

***Introduction to Probability and Statistics 15th edition
Mendenhall Solutions Manual***

SKU: 9781337554428-SOLUTIONS

Complete Solutions Manual
to Accompany

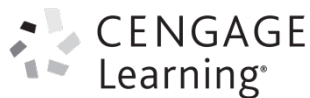
**Introduction to
Probability and Statistics**
15th Edition

William Mendenhall, III
1925-2009

Robert J. Beaver
University of California, Riverside, Emeritus

Barbara M. Beaver
University of California, Riverside, Emerita

Prepared by
Barbara M. Beaver



Australia • Brazil • Mexico • Singapore • United Kingdom • United States

© 2020 Cengage Learning

ALL RIGHTS RESERVED. No part of this work covered by the copyright herein may be reproduced, transmitted, stored, or used in any form or by any means graphic, electronic, or mechanical, including but not limited to photocopying, recording, scanning, digitizing, taping, Web distribution, information networks, or information storage and retrieval systems, except as permitted under Section 107 or 108 of the 1976 United States Copyright Act, without the prior written permission of the publisher except as may be permitted by the license terms below.

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

For permission to use material from this text or product, submit
all requests online at **www.cengage.com/permissions**
Further permissions questions can be emailed to
permissionrequest@cengage.com.

ISBN-13: 978-1-337-55829-7
ISBN-10: 1-337-55829-X

Cengage Learning
20 Channel Center Street
Boston, MA 02210
USA

Cengage Learning is a leading provider of customized learning solutions with office locations around the globe, including Singapore, the United Kingdom, Australia, Mexico, Brazil, and Japan. Locate your local office at:
www.cengage.com/global.

Cengage Learning products are represented in
Canada by Nelson Education, Ltd.

To learn more about Cengage Learning Solutions
or to purchase any of our products at our preferred
online store, 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.

Contents

Chapter 1: Describing Data with Graphs.....	1
Chapter 2: Describing Data with Numerical Measures.....	30
Chapter 3: Describing Bivariate Data.....	68
Chapter 4: Probability.....	93
Chapter 5: Discrete Probability Distribution.....	121
Chapter 6: The Normal Probability Distribution.....	156
Chapter 7: Sampling Distributions.....	186
Chapter 8: Large-Sample Estimation.....	210
Chapter 9: Large-Sample Test of Hypotheses.....	240
Chapter 10: Inference from Small Samples.....	271
Chapter 11: The Analysis of Variance.....	324
Chapter 12: Linear Regression and Correlation.....	364
Chapter 13: Multiple Regression Analysis.....	415
Chapter 14: The Analysis of Categorical Data.....	439
Chapter 15: Nonparametric Statistics.....	469

1: Describing Data with Graphs

Section 1.1

- 1.1.1 The experimental unit, the individual or object on which a variable is measured, is the student.
- 1.1.2 The experimental unit on which the number of errors is measured is the exam.
- 1.1.3 The experimental unit is the patient.
- 1.1.4 The experimental unit is the azalea plant.
- 1.1.5 The experimental unit is the car.
- 1.1.6 “Time to assemble” is a *quantitative* variable because a numerical quantity (1 hour, 1.5 hours, etc.) is measured.
- 1.1.7 “Number of students” is a *quantitative* variable because a numerical quantity (1, 2, etc.) is measured.
- 1.1.8 “Rating of a politician” is a *qualitative* variable since a quality (excellent, good, fair, poor) is measured.
- 1.1.9 “State of residence” is a *qualitative* variable since a quality (CA, MT, AL, etc.) is measured.
- 1.1.10 “Population” is a *discrete* variable because it can take on only integer values.
- 1.1.11 “Weight” is a *continuous* variable, taking on any values associated with an interval on the real line.
- 1.1.12 Number of claims is a *discrete* variable because it can take on only integer values.
- 1.1.13 “Number of consumers” is integer-valued and hence *discrete*.
- 1.1.14 “Number of boating accidents” is integer-valued and hence *discrete*.
- 1.1.15 “Time” is a *continuous* variable.
- 1.1.16 “Cost of a head of lettuce” is a *discrete* variable since money can be measured only in dollars and cents.
- 1.1.17 “Number of brothers and sisters” is integer-valued and hence *discrete*.
- 1.1.18 “Yield in bushels” is a *continuous* variable, taking on any values associated with an interval on the real line.
- 1.1.19 The statewide database contains a record of all drivers in the state of Michigan. The data collected represents the *population* of interest to the researcher.
- 1.1.20 The researcher is interested in the opinions of all citizens, not just the 1000 citizens that have been interviewed. The responses of these 1000 citizens represent a *sample*.
- 1.1.21 The researcher is interested in the weight gain of all animals that might be put on this diet, not just the twenty animals that have been observed. The responses of these twenty animals is a *sample*.
- 1.1.22 The data from the Internal Revenue Service contains the records of all wage earners in the United States. The data collected represents the *population* of interest to the researcher.
- 1.1.23
 - a The experimental unit, the item or object on which variables are measured, is the vehicle.
 - b Type (qualitative); make (qualitative); carpool or not? (qualitative); one-way commute distance (quantitative continuous); age of vehicle (quantitative continuous)
 - c Since five variables have been measured, this is *multivariate data*.
- 1.1.24
 - a The set of ages at death represents a population, because there have only been 38 different presidents in the United States history.
 - b The variable being measured is the continuous variable “age”.
 - c “Age” is a quantitative variable.

1.1.25 a The population of interest consists of voter opinions (for or against the candidate) at the time of the election for all persons voting in the election.

b Note that when a sample is taken (at some time prior or the election), we are not actually sampling from the population of interest. As time passes, voter opinions change. Hence, the population of voter opinions changes with time, and the sample may not be representative of the population of interest.

1.1.26 a-b The variable “survival time” is a quantitative continuous variable.

c The population of interest is the population of survival times for all patients having a particular type of cancer and having undergone a particular type of radiotherapy.

d-e Note that there is a problem with sampling in this situation. If we sample from all patients having cancer and radiotherapy, some may still be living and their survival time will not be measurable. Hence, we cannot sample directly from the population of interest, but must arrive at some reasonable alternate population from which to sample.

1.1.27 a The variable “reading score” is a quantitative variable, which is probably integer-valued and hence discrete.

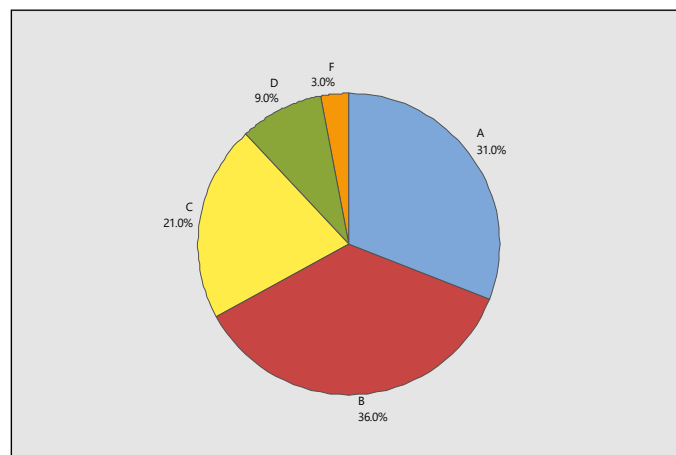
b The individual on which the variable is measured is the student.

c The population is hypothetical – it does not exist in fact – but consists of the reading scores for all students who could possibly be taught by this method.

