

Mind on Statistics 6th edition Utts Solutions Manual

SKU: 9781337793605-SOLUTIONS

Mind on Statistics

6th Edition

Complete Solutions Manual

© 2022 Cengage Learning, Inc.



© 2022 Cengage Learning, Inc.

Unless otherwise noted, all content is © Cengage.

ALL RIGHTS RESERVED. No part of this work covered by the copyright herein may be reproduced or distributed in any form or by any means, except as permitted by U.S. copyright law, without the prior written permission of the copyright holder.

For product information and technology assistance, contact us at
Cengage Customer & Sales Support,
1-800-354-9706 or support.cengage.com.

For permission to use material from this text or product, submit
all requests online at **www.cengage.com/permissions.**

ISBN: 978-133779484-8

Cengage
200 Pier 4 Boulevard
Boston, MA 02210
USA

Cengage is a leading provider of customized learning solutions with employees residing in nearly 40 different countries and sales in more than 125 countries around the world. Find your local representative at:
www.cengage.com.

To learn more about Cengage platforms and services, register or access your online learning solution, or purchase materials for your course, visit
www.cengage.com.

NOTE: UNDER NO CIRCUMSTANCES MAY THIS MATERIAL OR ANY PORTION THEREOF BE SOLD, LICENSED, AUCTIONED, OR OTHERWISE REDISTRIBUTED EXCEPT AS MAY BE PERMITTED BY THE LICENSE TERMS HEREIN.

READ IMPORTANT LICENSE INFORMATION

Dear Professor or Other Supplement Recipient:

Cengage Learning has provided you with this product (the "Supplement") for your review and, to the extent that you adopt the associated textbook for use in connection with your course (the "Course"), you and your students who purchase the textbook may use the Supplement as described below. Cengage Learning has established these use limitations in response to concerns raised by authors, professors, and other users regarding the pedagogical problems stemming from unlimited distribution of Supplements.

Cengage Learning hereby grants you a nontransferable license to use the Supplement in connection with the Course, subject to the following conditions. The Supplement is for your personal, noncommercial use only and may not be reproduced, or distributed, except that portions of the Supplement may be provided to your students in connection with your instruction of the Course, so long as such students are advised that they may not copy or distribute any portion of the Supplement to any third party. Test banks and other testing materials may be made available in the classroom and collected at the end of each class session, or posted electronically as described herein. Any material posted

electronically must be through a password-protected site, with all copy and download functionality disabled, and accessible solely by your students who have purchased the associated textbook for the Course. You may not sell, license, auction, or otherwise redistribute the Supplement in any form. We ask that you take reasonable steps to protect the Supplement from unauthorized use, reproduction, or distribution. Your use of the Supplement indicates your acceptance of the conditions set forth in this Agreement. If you do not accept these conditions, you must return the Supplement unused within 30 days of receipt.

All rights (including without limitation, copyrights, patents, and trade secrets) in the Supplement are and will remain the sole and exclusive property of Cengage Learning and/or its licensors. The Supplement is furnished by Cengage Learning on an "as is" basis without any warranties, express or implied. This Agreement will be governed by and construed pursuant to the laws of the State of New York, without regard to such State's conflict of law rules.

Thank you for your assistance in helping to safeguard the integrity of the content contained in this Supplement. We trust you find the Supplement a useful teaching tool.

Printed in the United States of America
Print Number: 01 Print Year: 2022

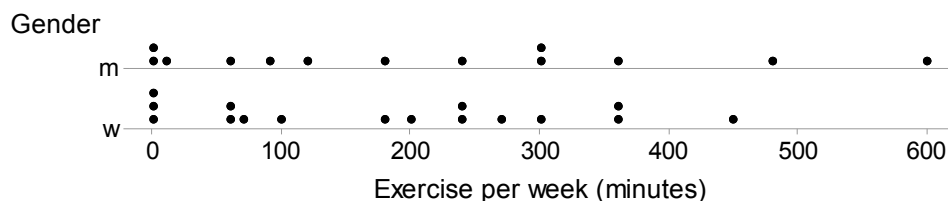
CHAPTER 1

EXERCISE SOLUTIONS

- 1.1**
- a. The fastest speed was 150 miles per hour.
 - b. The slowest speed driven by a male was 55 miles per hour.
 - c. 1/4 of the females reported having driven at 95 miles per hour or faster. Notice that 95 mph is the *upper quartile* for females. By definition, about 1/4 of the values in a dataset are greater than the upper quartile.
 - d. 1/2 of the females reported having driven 89 mph or faster. Notice that 89 mph is the *median* value.
 - e. 1/2 of 102 = 51 females have driven 89 mph or faster.
- Note:* For parts (d) and (e) the answer would have to be adjusted if there were any females who reported 89 as their value, but from the data on page 4 we can see that there were not. Because there were no “ties” with the median, we know that exactly half of the values fall above it and the other half fall below it.
- 1.2**
- a. The median height is 65 inches.
 - b. $\text{Range} = \text{tallest} - \text{shortest} = 71 - 59 = 12$ inches.
 - c. The interval from 59 to 63.5 inches contains the shortest 1/4 of these students. This interval is from the minimum to the lower quartile.
 - d. The interval from 63.5 to 67.5 inches contains the middle 1/2 of these students. This is an interval from the lower quartile to the upper quartile.
- 1.3**
- a. The observed rate of cervical cancer in Vietnamese American women is 86 per 200,000. This could also be expressed as 43 per 100,000 or 4.3 per 10,000, and so on. In decimal form, it is .00043.
 - b. The risk of developing cervical cancer for Vietnamese American women in the next year is $86/200,000 = .00043$.
 - c. The rate of 86 per 200,000 is based on past data and tells us the number of Vietnamese American women who developed cervical cancer out of a population of 200,000. The risk utilizes the rate from the past to tell us the future likelihood of cervical cancer in other Vietnamese American women.
- 1.4**
- a. The base rate is about 13 in 1000, or about .013.
 - b. The risk for men who smoke is just over 13 times the rate for non-smokers, or about .169.
- 1.5**
- a. All teens in the United States at the time the poll was taken.
 - b. All teens in the United States who had dated at the time the poll was taken.
- 1.6**
- A population is a collection of all individuals of interest while a sample is a subset of the population of interest, for which measurements are taken in a study. In Case Study 1.6, the population of interest is probably all men, and possibly women as well. However, the sample consisted of 22,071 male physicians who volunteered for the study, so the population to which the results apply is all men similar to them.
- 1.7**
- a. All adults in the United States at the time the poll was taken.
 - b. $\frac{1}{\sqrt{1048}} = .031$ or 3.1%.
 - c. $34\% \pm 3.1\%$, or 30.9% to 37.1%.
- 1.8**
- a. The population is probably all Canadians who were eligible to participate in the survey (which is probably all adults with telephones).
 - b. There were 2000 people in the sample.
 - c. $\frac{1}{\sqrt{n}} = \frac{1}{\sqrt{2000}} = .022$ or 2.2%.
 - d. In the sample, 16% viewed immigration as having a negative impact. The interval $16\% \pm 2.2\%$, or 13.8% to 18.2%, is 95% certain to cover the true percent of Canadians who viewed immigration as having a negative impact (at the time of the poll).

