures of variability—the range and the standard deviation. The range is the simpler to compute. The **range** is simply the difference between the highest and lowest scores in the distribution. For our sample distribution with five scores, it would be 85 minus 70, or 15. However, like the mean, unusually high or low scores distort the range. For example, if the 70 in the distribution had been a 20, the range would change to be 85 minus 20, or 65. This would not be a good measure of the distribution's variability because four of the five scores are 80 or 85, not very different.

The measure of variability used most often is the standard deviation. In general terms, the **standard deviation** is the average extent that the scores vary from the mean of the distribution. In other words, how spread out are the scores? If the scores do not vary much from the mean, the standard deviation will be small. If they vary a lot from the mean, the standard deviation will be larger. In our example of five test scores with a mean of 80, the scores (70, 80, 80, 85, and 85) didn't vary much from this mean, therefore the standard deviation would not be very large. However, if the scores had been 20, 40, 80, 120, and 140, the mean would still be 80;

range The difference between the highest and lowest scores in a distribution of scores.

standard deviation The average extent that the scores vary from the mean for a distribution of scores.

Table 1.4 Summary of Descriptive Statistics

Descriptive Statistic	Explanation of Statistic
Correlation coefficient	A number between -1.0 and $+1.0$ whose sign indicates the type ($+$ = positive and $-$ = negative) and whose absolute value (0 to 1.0) indicates the strength of the relationship between two variables
Mean	Numerical average for a distribution of scores
Median	Middle score in a distribution of scores when all scores are arranged in order from lowest to highest
Mode	Most frequently occurring score or scores in a distribution of scores
Range	Difference between highest and lowest scores in a distribution of scores
Standard deviation	Average extent to which the scores vary from the mean for a distribution of scores

but the scores vary more from the mean, therefore the standard deviation would be much larger.

The standard deviation and the various other descriptive statistics that we have discussed are summarized in Table 1.4. Review this table to make sure you understand each statistic. The standard deviation is especially relevant to the normal distribution, or bell curve. We will see in Chapter 6, on thinking and intelligence, that intelligence test scores are actually determined with respect to standard deviation units in the normal distribution. Next we will consider the normal distribution and the two types of skewed frequency distributions.

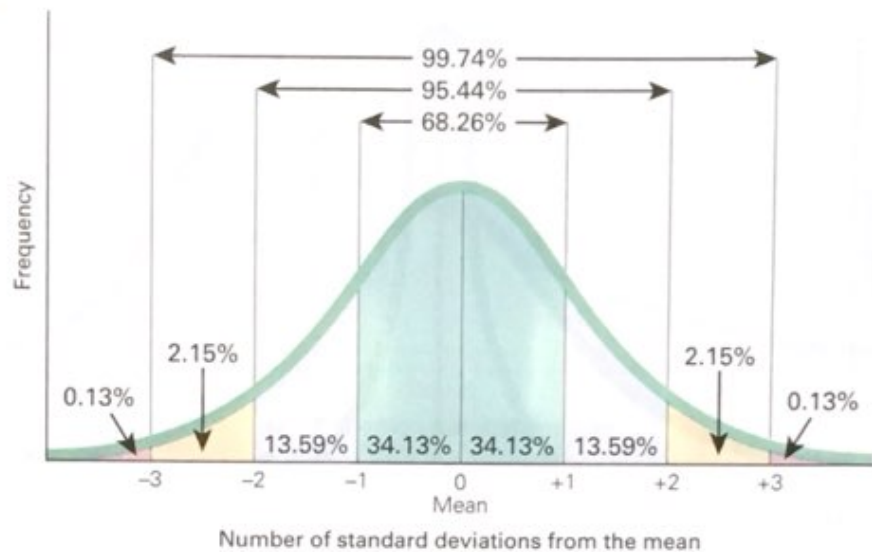
Frequency Distributions

A frequency distribution organizes the data in a score distribution so that we know the frequency of each score. It tells us how often each score occurred. These frequencies can be presented in a table or visually in a figure. We'll consider visual depictions. For many human traits (such as height, weight, and intelligence), the frequency distribution takes on the shape of a bell curve. For example, the heights of American adult men are distributed in a bell-shaped manner around a mean of 5 feet 10 inches (Wheeler, 2013). In fact, if a large number of people are measured on almost anything, the frequency distribution will visually approximate a bell-shaped curve. Statisticians call this bell-shaped frequency distribution, shown in Figure 1.3, the **normal distribution**.

normal distribution A frequency distribution that is shaped like a bell. About 68% of the scores fall within 1 standard deviation of the mean, about 95% within 2 standard deviations of the mean, and over 99% within 3 standard deviations of the mean.