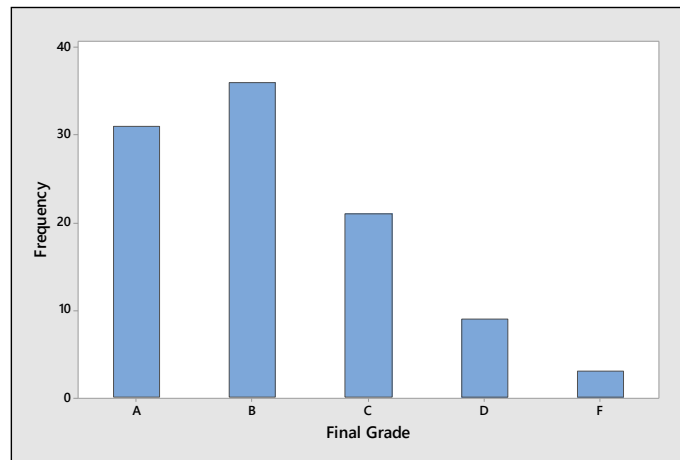
Section 1.2

1.2.1 The pie chart is constructed by partitioning the circle into five parts, according to the total contributed by each part. Since the total number of students is 100, the total number receiving a final grade of A represents $31/100 = 0.31$ or 31% of the total. Thus, this category will be represented by a sector angle of $0.31(360) = 111.6^\circ$. The other sector angles are shown next, along with the pie chart.

Final Grade	Frequency	Fraction of Total	Sector Angle
A	31	.31	111.6
B	36	.36	129.6
C	21	.21	75.6
D	9	.09	32.4
F	3	.03	10.8

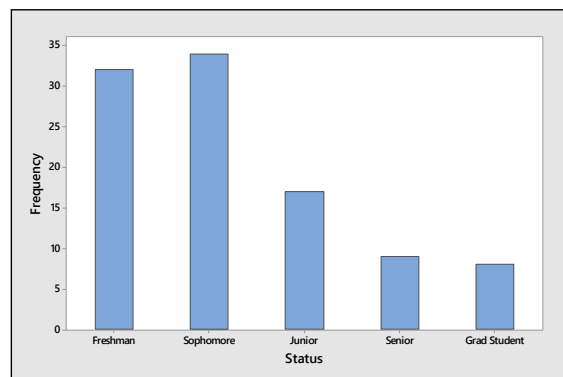
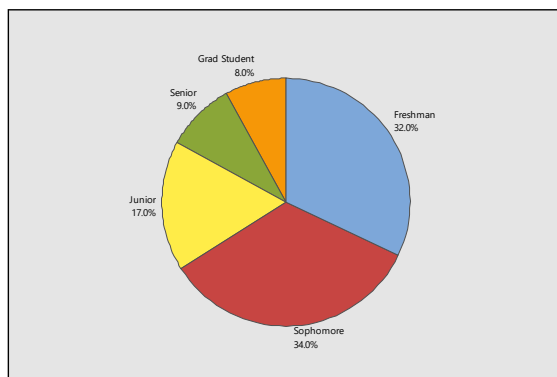


The bar chart represents each category as a bar with height equal to the frequency of occurrence of that category and is shown in the figure that follows.



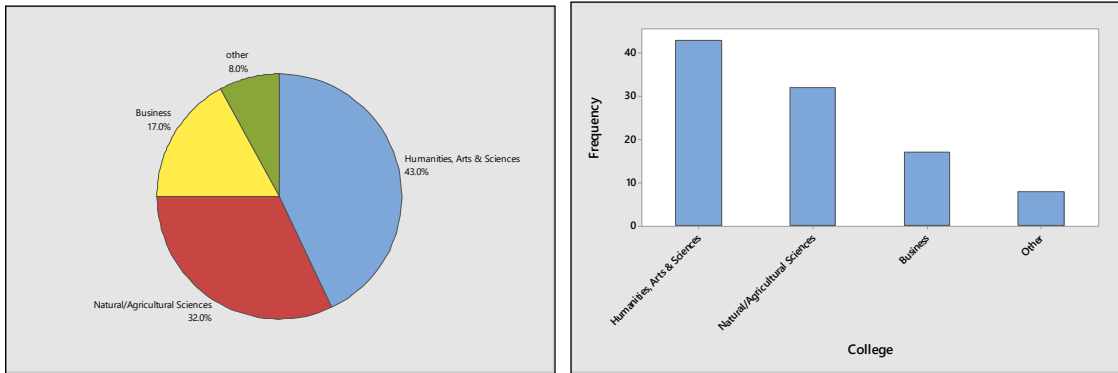
- 1.2.2 Construct a statistical table to summarize the data. The pie and bar charts are shown in the figures that follow.

Status	Frequency	Fraction of Total	Sector Angle
Freshman	32	.32	115.2
Sophomore	34	.34	122.4
Junior	17	.17	61.2
Senior	9	.09	32.4
Grad Student	8	.08	28.8



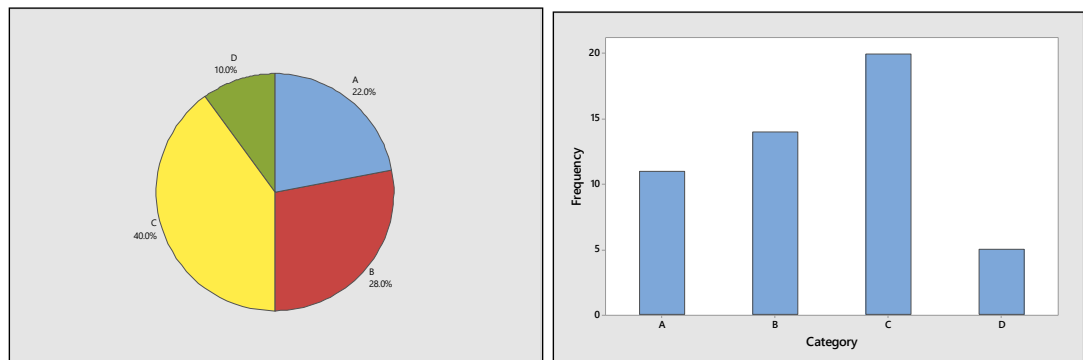
- 1.2.3 Construct a statistical table to summarize the data. The pie and bar charts are shown in the figures that follow.