- 1.9 Solve for n in the equation $\frac{1}{\sqrt{n}} = .05 = \frac{1}{20}$. Answer is $n = 400$ teenagers.
- 1.10 Solve for n in the equation $\frac{1}{\sqrt{n}} = .10$. Answer is $n = 100$.
- 1.11
- This is an example of a self-selected or volunteer sample. Magazine readers voluntarily responded to the survey and were not randomly selected.
 - These results may not represent the opinions of all readers of the magazine. The people who respond probably do so because they feel stronger about the issues (for example, violence on television or physical discipline) than the readers who do not respond. So, they may be likely to have a generally different point of view than those who do not respond.
- 1.12 The exercise did not specify what the survey is about, but no matter what, the survey is based on a self-selected sample, and people who feel strongly about the issues and/or who have extra time are more likely to respond. The results will not be representative of all students who use the cafeteria.
- 1.13
- Randomized experiment (because students were randomly assigned to the two methods).
 - Observational study (because people cannot be randomly assigned to smoke or not).
 - Observational study (because people cannot be randomly assigned to be a CEO or not).
- 1.14
- Randomized experiment, because students were randomly assigned to receive Vitamin C or placebo.
 - Observational study, because the patients were not randomly assigned to do anything. (Note that a random sample is not the same thing as random assignment.)
 - Randomized experiment, because participants were randomly assigned to meditation or low-fat diet.
- 1.15 Answers will vary, but one possibility is general level of activity. It is likely to differ for elderly people who attend church regularly and those who don't, and it is also likely to affect blood pressure. So it might partially explain the results of this study.
- 1.16
- Number of courses might be a confounding variable. Students taking many courses may sleep less due to the amount of work involved and may not do as well in school due to the load.
 - Weight is not likely to be a confounding variable. Weight is probably not related either to amount of sleep or to grades.
 - Hours spent partying might be a confounding variable. Students who party a lot may sleep less, and may also get lower grades because they're not studying.
- 1.17 You would need to know how large the difference in weight loss was for the two groups. If the difference in weight loss is very small (but not 0), it could be statistically significant, but not have much practical importance.
- 1.18 Statistical significance is when there is a relationship or difference that is large enough to be unlikely to have occurred in the sample if there was no relationship or difference in the population of interest. Practical significance occurs when the relationship or difference is large enough to be important or meaningful in a "real world" sense. A result can be statistically significant, but not practically significant. This is especially likely to occur in studies with very large sample sizes.
- 1.19 You would want to know how many different relationships were examined. If this result was the only one that was statistically significant out of many examined, it could easily be a false positive.
- 1.20 A false positive occurs when a relationship or difference is said to be statistically significant based on examining information from a sample, but in fact there is no relationship or difference in the population.
- 1.21 The placebo group estimates the baseline rate of heart attacks for men not taking aspirin. So, the estimated baseline rate of heart attacks is $189/11,034$, which is about 17 heart attacks per 1000 men or $17/1000$. (See Table 1.1 for the data.)
- 1.22
- The amount of exercise per week is similar for men and women except that there are a few high values for the men. The dotplot follows.

Figure for Exercise 1.22a



- b. *Women*, median = 190 min. *Men*, median = 180 min.

To find the median, put the data in order first.

For *women*, the ordered list of data is:

0, 0, 0, 60, 60, 70, 100, **180**, **200**, 240, 240, 270, 300, 360, 360, 450

The number of women is even (16), so the median is the average of the middle two values in the ordered list. These middle two values, underlined and bold in the above list, are 180 and 200 and their average is 190.

For *men*, the ordered list of data is:

0, 0, 14, 60, 90, 120, **180**, 240, 300, 300, 360, 480, 600

The number of men is odd (13), so the median is the middle value in the ordered data. This value, underlined and bold in the list above, is 180.

- c. Although the median response is different for women and men, the difference is only 10 minutes. The weekly amount of exercise is about the same for the samples of women and men.

1.23

a.

	Minutes of Exercise per Week		
Median		180	
Quartiles	37		330
Extremes	0		600

To determine the summary, first write the responses in order from smallest to largest.

The ordered list of data is:

0, 0, 14, 60, 90, 120, 180, 240, 300, 300, 360, 480, 600

Minimum = 0 min.

Maximum = 600 min.

Median = 180 min. (middle value in the ordered list)

Lower quartile = 42 min. It is the median of the values smaller than the median.

These are 0, 0, 14, 60, 90, 120.

Median of these six values is $(14 + 70)/2 = 42$.

Upper quartile = 330 min. It is the median of the values larger than the median.

Values larger than the median are 240, 300, 300, 360, 480, 600.

Median of these values is $(300 + 360)/2 = 330$.

- b. Reported exercise hours per week for the men in the sample ranged from a low of 0 to a high of 600 minutes per week. The median response was 180 minutes (3 hours). About 1/2 of the men (the middle

half) reported exercising between 37 and 330 minutes (5 and a half hours) per week. About 1/4 said they exercised less than 37 minutes per week while 1/4 said they exercised more than 330 minutes per week.

1.24 a.

	Minutes of Exercise per Week		
Median		190	
Quartiles	60		285
Extremes	0		450

To determine the summary, first write the responses in order from smallest to largest.

The ordered list of data is:

0, 0, 0, 60, 60, 70, 100, 180, 200, 240, 240, 270, 300, 360, 360, 450

Minimum = 0 min.

Maximum = 450 min.

Median = 190 min. (average of the middle two values in the ordered list, which are 180 and 200)

Lower quartile = 60 min. It is the median of the values that are smaller than the median.

These are 0, 0, 0, 60, 60, 70, 100, 180.

Median of these eight values is $(60 + 60)/2 = 60$.

Upper quartile = 285 min. It is the median of the values that are larger than the median.

Values larger than the median are 200, 240, 240, 270, 300, 360, 360, 450.

Median of these values is $(270 + 300)/2 = 285$.