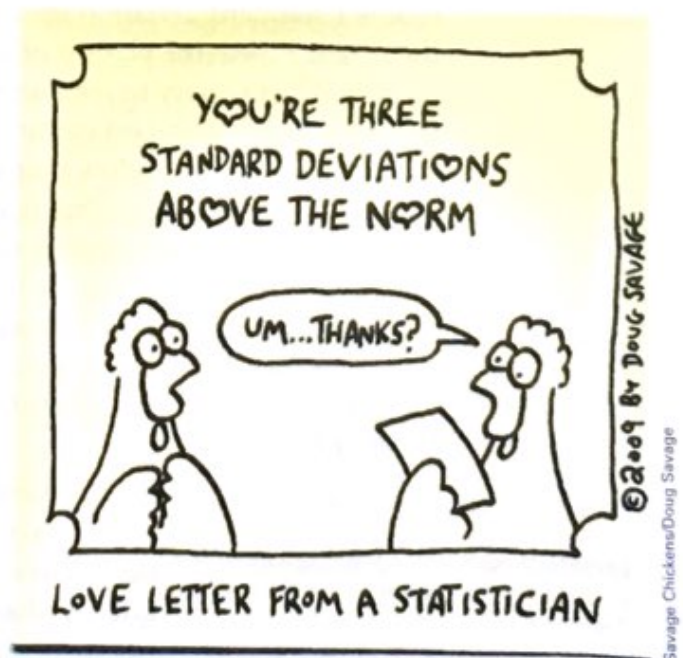
Figure 1.3 | The Normal Distribution

In a normal distribution, the mean, the median, and the mode are all equal because the distribution is perfectly symmetrical about its center. In addition, about 68% of the scores fall within 1 standard deviation of the mean, about 95% within 2 standard deviations of the mean, and over 99% within 3 standard deviations of the mean.



Normal distributions. There are two main aspects of a normal distribution. First, the mean, the median, and the mode are all equal because the normal distribution is symmetric about its center. You do not have to worry about which measure of central tendency to use because all of them are equal. The same number of scores fall below the center point as above it. Second, the percentage of scores falling within a certain number of standard deviations of the mean is set. About 68% of the scores fall within 1 standard deviation of the mean; about 95% within 2 standard deviations of the mean; and over 99% within 3 standard deviations of the mean. So what does this mean for the normal distribution of the heights of American adult men with a mean of 5 feet 10 inches? First, we have to know the standard deviation for this distribution. It is 3 inches. Thus, 68% of the heights of American men fall between 5 feet 7 inches (5 feet 10 inches - 3 inches) and 6 feet 1 inches (5 feet 10 inches + 3 inches), 95% between 5 feet 4 inches (5 feet 10 inches - 6 inches) and 6 feet 4 inches (5 feet 10 inches + 6 inches), and 99% between 5 feet 1 inch (5 feet 10 inches - 9 inches) and 6 feet 7 inches (5 feet 10 inches + 9 inches).

These percentages falling within a certain number of standard deviations are what give the normal distribution its bell shape. The percentages hold regardless of the size of the standard deviation for a normal distribution. Figure 1.4 (page 32) shows two normal distributions with the same mean but different