Status	Frequency	Fraction of Total	Sector Angle
Humanities, Arts & Sciences	43	.43	154.8
Natural/Agricultural Sciences	32	.32	115.2
Business	17	.17	61.2
Other	8	.08	28.8



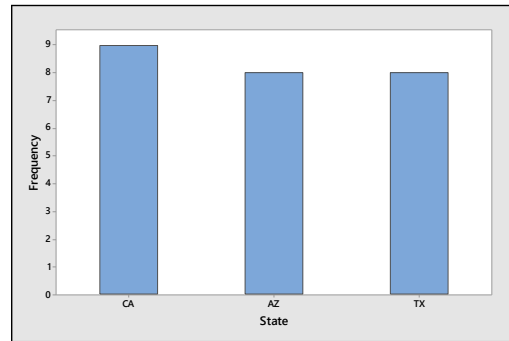
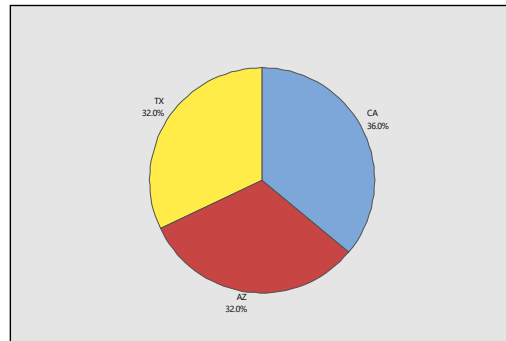
- 1.2.4 a** The pie chart is constructed by partitioning the circle into four parts, according to the total contributed by each part. Since the total number of people is 50, the total number in category A represents $11/50 = 0.22$ or 22% of the total. Thus, this category will be represented by a sector angle of $0.22(360) = 79.2^\circ$. The other sector angles are shown below. The pie chart is shown in the figure that follows.

Category	Frequency	Fraction of Total	Sector Angle
A	11	.22	79.2
B	14	.28	100.8
C	20	.40	144.0
D	5	.10	36.0



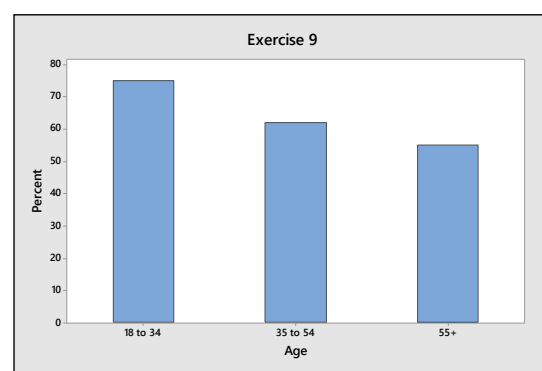
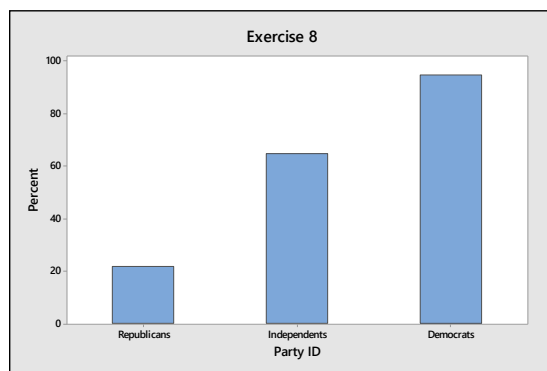
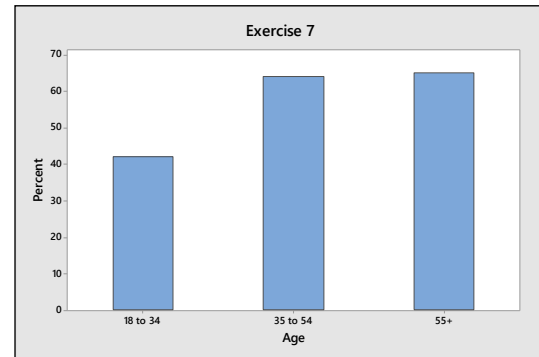
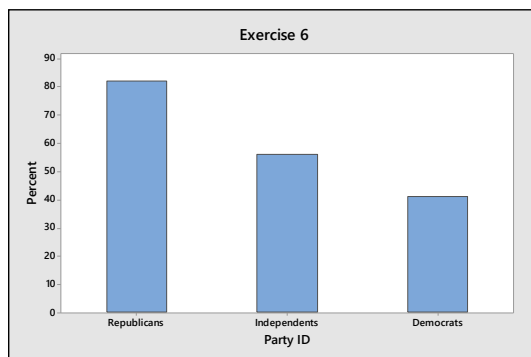
- b** The bar chart represents each category as a bar with height equal to the frequency of occurrence of that category and is shown in the figure above.
- c** Yes, the shape will change depending on the order of presentation. The order is unimportant.
- d** The proportion of people in categories B, C, or D is found by summing the frequencies in those three categories, and dividing by $n = 50$. That is, $(14 + 20 + 5)/50 = 0.78$.
- e** Since there are 14 people in category B, there are $50 - 14 = 36$ who are not, and the percentage is calculated as $(36/50)100 = 72\%$.
- 1.2.5 a-b** Construct a statistical table to summarize the data. The pie and bar charts are shown in the figures that follow.

State	Frequency	Fraction of Total	Sector Angle
CA	9	.36	129.6
AZ	8	.32	115.2
TX	8	.32	115.2

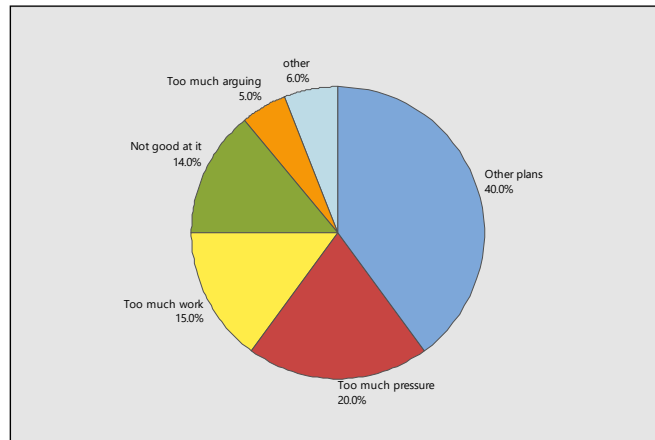


- c** From the table or the chart, Texas produced $8/25 = 0.32$ of the jeans.
- d** The highest bar represents California, which produced the most pairs of jeans.
- e** Since the bars and the sectors are almost equal in size, the three states produced roughly the same number of pairs of jeans.

1.2.6-9 The bar charts represent each category as a bar with height equal to the frequency of occurrence of that category.