- b. Reported exercise hours for the women in the sample ranged from a low of 0 to a high of 450 minutes per week. The median response was 190 minutes (3 hours and 10 minutes). About 1/2 of the women (the middle half) reported exercising between 60 and 285 minutes per week. About 1/4 said they exercised less than 60 minutes per week while 1/4 said they exercised more than 285 minutes per week.
- 1.25 a. This is an observational study because vegetarians and non-vegetarians are compared and these groups occur naturally. People were not assigned to treatment groups.
- b. Since this is an observational study and not a randomized experiment, we cannot conclude that a vegetarian diet causes lower death rates from heart attacks and cancer. Other variables not accounted for may be causing this reduction.
- c. This answer will differ for each student. One potential confounding variable is amount of exercise. This is a confounding variable because it may be that vegetarians also exercise more on average and this led to lower death rates from heart attacks and cancer.
- 1.26 Base rates were not given. In this study, a base rate would be the actual rate (risk) of a particular cause of death for people who are not vegetarians.
- 1.27 The base rate or baseline risk is missing from the report. You need to know the base rate of cancer of the rectum for men to decide if the increased risk from drinking beer is large or small.
- 1.28 a. This was a randomized experiment because volunteers were randomly assigned to wear either a nicotine patch or a placebo patch.
- b. You can conclude that use of nicotine patches leads to a higher success rate for those trying to quit smoking than use of placebo patches.
- c. It was advisable to assign some of the patients to wear a placebo patch because then you can compare the success rate of those patients to the success rate of the patients wearing nicotine patches. You will also learn in a future chapter that even though they have no active ingredients, placebos can have a large psychological effect. Also, presumably people in the experiment want to quit smoking, so some will succeed regardless of treatment method.

- 1.29 For Caution 1: Because the difference given in the previous exercise is a large difference (46% with nicotine patch and 20% with placebo), it has practical importance as well as statistical significance. For Caution 2: Because the result is based on a randomized experiment, it is not possible that whether someone quit or not influenced the type of patch they were assigned.
- 1.30 The definition is “**Statistics** is a collection of procedures and principles for gathering data and analyzing information to help people make decisions when faced with uncertainty.” In the study described in Exercise 1.28, data were gathered and analyzed in order to make a decision about the effectiveness of the nicotine patch. This information will help individuals decide whether to wear nicotine patches when trying to quit smoking. Although the observed result was based only on a sample from a larger population, the data collection and analysis methods make it reasonable to conclude that the patch is more effective than a placebo for the population represented by this sample.
- 1.31 Neither caution applies. The magnitude of the difference is given in Case Study 1.6, and considering the number of men between the ages of 40 and 84 in the U.S. population, the given difference has practical importance. Because men were randomly assigned to take aspirin (or not) we can conclude that the correct direction of the cause and effect is that taking aspirin caused the reduction in heart attacks.
- 1.32
- a. The population is all University of California faculty members in 1995.
 - b. The margin of error is approximately $\frac{1}{\sqrt{n}} = \frac{1}{\sqrt{1000}} = .032$, or roughly 3%.
 - c. It cannot be concluded that a majority of all University faculty favored the criteria. In the sample a slight majority (52%) was in favor, but the margin of error was about 3%. So, in the population the percent in favor could possibly be a minority (below 50%). An interval that is 95% certain to contain the percent in the population is $52\% \pm 3\%$, which is 49% to 55%. Because some values within this interval are below 50%, we are not able to rule out the possibility that the percent in favor in the population is a minority.
- 1.33
- a. The margin of error is about $\frac{1}{\sqrt{n}} = \frac{1}{\sqrt{2003}} = .022$.
 - b. The sample proportion is $360/2003 = .180$. So, the interval is $.180 \pm .022$, which is .158 to .202.
This is *sample proportion $\pm$ margin of error*.
- 1.34 For this population, the risk of seeing a ghost is $360/2003 = .180$.
- 1.35
- a. This is a self-selected (or volunteer) sample.
 - b. Probably higher, because people who would say they have seen a ghost would be more likely to call the late-night radio talk show than others. They might even be more likely to be listening to such a show.
- 1.36
- a. Self-selected sample or volunteer sample.
 - b. The results cannot be extended to any larger population because the sample was not selected to be representative of any population.
- 1.37 The term *data snooping* refers to looking at the data in a variety of ways until something interesting to report emerges.
- 1.38
- a. This statement has to be based on an observational study. The researchers observed who was breast-fed and who was not. It would not be possible to randomly assign mothers to breast-feed or not. It would not be ethical to randomly assign this treatment.
 - b. The better headline is “Link found between breast-feeding and school performance” because it does not imply that there is a cause-and-effect conclusion, while the other headline does. Because the study was observational, a cause-and-effect conclusion cannot be made. It is likely that children who were breast-fed as infants have other (confounding) factors in their lives that differ from children who were not, and that may influence school performance. For example, perhaps they are more likely to have mothers who don’t work, perhaps they are more likely to be first-born children, and so on. We can say that a link was found, but cannot say that breast-feeding *leads* to better performance.
- 1.39 In some situations it is not practical or even possible to conduct a randomized experiment. For example, a researcher may wish to study whether occupational exposure to asbestos affects the risk of lung disease. It

would not be possible, or ethical, to assign people to occupations that involve differing amounts of exposure to asbestos.