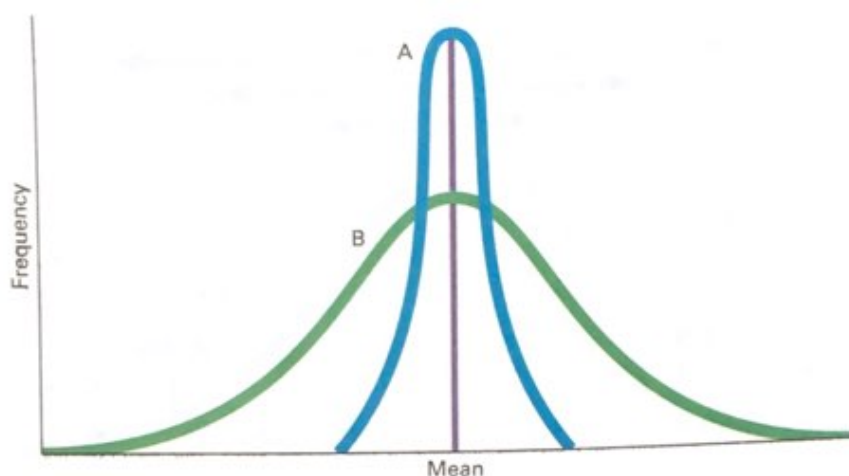


Figure 1.4 | Normal Distributions with Different Standard Deviations | These are normal distributions with the same mean but different standard deviations. Normal distribution A has a smaller standard deviation than normal distribution B. As the standard deviation for a normal distribution gets smaller, its bell shape gets narrower and taller.

standard deviations. Both have bell shapes, but the distribution with the smaller standard deviation (A) is taller. As the size of the standard deviation increases, the bell shape becomes shorter and wider (like B).

The percentages of scores and the number of standard deviations from the mean always have the same relationship in a normal distribution. This allows you to compute percentile ranks for scores. A **percentile rank** is the percentage of scores below a specific score in a distribution of scores. If you know how many standard deviation units a specific score is above or below the mean in a normal distribution, you can compute that score's percentile rank. For example, the percentile rank of a score that is 1 standard deviation above the mean is roughly 84%. Remember, a normal distribution is symmetric about the mean so that 50% of the scores are above the mean and 50% are below the mean. This means that the percentile rank of a score that is 1 standard deviation above the mean is greater than 50% (the percent below the mean) + 34% (the percent of scores from the mean to +1 standard deviation).

Now I'll let you try to compute a percentile rank. What is the percentile rank for a score that is 1 standard deviation below the mean? Remember that it is the percentage of the scores below that score. Look at Figure 1.3. What percentage of the scores is less than a score that is 1 standard deviation below the mean? The answer is about 16%. You can never have a percentile rank of 100% because you cannot outscore yourself, but you can have a percentile rank of 0% if you have the lowest score in the distribution. The scores on intelligence tests and the SAT are based on normal distributions, so percentile ranks can be calculated for these scores. We will return to the normal distribution when we discuss intelligence test scores in Chapter 6.

percentile rank The percentage of scores below a specific score in a distribution of scores.

Skewed distributions. In addition to the normal distribution, two other types of frequency distributions are important. They are called skewed distributions, which are frequency distributions that are asymmetric in shape. The two major types of skewed distributions are illustrated in

Figure 1.5 | Sample Right-Skewed and Left-Skewed Frequency Distributions

(a) This is an example of a right-skewed frequency distribution in which the tail of the distribution goes off to the right. In a right-skewed distribution, the mean is greater than the median because the unusually high scores distort it.

(b) This is an example of a left-skewed frequency distribution in which the tail of the distribution goes off to the left. The mean is less than the median because the unusually low scores distort it.