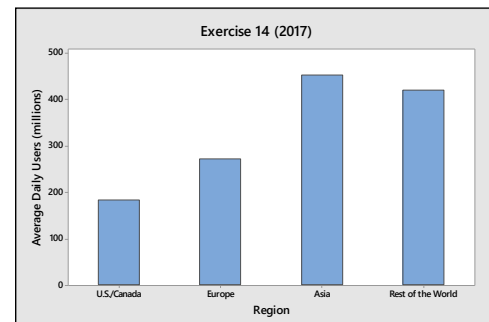
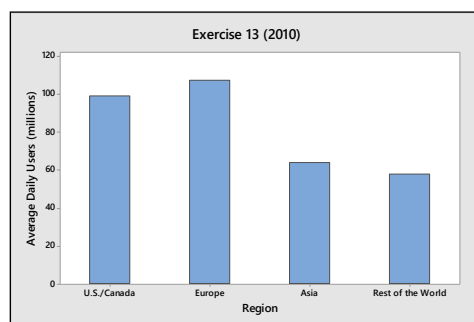
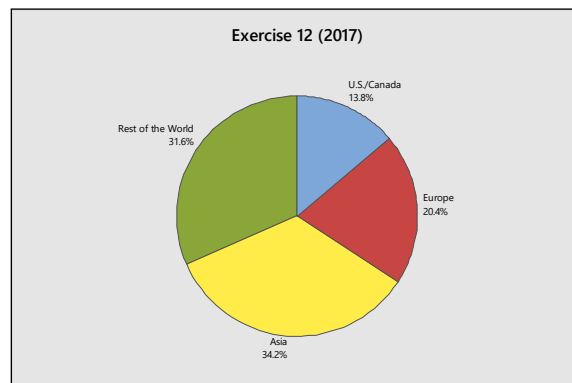
- 1.2.10** Answers will vary.
- 1.2.11 a** The percentages given in the exercise only add to 94%. We should add another category called “Other”, which will account for the other 6% of the responses.
- b** Either type of chart is appropriate. Since the data is already presented as percentages of the whole group, we choose to use a pie chart, shown in the figure that follows.



c-d Answers will vary.

1.2.12-14 The percentages falling in each of the four categories in 2017 are shown next (in parentheses), and the pie chart for 2017 and bar charts for 2010 and 2017 follow.

Region	2010	2017
United States/Canada	99	183 (13.8%)
Europe	107	271 (20.4%)
Asia	64	453 (34.2%)
Rest of the World	58	419 (31.6%)
Total	328	1326 (100%)

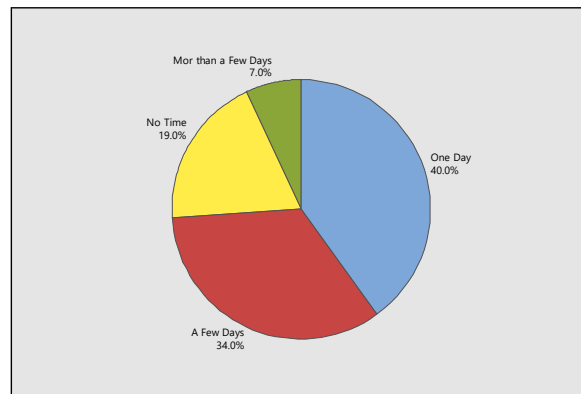


1.2.15 Users in Asia and the rest of the world have increased more rapidly than those in the U.S., Canada or Europe over the seven-year period.

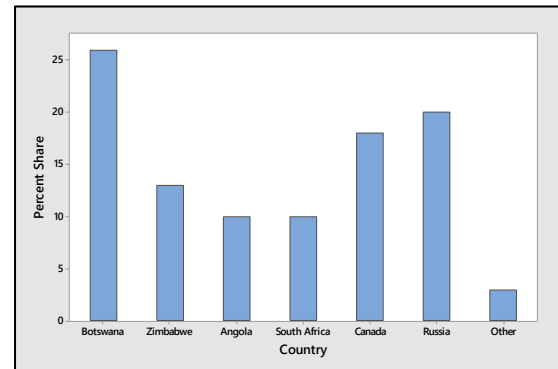
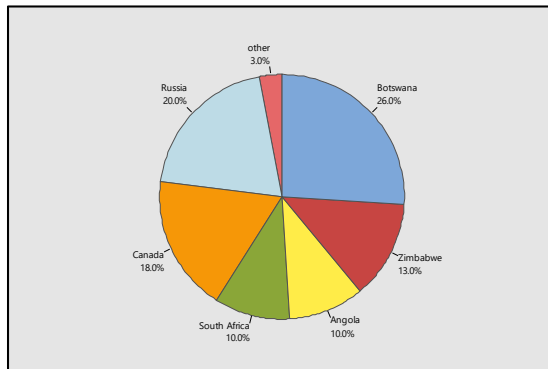
1.2.16 a The total percentage of responses given in the table is only $(40 + 34 + 19)\% = 93\%$. Hence there are 7% of the opinions not recorded, which should go into a category called “Other” or “More than a few days”.

b Yes. The bars are very close to the correct proportions.

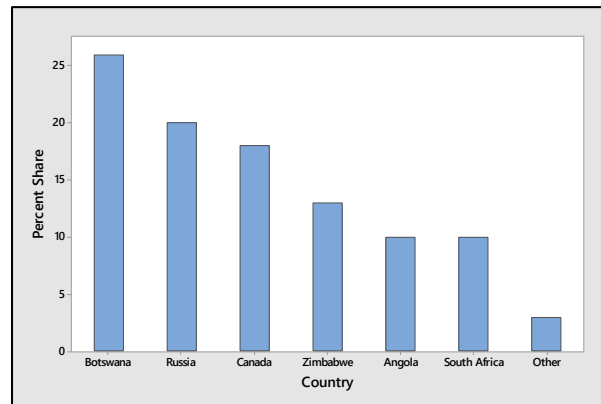
c Similar to previous exercises. The pie chart is shown next. The bar chart is probably more interesting to look at.



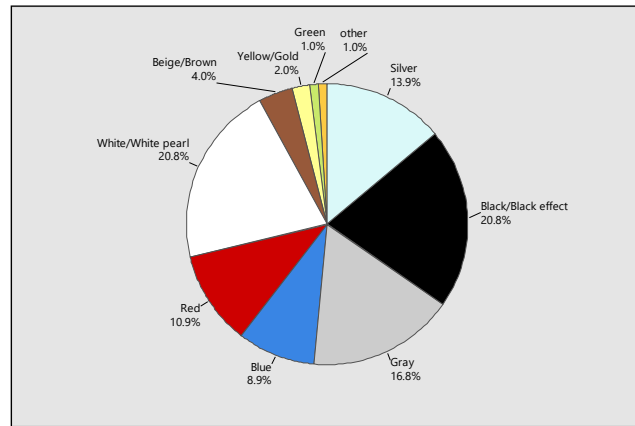
1.2.17-18 Answers will vary from student to student. Since the graph gives a range of values for Zimbabwe’s share, we have chosen to use the 13% figure, and have used 3% in the “Other” category. The pie chart and bar charts are shown next.