- 1.40** An observational study was done instead of an experiment because the researchers could not assign individuals either to attend a religious service once a week and pray regularly or to not engage in these practices.
- 1.41** The answer will differ for students, but here is an example. Randomly assign volunteers to either eat lots of chocolate or not eat any chocolate for a period of time, and give them a questionnaire about depression at the beginning and the end of the time period. Then compare the change in depression scores for the two groups.
- 1.42**
- a. Randomized experiment. People would not take a placebo as part of an observational study.
 - b. Maybe not, because 20 side effects were tested so one or a few could appear to be a problem (statistically significant) just by chance even if none of them were a problem. In other words, it is possible that the observed relationship between taking aspirin and having headaches was a false positive.
- 1.43** *USA Today* made the mistake of making a cause-and-effect conclusion about the relationship between prayer and blood pressure. This conclusion is not justified because the data were from an observational study. Specifically, they neglected to consider possible confounding variables like lifestyle choices, social networks, and health of the people between the two groups. As a result, people may be led to believe that if they pray more often, they will have lower blood pressure. That conclusion is not justified based on this observational study.
- 1.44**
- Step 1:** The investigators asked if taking aspirin reduces the risk of a heart attack and then in the conclusion, asked to what population the results of this study apply.
 - Step 2:** They collected data from a 5-year randomized experiment, including what treatment each doctor received (aspirin or placebo), and whether each doctor had a heart attack or not.
 - Step 3:** They summarized the data from Step 2 by categorizing the doctors according to which type of pill they took (aspirin or placebo). They then counted how many doctors in each treatment group had a heart attack. Further data analysis included calculating the “Attacks Per 1000 Doctors” for each group.
 - Step 4:** The investigators proceeded to make conclusions from the data given in Table 1.1. Specifically, they stated, “The results ... support the conclusion that taking aspirin does indeed help reduce the risk of having a heart attack. The rate of heart attacks in the group taking aspirin was only about half the rate of heart attacks in the placebo group. In the aspirin group, there were 9.42 heart attacks per 1000 participating doctors, while in the placebo group, there were 17.13 heart attacks per 1000 participants.” They then continued to question if there were any other important risk factors or differences between the two groups.
 - Step 5:** The new knowledge is that this study has provided support for the benefit of aspirin in broader populations concerning the reduction of the risk in heart attacks. As a result, millions of people take aspirin with the hope that it will prevent a heart attack.
- 1.45**
- a. Categorizing the adults by sex would be more appropriate because cancer type is likely to differ based on biological sex rather than by the social construct of gender. For instance, biological females cannot get prostate cancer and biological males cannot get ovarian cancer.
 - b. Categorizing by gender would be more appropriate because the question of interest is whether there is a relationship between gender self-identity and opinion on the death penalty.
 - c. Categorizing by gender would be more appropriate because age at marriage is a social choice, not a biological trait.
 - d. Categorizing by sex would be more appropriate because body temperature is likely to differ based on biological sex and not based on social identity.

CHAPTER 2

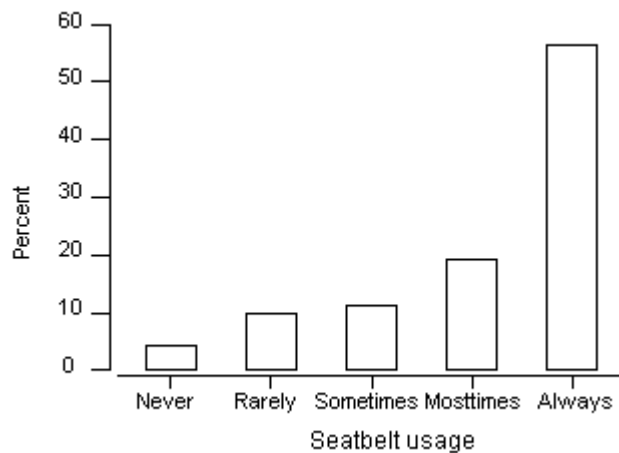
EXERCISE SOLUTIONS

- 2.1 a. 4.
 b. A state in the United States.
 c. $n = 50$.
- 2.2 a. 2.
 b. A randomly selected person.
 c. $n = 620$.
- 2.3 a. Whole population.
 b. Sample.
- 2.4 a. Whole population.
 b. Sample.
- 2.5 a. Population parameter.
 b. Sample statistic.
 c. Sample statistic.
- 2.6 a. Sample statistic.
 b. Population parameter.
 c. Sample statistic.
- 2.7 a. Sex and self-reported fastest ever driven speed.
 b. Students in a statistics class.
 c. The answer will vary. If you think the students represent a larger group of individuals, it is sample data. If interest is only in this group of students, or if you think these students do not represent any larger group, it is population data.
- 2.8 a. $n = 2391$.
 b. Individuals aged 65 years or older.
 c. Frequency of attending religious services and frequency of praying or reading the bible were related to blood pressure.
 d. Sample data. They used the data to make generalizations about a larger population.
- 2.9 This is a population summary if we restrict our interest only to the fiscal year 1998. (If we were to use this value to represent errors in other years, it could be considered to be a sample summary.)
- 2.10 a. Treatment used (placebo or aspirin) and whether individual died from heart attack or not.
 b. Male physicians between 40 and 84 years old.
 c. $n = 22,071$.
 d. Sample data. They used the data to make generalizations about a larger population.
- 2.11 a. Categorical.
 b. Quantitative.
 c. Quantitative.
 d. Categorical.
- 2.12 a. Quantitative.
 b. Categorical.
 c. Quantitative.
 d. Categorical.
- 2.13 a. Not ordinal. It's categorical but the categories are not ordered.
 b. Ordinal. Grades are ordered categories.
 c. Not ordinal. It's quantitative.

- 2.14**
- a. Continuous. All weights are possible within an interval of possibilities (although we can't measure accurately enough to observe all possibilities).
 - b. Not continuous. The number of text messages must be an integer.
 - c. Not continuous. The number of colds in a year would be an integer.
- 2.15**
- a. Explanatory variable is score on the final exam; response variable is final course grade.
 - b. Explanatory variable is gender; response variable is opinion about the death penalty.
- 2.16**
- a. Not ordinal. It's categorical but the categories are not ordered.
 - b. Ordinal. The ratings are ordered.
 - c. Not ordinal. It's quantitative.
- 2.17**
- a. Not continuous. A student could not miss 4.631 classes, for example.
 - b. Continuous. With an accurate enough measuring instrument, any measurement is possible.
 - c. Continuous. With an accurate enough time piece, any length of time is possible.
- 2.18**
- a. Explanatory variable is amount person walks or runs per day; response variable is the performance on the lung test.
 - b. Explanatory variable is age of the respondent; response variable is feeling about religious importance.
- 2.19**
- a. Whether a person supports the smoking ban or not is a categorical variable.
 - b. Gains on verbal and math SATs are quantitative variables.
- 2.20** The explanatory variable is smoker or not. The response variable is Alzheimer sufferer or not. Both variables are categorical.
- 2.21**
- a. Sex and pulse rate.
 - b. Sex is categorical, pulse rate is quantitative.
 - c. Is there a difference between the mean pulse rates of men and women? The sample mean pulse rate for each sex would be useful.
- 2.22** This will differ for each student. As an example, suppose a survey question about income only allowed the response categories 1 = under \$20,000, 2 = \$20,000 to \$49,999, and 3 = more than \$49,999. The income categories are ordered, but so little is known about actual income that the mean response is meaningless.
- 2.23** This will differ for each student. One example where numerical summaries would make sense for an ordinal variable is the response to the question "What grade do you expect in this class? 1 = A, 2 = B, 3 = C, 4 = D, 5 = F." The mean numerical response is an expected class GPA.
- 2.24** This will differ for each student.
- 2.25**
- a. A unit is a person. Dominant hand is a categorical variable and IQ is a quantitative variable. Explanatory variable is dominant hand and response variable is IQ.
 - b. A unit is a married couple. Eventual divorce status and pet ownership are both categorical variables. Explanatory variable is pet ownership and response variable is eventual divorce status.
- 2.26**
- a. A unit is a college student. GPA and hours of study each week are both quantitative variables. Explanatory variable is hours of study and response variable is GPA.
 - b. A unit is a tax-paying individual in the United States. Tax bracket is an ordinal variable and percentage donated to charities is a quantitative variable. Explanatory variable is tax bracket and response variable is percentage donated to charities.
- 2.27**
- a. $1427/2530 = .564$, which is 56.4%.
 - b. $1 - (1427/2530) = .436$, which is 43.6%.
 - c. Never: $105/2530 = .042$ (4.2%); Rarely: $248/2530 = .098$ (9.8%); Sometimes: $286/2530 = .113$ (11.3%); Most times: $464/2530 = .183$ (18.3%); Always: $1427/2530 = .564$ (56.4%).