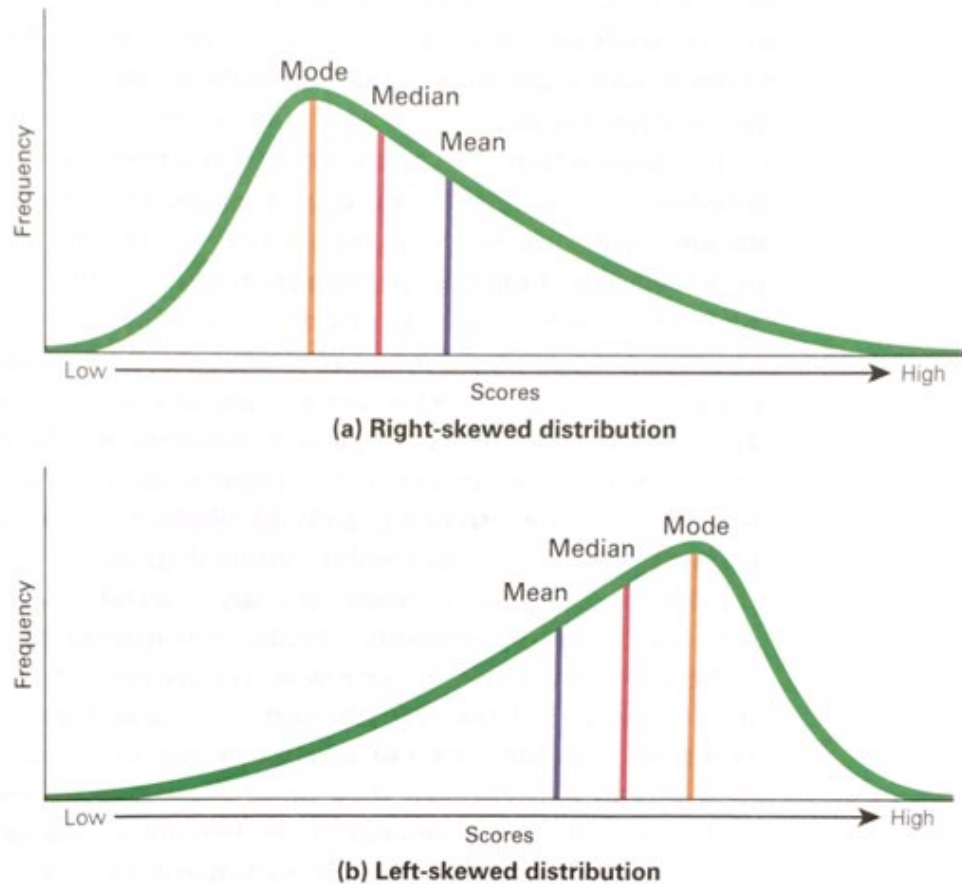


Figure 1.5. A **right-skewed distribution** is a frequency distribution in which there are some unusually high scores [shown in Figure 1.5(a)]. A **left-skewed distribution** is a frequency distribution in which there are some unusually low scores [shown in Figure 1.5(b)]. An easy way to remember the difference is that the tail of the right-skewed distribution goes off to the right, and the tail of the left-skewed distribution goes off to the left. A right-skewed distribution is also called a positively skewed distribution (the tail goes toward the positive end of the number line); a left-skewed distribution is also called a negatively skewed distribution (the tail goes toward the negative end of the number line).

Now that we have defined right-skewed and left-skewed distributions, let's consider some examples of each type of distribution so that we get a better understanding of these distributions. As you read these examples, visually think about what the distributions would look like. Remember, the tail of a right-skewed distribution goes to the right (the high end of the scale), and the tail of a left-skewed distribution goes to the left (the low end of the scale). Also, remember that the tails of these skewed distributions can be very long. Age at retirement is a good example of a left-skewed distribution.

right-skewed distribution An asymmetric frequency distribution in which there are some unusually high scores that distort the mean to be greater than the median.

left-skewed distribution An asymmetric frequency distribution in which there are some unusually low scores that distort the mean to be less than the median.

Most people retire in their mid to late 60s or early 70s, some retire in their 50s, and relatively few in their 40s or earlier. Another example would be scores on a relatively easy exam. Most students would get As or Bs (high scores), some would get Cs, a few Ds, and hardly any Fs (low scores).

The distribution of people's incomes is a good example of a right-skewed distribution. The incomes of most people tend to be on the lower end of possible incomes, but some people make a lot of money, with very high incomes increasingly rare. The "long tail" phenomenon in the digital business world (Anderson, 2008) also involves right-skewed distributions. The "long tail" referred to is that for the distribution depicting the sales of all the available items in a specific market (e.g., music tracks or book titles) as a function of their popularity rank. The small number of very popular items make up the head of the distribution, and the very large number of less popular items make up the long tail. Given their much lower stocking and distribution costs versus physical retailers, e-businesses can essentially sell the entire distribution; and the aggregate of all of the items in the long tail make up a large market for sales. Figure 1.6 shows the long tail for music downloads for online music retailer Rhapsody versus Walmart.