1.2.19-20 The Pareto chart is shown below. The Pareto chart is more effective than the bar chart or the pie chart.

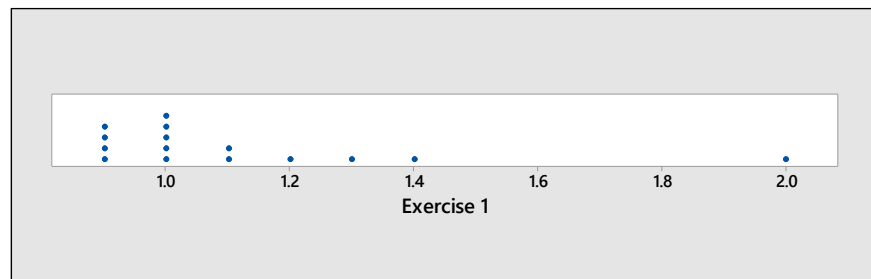


1.2.21 The data should be displayed with either a bar chart or a pie chart. The pie chart is shown next.

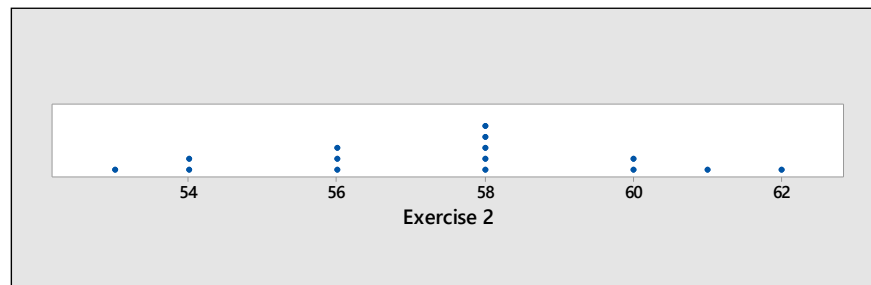


Section 1.3

1.3.1 The dotplot is shown next; the data is skewed right, with one outlier, $x = 2.0$.



1.3.2 The dotplot is shown next; the data is relatively mound-shaped, with no outliers.



1.3.3-5 The most obvious choice of a stem is to use the ones digit. The portion of the observation to the right of the ones digit constitutes the leaf. Observations are classified by row according to stem and also within each stem according to relative magnitude. The stem and leaf display is shown next.

```

1  6 8
2  1 2 5 5 5 7 8 8 9 9
3  1 1 4 5 5 6 6 6 7 7 7 7 8 9 9 9      leaf digit = 0.1
4  0 0 0 1 2 2 3 4 5 6 7 8 9 9 9      1 2 represents 1.2
5  1 1 6 6 7
6  1 2

```

- The stem and leaf display has a mound shaped distribution, with no outliers.
- From the stem and leaf display, the smallest observation is 1.6 (1 6).
- The eight and ninth largest observations are both 4.9 (4 9).

- 1.3.6** The stem is chosen as the ones digit, and the portion of the observation to the right of the ones digit is the leaf.

3 | 2 3 4 5 5 5 6 6 7 9 9 9 9

4 | 0 0 2 2 3 3 3 4 4 5 8 leaf digit = 0.1 1 2 represents 1.2

- 1.3.7-8** The stems are split, with the leaf digits 0 to 4 belonging to the first part of the stem and the leaf digits 5 to 9 belonging to the second. The stem and leaf display shown below improves the presentation of the data.

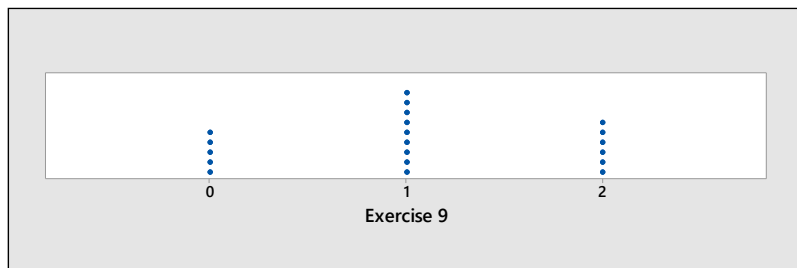
3 | 2 3 4

3 | 5 5 5 6 6 7 9 9 9 9 leaf digit = 0.1 1 2 represents 1.2

4 | 0 0 2 2 3 3 3 4 4

4 | 5 8

- 1.3.9** The scale is drawn on the horizontal axis and the measurements are represented by dots.



- 1.3.10** Since there is only one digit in each measurement, the ones digit must be the stem, and the leaf will be a zero digit for each measurement.

0 | 0 0 0 0 0

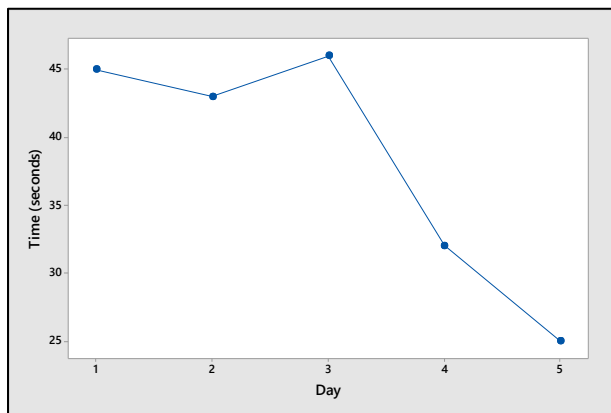
1 | 0 0 0 0 0 0 0 0 0

2 | 0 0 0 0 0 0

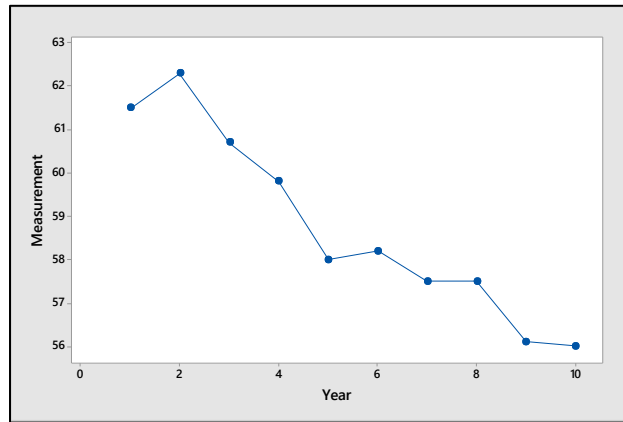
- 1.3.11** The distribution is relatively mound-shaped, with no outliers.

- 1.3.12** The two plots convey the same information if the stem and leaf plot is turned 90° and stretched to resemble the dotplot.