d.

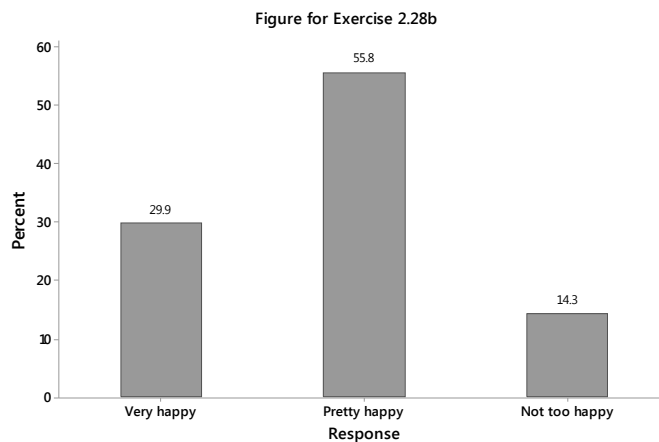
Figure for Exercise 2.27d



2.28 a.

Response	Frequency	Relative Frequency
Very happy	701	$701/2344 = .299$ (29.9%)
Pretty happy	1307	$1307/2344 = .558$ (55.8%)
Not too happy	336	$336/2344 = .143$ (14.3%)
Total	2344	1 (100%)

b. A bar graph follows.

c. $29.9\% + 55.8\% = 85.7\%$.

2.29 a.

	Preferred Use of Cell Phone		Total
	To Talk	To Text	
Women	22 (20.8%)	84 (79.2%)	106 (100%)
Men	34 (41.0%)	49 (59.0%)	83 (100%)

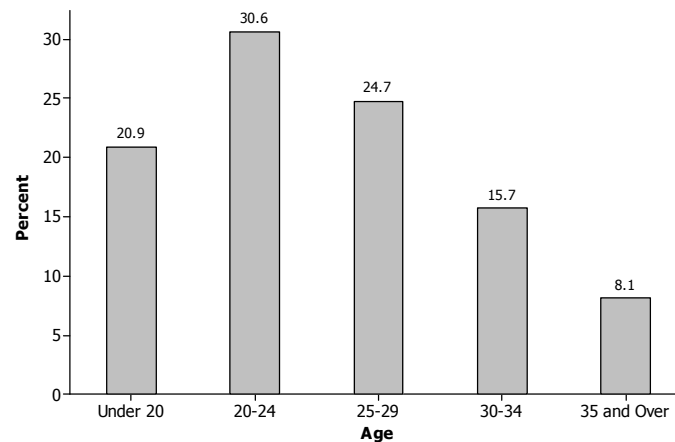
b. Self-identified women: 20.8% to talk, 79.2% to text.

c. Self-identified men: 41.0% to talk, 59.0% to text.

d. Women were more likely to say “to text” than men, whereas men were more likely to say “to talk.”

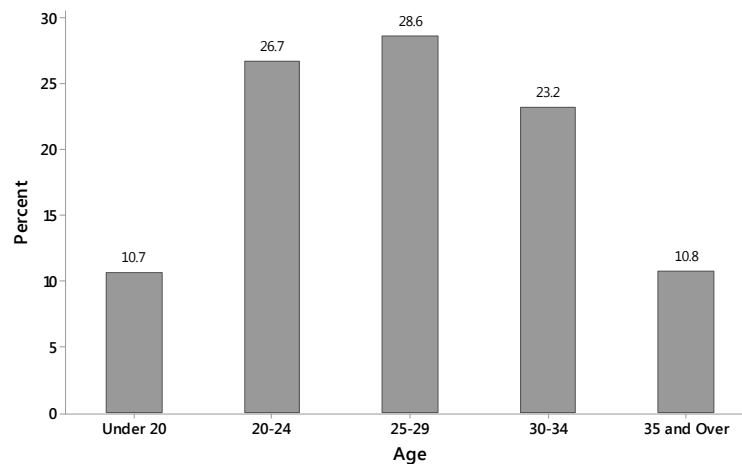
- 2.30**
- a. $1700/2470 = .688$, or 68.8%.
 - b. $1056/1700 = .621$, or 62.1%.
 - c. $300/657 = .457$, or 45.7%.
 - d. $41/113 = .363$, or 36.3%.
- 2.31**
- a. Explanatory variable is whether a person smoked or not. Response variable is whether they developed Alzheimer's or not.
 - b. Explanatory variable is political party. Response variable is whether a person voted or not.
 - c. Explanatory variable is income level. Response variable is whether a person has been subjected to a tax audit or not.
- 2.32**
- a. The bar graph for 2006 is as follows.

Figure for Exercise 2.32a



- b. The bar graph for 2018 is as follows.

Figure for Exercise 2.32b



- c. There are higher proportions in the three oldest age groups in 2018 compared to 2006.