Because unusually high or low scores distort a mean, such distortion occurs for the means of skewed distributions. The mean for a right-skewed distribution is distorted toward the tail created by the few high scores and therefore is greater than the median. The mean for the left-skewed distribution is distorted toward the tail created by the few low scores and therefore is less than the median. When you have a skewed distribution, you should use the median because atypical scores in the distribution do not distort the median. Consider

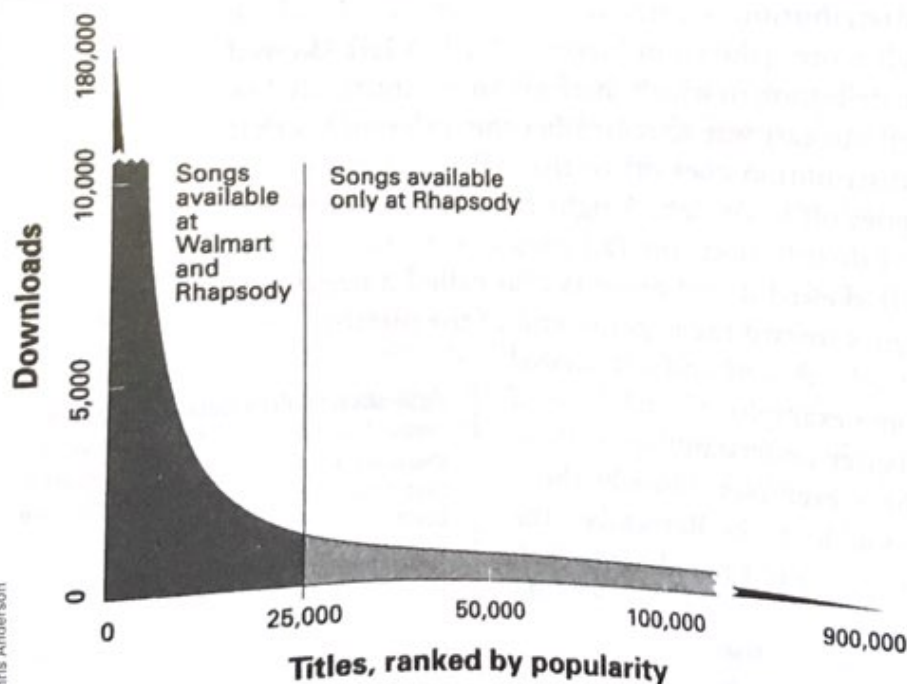


Figure 1.6 | An Example of a Right-Skewed Distribution

In this case, it is the "long tail" phenomenon involving music downloads from Walmart and online music retailer Rhapsody.

this example from Wheelan (2013). The median annual income for 10 guys sitting at a bar is \$35,000 (each earns \$35,000 a year). When multibillionaire Bill Gates walks in and sits at the bar, the median remains \$35,000. If another multibillionaire comes in and sits at the bar, the median still does not change. Why? Atypical scores do not distort the median. This means that you need to know the type of frequency distribution before deciding which measure of central tendency—mean or median—is more appropriate. Beware, sometimes the inappropriate measure of central tendency for skewed distributions (the mean) is used to mislead you (Huff, 1954).

Skewed distributions are also important to understand because various aspects of everyday life, such as medical trends (mortality rates for various diseases), are often skewed. Let's consider a famous example of the importance of understanding skewed distributions (Gould, 1985). Stephen Jay Gould, a noted Harvard scientist, died of cancer in 2002. However, this was 20 years after he was diagnosed with abdominal mesothelioma cancer and told that this type of cancer had "a median mortality rate of 8 months after diagnosis." Most people would think that they would only have about 8 months to live if given this median statistic. However, Gould realized that his expected chances depended upon the type of frequency distribution for the deaths from this disease. Because the statistic is reported as a median rather than a mean, the distribution is skewed. Now, if you were Gould, which type of skewed distribution would you want—right or left? Many people at first think they would want left-skewed, but you wouldn't want this distribution because everyone would be dead within less than a year. Look at its shape in Figure 1.5(b). If it is 8 months from the origin to the median, then it is less than 4 months from the median to the end of the distribution. You would want a severely right-skewed distribution with a long tail to the right, going on for years. This is what Gould found the distribution to be when he examined the medical literature on the disease. The distribution had a tail that stretched out to the right for many years beyond the median, and Gould was fortunate to be out in this long tail, living for 20 more years after getting the diagnosis. When he did die, it was from a different type of cancer (Blastland & Dilnot, 2009).