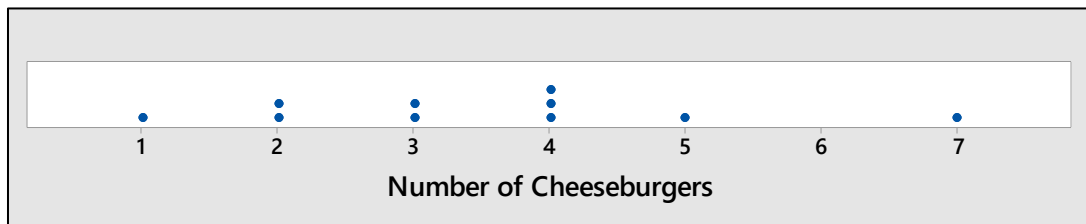
- 1.3.13** The line chart plots “day” on the horizontal axis and “time” on the vertical axis. The line chart shown next reveals that learning is taking place, since the time decreases each successive day.



- 1.3.14** The line graph is shown next. Notice the change in y as x increases. The measurements are decreasing over time.

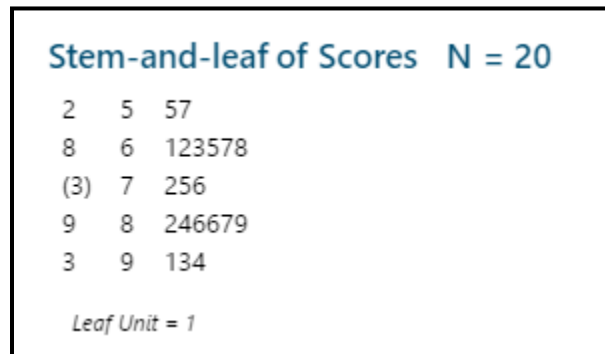


1.3.15 The dotplot is shown next.



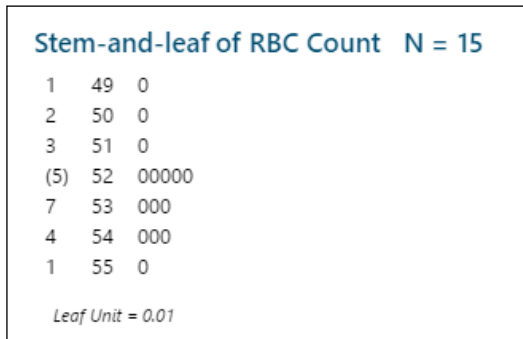
- a** The distribution is somewhat mound-shaped (as much as a small set can be); there are no outliers.
b $2/10 = 0.2$

1.3.16 a The test scores are graphed using a stem and leaf plot generated by *Minitab*.



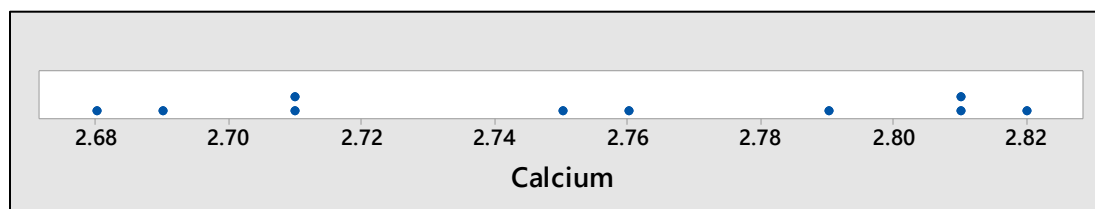
b-c The distribution is not mound-shaped, but is rather has two peaks centered around the scores 65 and 85. This might indicate that the students are divided into two groups – those who understand the material and do well on exams, and those who do not have a thorough command of the material.

1.3.17 a We choose a stem and leaf plot, using the ones and tenths place as the stem, and a zero digit as the leaf. The *Minitab* printout is shown next.



- b** The data set is relatively mound-shaped, centered at 5.2.
- c** The value $x = 5.7$ does not fall within the range of the other cell counts, and would be considered somewhat unusual.

1.3.18 a-b The dotplot and the stem and leaf plot are drawn using *Minitab*.



Stem-and-leaf of Calcium N = 10

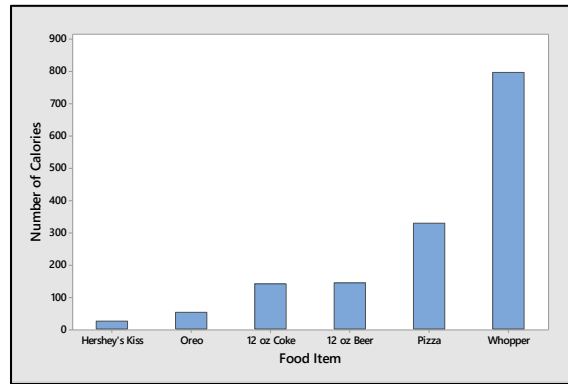
2	26	89
(5)	27	11569
3	28	112

Leaf Unit = 0.01

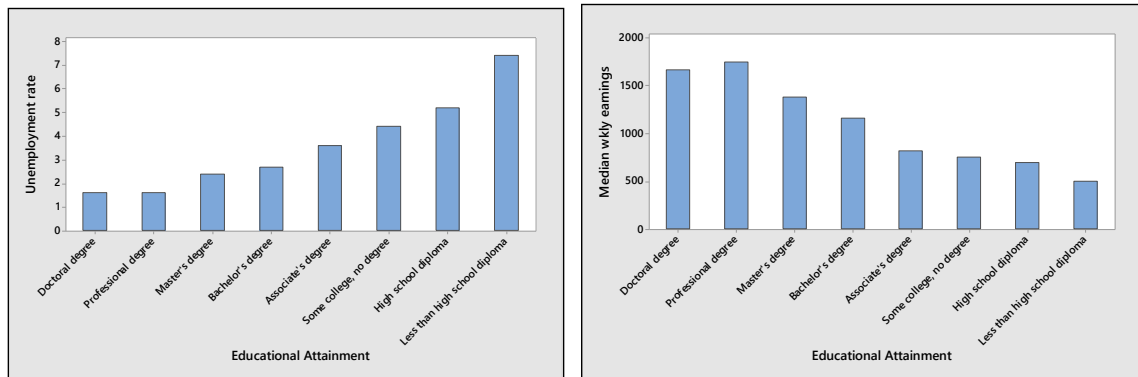
- c** The measurements all seem to be within the same range of variability. There do not appear to be any outliers.
- 1.3.19 a** Stem and leaf displays may vary from student to student. The most obvious choice is to use the tens digit as the stem and the ones digit as the leaf.

7		8 9
8		0 1 7
9		0 1 2 4 4 5 6 6 6 8 8
10		1 7 9
11		2

- b** The display is fairly mound-shaped, with a large peak in the middle.
- 1.3.20 a** The sizes and volumes of the food items do increase as the number of calories increase, but not in the correct proportion to the actual calories. The differences in calorie content are not accurately portrayed in the graph.
- b** The bar graph which accurately portrays the number of calories in the six food items is shown next.