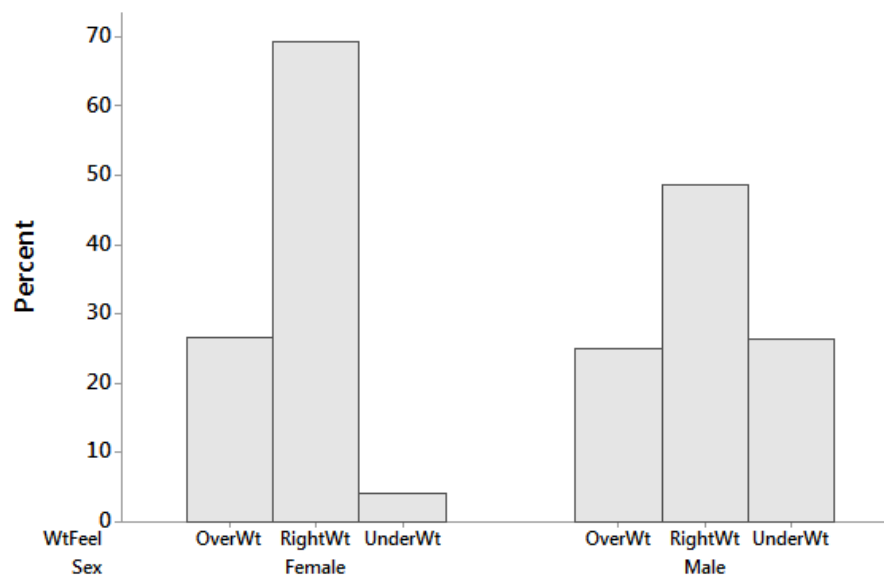
- 2.33 a. The explanatory variable is gender and the response variable is how they feel about their weight.
b.

Feelings About Weight				
Gender	Overweight	About Right	Underweight	Total
Women	38 (26.6%)	99 (69.2%)	6 (4.2%)	143
Men	18 (23.1%)	35 (44.9%)	25 (32.1%)	78

- c. Feeling overweight: $38/143 = .266$, or 26.6%; about right: $99/149 = .692$, or 69.2%; underweight: $6/149 = .042$, or 4.2%.
d. Feeling overweight: $18/78 = .231$, or 23.1%; about right: $35/78 = .449$, or 44.9%; underweight: $25/78 = .321$ or 32.1%.
e. Men are more likely than women to feel that they are underweight; women are more likely than men to say that their weight is about right.

- 2.34 The bar graph follows.

Figure for Exercise 2.34

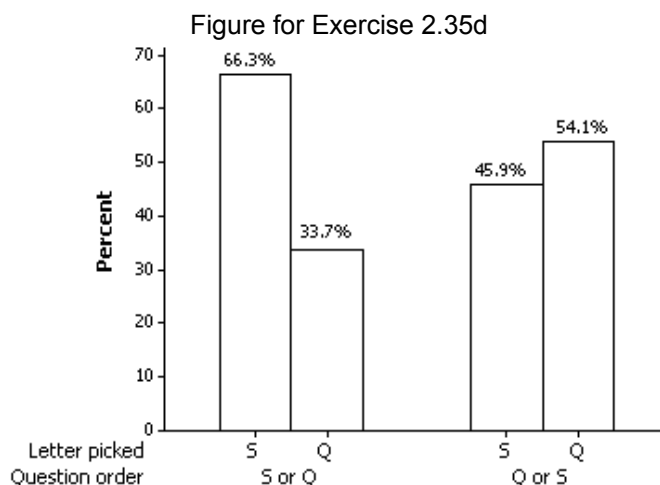


- 2.35 a.

	Picked S	Picked Q	Total
S Listed First	61	31	92
Q Listed First	45	53	98
Total	106	84	190

- b. Picked S = $(61/92) \times 100\% = 66.3\%$; Picked Q = $(31/92) \times 100\% = 33.7\%$.
c. Picked S = $(45/98) \times 100\% = 45.9\%$; Picked Q = $(53/98) \times 100\% = 54.1\%$.

d.

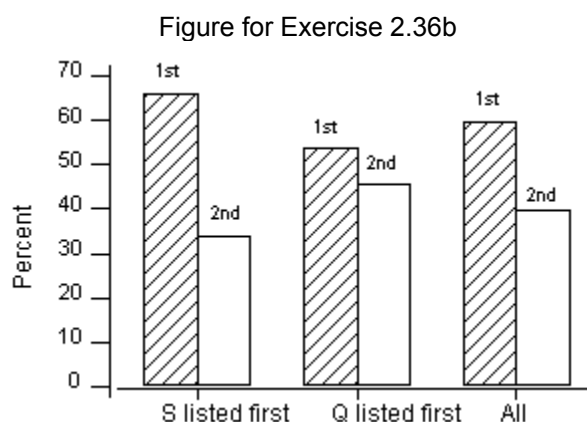


e. Parts (b) and (c) show that the percentage picking S was higher when S was listed first than when Q was listed first. It looks like the letter picked was influenced by the letter listed first.

- 2.36 a. The columns can be labeled “Picked 1st letter” and “Picked 2nd letter.” In the “S listed first” row, list the counts in the same order as in the table for Exercise 2.35(a). In the “Q listed first” row, the count for “Q picked” should be in the “Picked 1st letter” column (because Q was the first letter). The table is

	Picked 1st Letter	Picked 2nd Letter	Total
S Listed First	61 (66%)	31 (34%)	92
Q Listed First	53 (54%)	45 (46%)	98
All	114 (60%)	76 (40%)	190

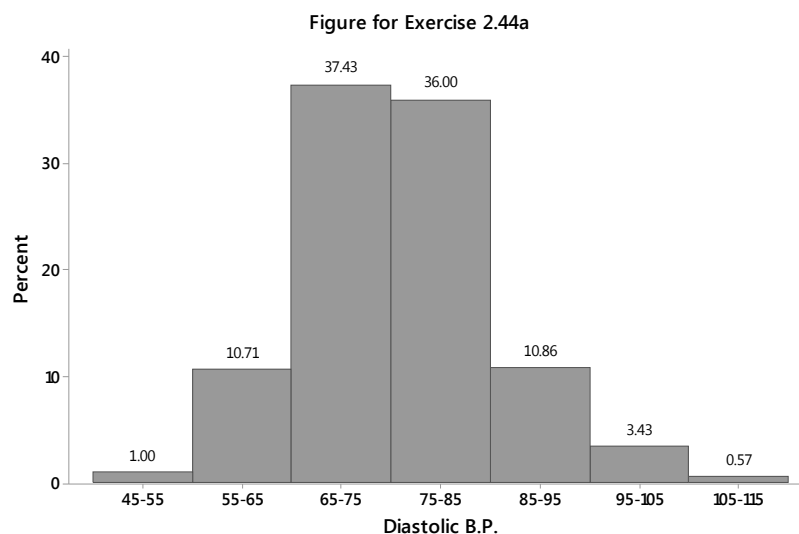
b. A bar chart of percentages picking first and second letters given on each form of the question (row percentages in the table above) along with these percentages for the overall sample follows:



c. The variables used in this exercise are more appropriate for illustrating the point of this dataset. The question of interest is whether participants might be more likely to pick the first letter given than the second regardless of whether it was an S or Q. The bar chart given for part (b) illustrates that the first letter listed was picked more often for both forms.