Gould wrote a famous article called "The Median Isn't the Message," in which he argued that his knowledge of statistics saved him from the erroneous conclusion that he would necessarily be dead in 8 months. In confronting his illness, Gould was thinking like a scientist. Such thinking provided him and the many readers of his article with a better understanding of a very difficult medical situation. Thinking like a scientist allows all of us to gain a better understanding of ourselves, others, and the world we all inhabit. Such thinking, along with the accompanying research, has enabled psychological scientists to gain a much better understanding of human behavior and mental processing. We describe the basic findings of this research in the remainder of the book. You will benefit not only from learning about these findings but also from thinking more like a scientist in your daily life.

Section Summary

To understand research findings, psychologists use statistics—a branch of mathematics that provides procedures for the description and analysis of data. In this section, we were concerned with descriptive statistics. Measures of central tendency allow a researcher to describe the “typical” score for a distribution of scores concisely. There are three such measures: mean, median, and mode. The mean is merely the arithmetical average. The median is the middle score when the distribution is arranged in ascending or descending order. The mode is the most frequently occurring score. Of these three measures, the mean is used most often. However, if unusually high or low scores in the distribution distort the mean, then the median should be used. In addition to describing the typical score, we need to determine the variability of the scores. We could use the range—the difference between the highest and lowest scores—but unusually low or high scores distort it. The measure of variability most often used is the standard deviation, the average extent that the scores vary from the mean of the distribution.

The standard deviation is especially relevant to the normal (bell-shaped) frequency distribution. Sixty-eight percent of the scores in a normal distribution fall within 1 standard deviation of the mean, 95% within 2 standard deviations, and over 99% within 3 standard deviations. These percentages hold true regardless of the value of the standard deviation. They also enable us to compute the percentile rank for a specific score in a normal distribution. The percentile rank for a score is the percentage of the scores below it in the distribution of scores. All distributions are not symmetric like the normal distribution. Two important nonsymmetric distributions are the right-skewed and left-skewed distributions. In a right-skewed distribution, there are some unusually high scores, leading to the mean being greater than the median; in a left-skewed distribution, there are some unusually low scores, leading to the mean being less than the median. In both cases, the median should be used because the mean is distorted toward the tail of the distribution.

ConceptCheck | 3

Explain what measures of central tendency and measures of variability tell us about a distribution of scores.

Explain why the normal distribution has a bell shape.

Explain the relationship between the mean and median in a right-skewed distribution and in a left-skewed distribution.

Study Guide

Chapter Key Terms

You should know the definitions of the following key terms from the chapter. They are listed in the order in which they appear in the chapter. For those you do not know, return to the relevant section of the chapter to learn them. When you think that you know all of the terms, complete the matching exercise based on these key terms.

psychology	correlation coefficient	nocebo effect
biological perspective	positive correlation	placebo group
cognitive perspective	negative correlation	inferential statistical analyses
behavioral perspective	regression toward the mean	double-blind procedure
sociocultural perspective	mean	meta-analysis
hindsight bias (I-knew-it-all-along phenomenon)	scatterplot	descriptive statistics
descriptive methods	third-variable problem	frequency distribution
naturalistic observation	spurious correlation	median
participant observation	random assignment	mode
case study	independent variable	range
survey research	dependent variable	standard deviation
population	experiment	normal distribution
sample	experimental group	percentile rank
random sampling	control group	right-skewed distribution
correlational study	operational definition	left-skewed distribution
variable	placebo effect	
	placebo	

Key Terms Exercise

Identify the correct term for each of the following definitions. The answers to this exercise follow the answers to the ConceptChecks at the end of the chapter.

1. An explanation of a correlation between two variables in terms of another variable that could possibly be responsible for the observed relationship between the two variables.

2. A control measure in an experiment in which neither the experimenters nor the participants know which participants are in the experimental and control groups.

3. The score positioned in the middle of a distribution of scores when all of the scores are arranged from lowest to highest.

4. An asymmetric frequency distribution in which there are some unusually high scores that distort the mean to be greater than the median.

5. Improvement due to the expectation of improving because of receiving treatment.

6. An inverse relationship between two variables.

7. A control measure in an experiment in which participants are randomly assigned to groups in order to equalize participant characteristics across the various groups in the experiment.

 8. The percentage of scores below a specific score in a distribution of scores.

 9. A research perspective whose major explanatory focus is how the brain, nervous system, and other physiological mechanisms produce our behavior and mental processes.

 10. A description of the operations or procedures that a researcher uses to manipulate or measure a variable.

 11. The tendency, after learning about an outcome, to be overconfident in one's ability to have predicted it.

 12. A visual depiction of correlational data in which each data point represents the scores on the two variables for each participant.