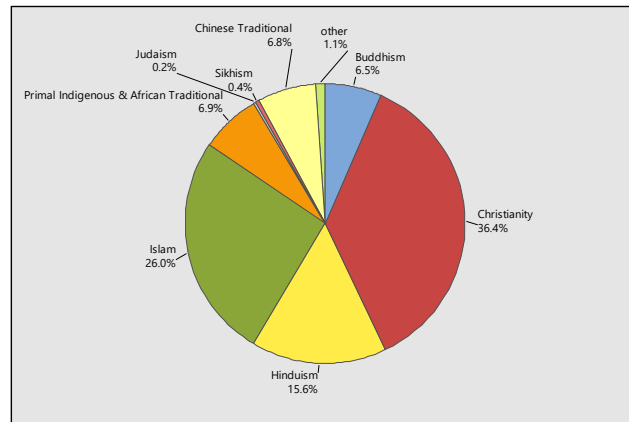


1.3.21 a-b The bar charts for the median weekly earnings and unemployment rates for eight different levels of education are shown next.

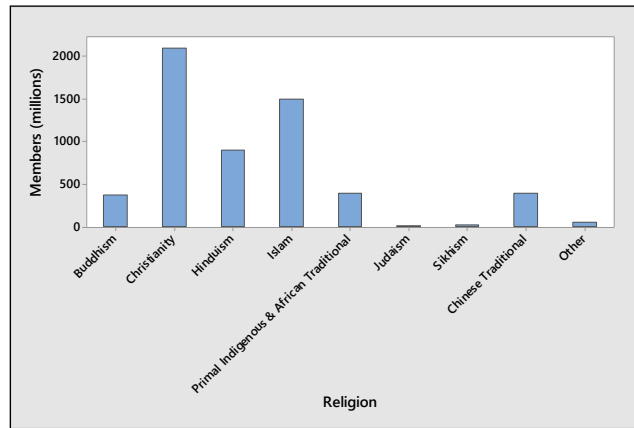


c The unemployment rate drops and the median weekly earnings rise as the level of educational attainment increases.

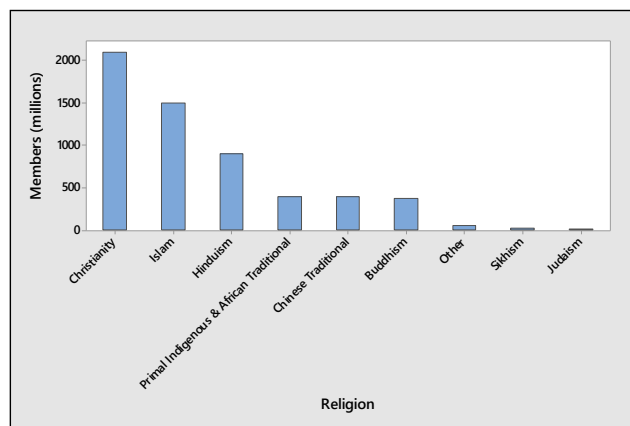
1.3.22 a Similar to previous exercises. The pie chart is shown next.



b The bar chart is shown next.



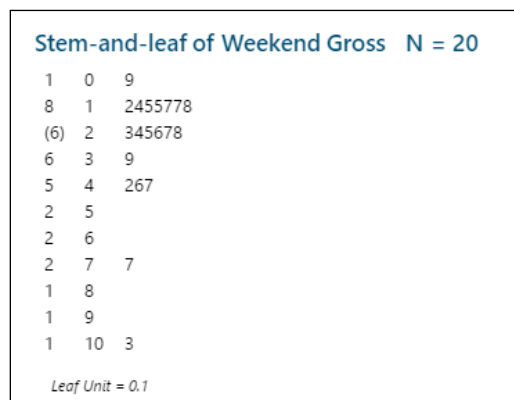
- c** The Pareto chart is a bar chart with the heights of the bars ordered from large to small. This display is more effective than the pie chart.



- 1.3.23 a** The distribution is skewed to the right, with a several unusually large measurements. The five states marked as HI are California, New Jersey, New York and Pennsylvania.

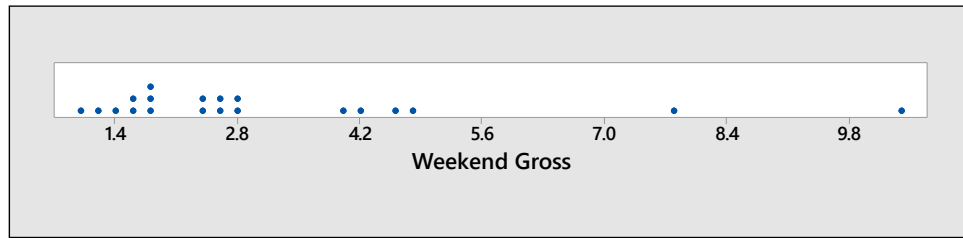
b Three of the four states are quite large in area, which might explain the large number of hazardous waste sites. However, New Jersey is relatively small, and other large states do not have unusually large number of waste sites. The pattern is not clear.

- 1.3.24 a**



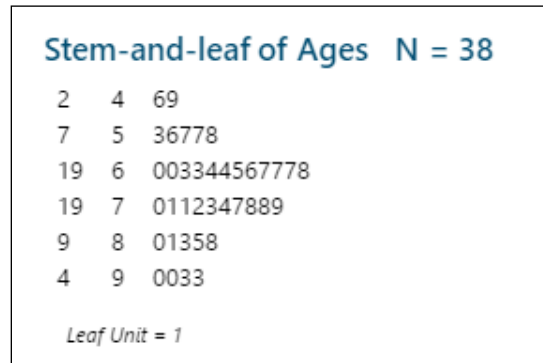
The distribution is skewed to the right, with two outliers.

b The dotplot is shown next. It conveys nearly the same information, but the stem-and-leaf plot may be more informative.



1.3.25 a Answers will vary.

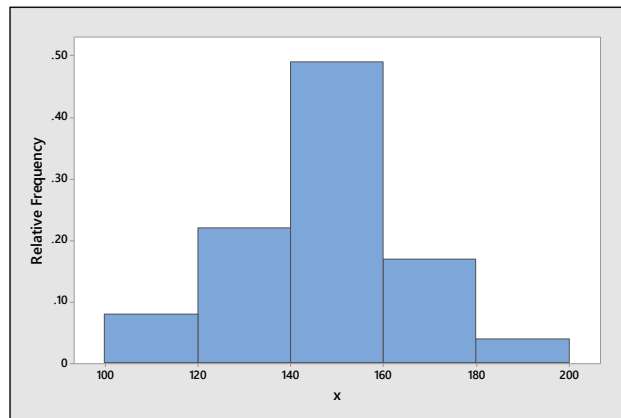
b The stem and leaf plot is constructed using the tens place as the stem and the ones place as the leaf. Notice that the distribution is roughly mound-shaped.



c-d Three of the five youngest presidents – Kennedy, Lincoln and Garfield – were assassinated while in office. This would explain the fact that their ages at death were in the lower tail of the distribution.