- 2.37
- a. The fastest speed was 150 miles per hour.
 - b. The slowest speed driven by a male student was 55 miles per hour.
 - c. $1/4$ of the female students reported having driven at 95 miles per hour or faster. Notice that 95 mph is the *upper quartile* for female students. By definition, about $1/4$ of the values in a dataset are greater than the upper quartile.
 - d. $1/2$ of the female students reported having driven 89 mph or faster. Notice that 89 mph is the *median* value.
 - e. $1/2$ of $102 = 51$ female students have driven 89 mph or faster.
- 2.38
- a. The median value is 110 mph for male students, compared to 89 mph for female students.
 - b. The spread is about the same for the two genders. The spread of the extremes is $150 - 55 = 95$ mph for the male students, compared to $130 - 30 = 100$ mph for the female students. The spread of the quartiles is slightly greater for males ($120 - 95 = 25$ for the male students, compared to $95 - 80 = 15$ for the female students).
- 2.39
- a. The center for the female server is at a greater percentage than it is for the male server. For the female server the center looks to be somewhere around 27%. For the male server, the center looks to be a bit less than 18%.
 - b. The data are more spread for the female server.
 - c. The greatest two female server percentages are set apart from the bulk of the data. The values are about 65% and 72%.
- 2.40
- a. Median height = 65 inches.
 - b. Range = Tallest – Shortest = $71 - 59 = 12$ inches.
 - c. The interval from 59 to 63.5 inches contains the shortest $1/4$ of these students. This interval is from the minimum to the lower quartile.
 - d. The interval from 63.5 to 67.5 inches contains the middle $1/2$ of these students. This is an interval from the lower quartile to the upper quartile.
- 2.41
- a. The median value, 65 inches, describes the location.
 - b. The interval described by the extremes, 59 to 71 inches, describes spread. We might also describe spread using the interval 63.5 to 67.5, the spread of the middle 50% of the data.
- 2.42
- a. The dataset is skewed to the right (it stretches in that direction).
 - b. The value 13 looks to be an outlier. It is separate from the bulk of the data.
 - c. 2 ear pierces was the most reported value. About 44 or so of the students said they had this many ear pierces.
 - d. About 32 or so of the students said they had 4 ear pierces.
- 2.43
- a. The dataset looks approximately symmetric and bell-shaped.
 - b. There are no noticeable outliers.
 - c. The most frequently reported value for sleep was 7 hours.
 - d. Roughly 14 or so students said they slept 8 hours the previous night.

- 2.44 a.** A relative frequency histogram with percentages follows. The intervals of diastolic blood pressure include the lower value but not the upper value.



- b.** The majority of values fall between 65 and 84. The percentage of individuals with values between 65 and 84 is $37.43\% + 36.00\% = 73.43\%$. With regard to shape, there may be a slight skew to the right.
- c.** A stem-and-leaf plot presents all individual values, but we do not have the data in this form. It has already been tallied for specified intervals.
- 2.45 a.** The dataset looks approximately symmetric and bell-shaped.
- b.** The highest temperature is 92°F .
- c.** The lowest temperature is 64°F .
- d.** $5/20 = .25$, which is 25%.
- 2.46 a.** The following stem-and-leaf uses “hundreds” for stems (row labels), “tens” for leaves within the rows and truncates the “ones” digit. Each “hundreds” is represented by two separate stems, one for leaves 0–4 and one for leaves 5–9.

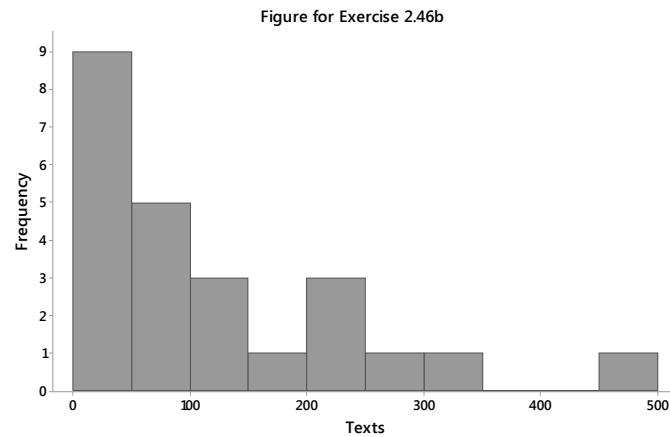
Figure for Exercise 2.46a

```

| 0 | 001222233
| 0 | 55569
| 1 | 002
| 1 | 5
| 2 | 002
| 2 | 5
| 3 | 0
| 3 |
| 4 |
| 4 | 5

```

- b.** In the following histogram, values tied with the lower value of an interval are counted into that interval. For example, 50 is counted into the interval that spans from 50 to 100.



c. The data are skewed to the right with a possible outlier (at 450).

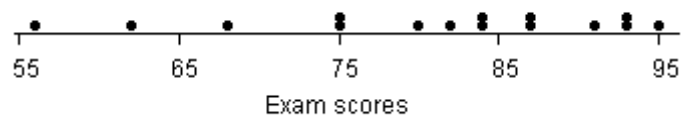
- 2.47 a. The figure below uses separate stems for last digits 0–4 and 5–9. That's not imperative, although doing so gives more detail of the shape of the data.

Figure for Exercise 2.47a

5	6
6	2
6	8
7	
7	55
8	02344
8	77
9	13
9	5

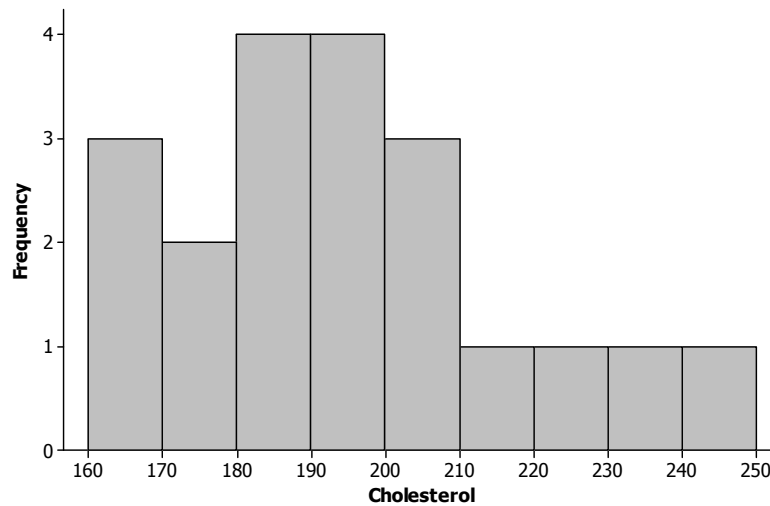
b.