 13. The entire group of people that a researcher is studying.

 14. The difference between the highest and lowest scores in a distribution of scores.

 15. Statistical analyses that allow researchers to draw conclusions about the results of a study by determining the probability the results are due to random variation (chance).

1. Which of the following major research perspectives focuses on conditioning by external environmental events as the major cause of our behavior?
 - a. biological
 - b. cognitive
 - c. behavioral
 - d. sociocultural
 2. Which of the following would be the best procedure for obtaining a representative sample of the students at your school?
 - a. sampling randomly among students in the student union
 - b. sampling randomly among students studying in the library
 - c. sampling randomly among the students who belong to Greek organizations
 - d. sampling randomly from a list of all the students enrolled at your school
 3. Which of the following research methods allow(s) the researcher to draw cause-effect conclusions?
 - a. descriptive
 - b. correlational
 - c. experimental
 - d. all of the above
 4. Height and weight are _____ correlated; elevation and temperature are _____ correlated.
 - a. positively; positively
 - b. positively; negatively
 - c. negatively; positively
 - d. negatively; negatively
 5. Which of the following correlation coefficients indicates the STRONGEST relationship?
 - a. +.75
 - b. -.81
 - c. +1.25
 - d. 0.00
 6. Manipulate is to measure as _____ is to _____.
 - a. positive correlation; negative correlation
 - b. negative correlation; positive correlation
 - c. independent variable; dependent variable
 - d. dependent variable; independent variable

Practice Test Questions

The following are practice multiple-choice test questions on some of the chapter content. The answers are given after the Key Terms Exercise answers at the end of the chapter. If you guessed on a question or incorrectly answered a question, restudy the relevant section of the chapter.

7. In an experiment, the _____ group participants receive an inactive treatment but are told that the treatment will help them.
 - a. experimental
 - b. control
 - c. placebo
 - d. third-variable
8. The most frequently occurring score in a distribution of scores is the _____, and the average score is the _____.
 - a. mode; mean
 - b. mean; mode
 - c. median; mean
 - d. mean; median
9. About _____ % of the scores in a normal distribution are between -1 standard deviation and $+1$ standard deviation of the mean.
 - a. 34
 - b. 68
 - c. 95
 - d. 99
10. In a left-skewed distribution, the mean is _____ than the median; in a right-skewed distribution, the mean is _____ than the median.
 - a. greater; greater
 - b. greater; less
 - c. less; greater
 - d. less; less
11. Shere Hite's failure to use _____ resulted in misleading findings for her women and love survey study.
 - a. a placebo group
 - b. a double-blind procedure
 - c. random assignment
 - d. random sampling
12. Professor Jones noticed that the distribution of students' scores on his last biology exam had an extremely small standard deviation. This indicates that:
 - a. the exam was given to a very small class of students.
 - b. the exam was a poor measure of the students' knowledge.
 - c. the students' scores tended to be very similar to one another.
 - d. the students' mean exam score was less than the median exam score.
13. In a normal distribution, the percentile rank for a score that is 1 standard deviation below the mean is roughly _____.
 - a. 16
 - b. 34
 - c. 68
 - d. 84
14. Dian Fossey's study of gorillas is an example of _____.
 - a. naturalistic observation
 - b. participant observation
 - c. naturalistic observation that turned into participant observation
 - d. a case study
15. Which of the following types of scatterplots depicts a weak, negative correlation?
 - a. a lot of scatter with data points going from top left to bottom right
 - b. very little scatter with data points going from top left to bottom right
 - c. a lot of scatter with data points going from bottom left to top right
 - d. very little scatter with data points going from bottom left to top right

Chapter ConceptCheck Answers

ConceptCheck | 1

- Both of these research perspectives emphasize internal causes in their explanations of human behavior and mental processing. The biological perspective emphasizes the role of our actual physiological hardware, especially the brain and nervous system, while the cognitive perspective emphasizes the role of our mental processes, the "programs" of the brain. For example, biological explanations will involve actual parts of the brain or chemicals in the brain. Cognitive explanations, however, will involve mental processes such as perception and memory without specifying the parts of the brain or chemicals involved in these processes. Thus, the biological and cognitive perspectives