Figure for Exercise 2.47b



- 2.48 a. The following histogram uses nine intervals. Others are possible.

Figure for Exercise 2.48a



- b. In the following stem-and-leaf, the first two digits of the values are used for stems (row labels).

Figure for Exercise 2.48b

```

| 16 | 046
| 17 | 88
| 18 | 2248
| 19 | 6888
| 20 | 046
| 21 | 2
| 22 | 2
| 23 | 0
| 24 | 2

```

- c. There are no obvious outliers.
d. The shape is difficult to judge—perhaps a slight skew to the right.
- 2.49 a. In the figure shown here, two stems have been used for each possible “tens” place in the number (values under 10 are an exception because the lowest value is about 6 inches). We rounded the 1995 total of 24.5 inches up to 25.

Figure for Exercise 2.49a

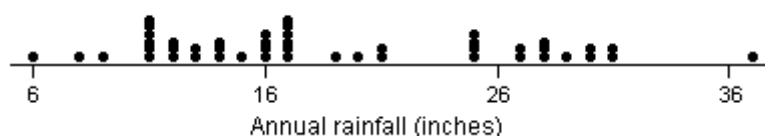
```

| 0 | 689
| 1 | 11111122233444
| 1 | 566667777779
| 2 | 011
| 2 | 5555778889
| 3 | 0011
| 3 | 7

```

- b. The dotplot follows.

Figure for Exercise 2.49b



- c. The data are skewed (but only slightly) to the right.

2.50 *Note:* The data for males are in the dataset **pennstate1M** on the companion website.

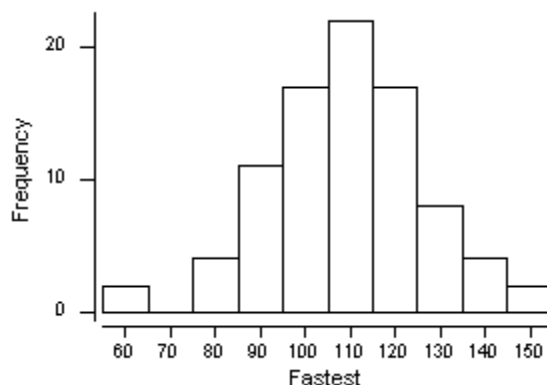
a.

Figure for Exercise 2.50a

5	5
6	0
7	
8	00005555
9	0000024555555
10	00000000012555555559
11	0000000000002555555
12	00000000004555555
13	00
14	00005
15	0

- b. The histogram shown here was done with Minitab. Notice that the midpoints of intervals are shown under the bars.

Figure for Exercise 2.50b

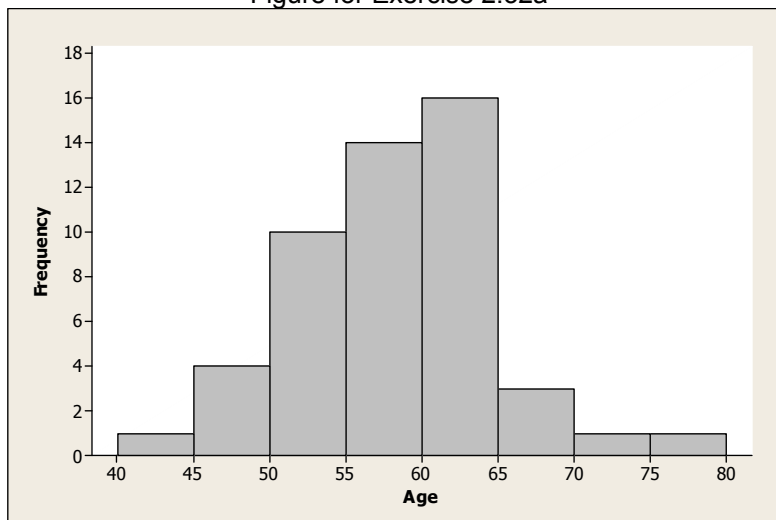


- c. The histogram may provide the best view of the overall shape and characteristics of the distribution. The information given in the stem-and-leaf plot is equivalent to that given in the histogram, but is somewhat harder to read. An advantage of the stem-and-leaf is that it's possible to determine individual values. The dotplot gives much information about individual values as well as the range and general location. It is, however, more difficult to judge the shape of the distribution with a dotplot.
- d. These data are symmetric in shape (and there may be two outliers).

2.51 Yes, a stem-and-leaf plot provides sufficient information to determine whether a dataset contains an outlier. Because all individual values are shown, it is possible to see whether any values are inconsistent with the bulk of the data.

- 2.52 a. The answer will vary due to the flexibility possible for deciding on the endpoints of intervals. The histogram shown here is based on 5-year age groups and the age under the left edge of a bar is included in that interval while the age under the right edge is not.

Figure for Exercise 2.52a



- b. Two stems should be used per decade of ages because there should be 6 to 15 stem values. [Note: If only one stem value is used per decade, there would be only 5 values, and if 5 were used per decade, there would be 22 (because only 4 would be needed for the 30s and 3 for the 70s).] Notice that if this stem-and-leaf plot were turned on its side, it would have the same shape as the histogram shown for part (a).

Figure for Exercise 2.52b

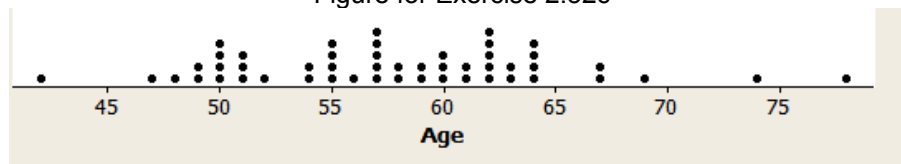
```

4 | 2
4 | 7899
5 | 0000111244
5 | 55556777778899
6 | 0001122222334444
6 | 779
7 | 4
7 | 8

```

c.

Figure for Exercise 2.52c



- d. The CEO ages may be skewed slightly to the left although this is not a pronounced skew. It may be fair to say that the data are roughly bell-shaped.
- e. There do not seem to be any obvious outliers although there is a bit of a gap between the oldest two CEOs and the third oldest and also a small gap between the youngest and the others.
- f. This is a conjecture, but it seems more likely that an outlier would occur in the salaries. A person's salary can become very high, and a company might conceivably pay an extraordinary salary to a successful executive.
- 2.53 Skewed to the left. Along a number line, values stretch more toward the left (small values).