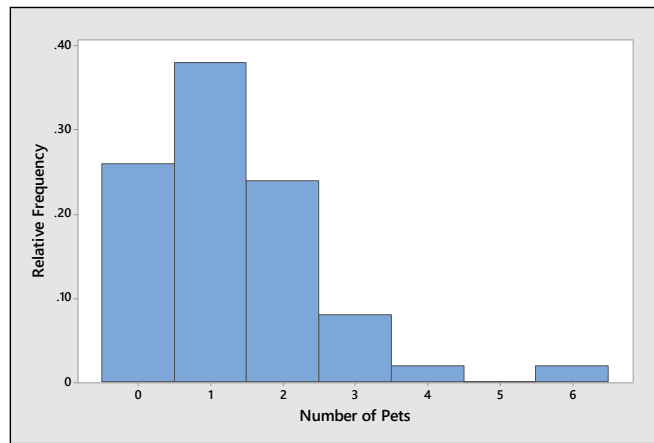
Section 1.4

1.4.1 The relative frequency histogram displays the relative frequency as the height of the bar over the appropriate class interval and is shown next. The distribution is relatively mound-shaped.



1.4.2 Since the variable of interest can only take integer values, the classes can be chosen as the values 0, 1, 2, 3, 4, 5 and 6. The table containing the classes, their corresponding frequencies and their relative frequencies and the relative frequency histogram are shown next. The distribution is skewed to the right.

Number of Household Pets	Frequency	Relative Frequency
0	13	$13/50 = .26$
1	19	$19/50 = .38$
2	12	$12/50 = .24$
3	4	$4/50 = .08$
4	1	$1/50 = .02$
5	0	$0/50 = .00$
6	1	$1/50 = .02$
Total	50	$50/50 = 1.00$



1.4.3-8 The proportion of measurements falling in each interval is equal to the sum of the heights of the bars over that interval. Remember that the lower class boundary is included, but not the upper class boundary.

3. $.20 + .40 + .15 = .75$ 4. $.05 + .15 + .20 = .40$
5. $.05$ 6. $.40 + .15 = .55$
7. $.15$ 8. $.05 + .15 + .20 = .40$

1.4.9 Answers will vary. The range of the data is $110 - 10 = 90$ and we need to use seven classes. Calculate $90/7 = 12.86$ which we choose to round up to 15. Convenient class boundaries are created, starting at 10: 10 to < 25, 25 to < 40, ..., 100 to < 115.

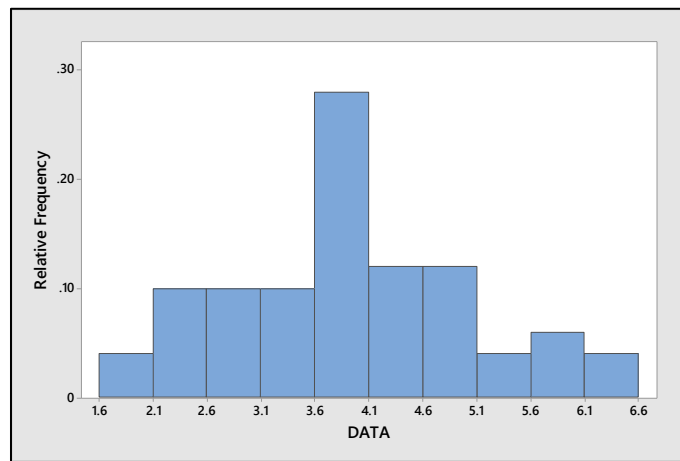
1.4.10 Answers will vary. The range of the data is $76.8 - 25.5 = 51.3$ and we need to use six classes. Calculate $51.3/6 = 8.55$ which we choose to round up to 9. Convenient class boundaries are created, starting at 25: 25 to < 34, 34 to < 43, ..., 70 to < 79.

1.4.11 Answers will vary. The range of the data is $1.73 - .31 = 1.42$ and we need to use ten classes. Calculate $1.42/10 = .142$ which we choose to round up to .15. Convenient class boundaries are created, starting at .30: .30 to < .45, .45 to < .60, ..., 1.65 to < 1.80.

1.4.12 Answers will vary. The range of the data is $192 - 0 = 192$ and we need to use eight classes. Calculate $192/8 = 24$ which we choose to round up to 25. Convenient class boundaries are created, starting at 0: 0 to < 25, 25 to < 50, ..., 175 to < 200.

1.4.13-16 The table containing the classes, their corresponding frequencies and their relative frequencies and the relative frequency histogram are shown next.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	1.6 to < 2.1	11	2	.04
2	2.1 to < 2.6	11111	5	.10
3	2.6 to < 3.1	11111	5	.10
4	3.1 to < 3.6	11111	5	.10
5	3.6 to < 4.1	11111 11111 1111	14	.28
6	4.1 to < 4.6	11111 11	7	.14
7	4.6 to < 5.1	11111	5	.10
8	5.1 to < 5.6	11	2	.04
9	5.6 to < 6.1	111	3	.06
10	6.1 to < 6.6	11	2	.04



13. The distribution is roughly mound-shaped.

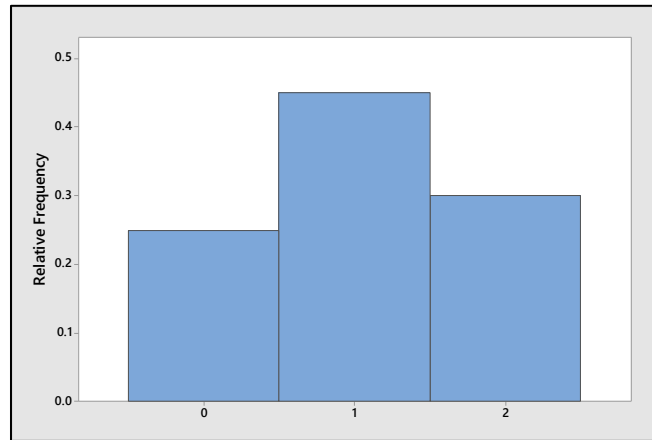
14. The fraction less than 5.1 is that fraction lying in classes 1-7, or $(2+5+\cdots+7+5)/50 = 43/50 = 0.86$.

15. The fraction larger than 3.6 lies in classes 5-10, or $(14+7+\cdots+3+2)/50 = 33/50 = 0.66$.

16. The fraction from 2.6 up to but not including 4.6 lies in classes 3-6, or $(5+5+14+7)/50 = 31/50 = 0.62$.

1.4.17-20 Since the variable of interest can only take the values 0, 1, or 2, the classes can be chosen as the integer values 0, 1, and 2. The table shows the classes, their corresponding frequencies and their relative frequencies. The relative frequency histogram follows the table.

Value	Frequency	Relative Frequency
0	5	.25
1	9	.45
2	6	.30



17. Using the table above, the proportion of measurements greater than 1 is the same as the proportion of “2”s, or 0.30.

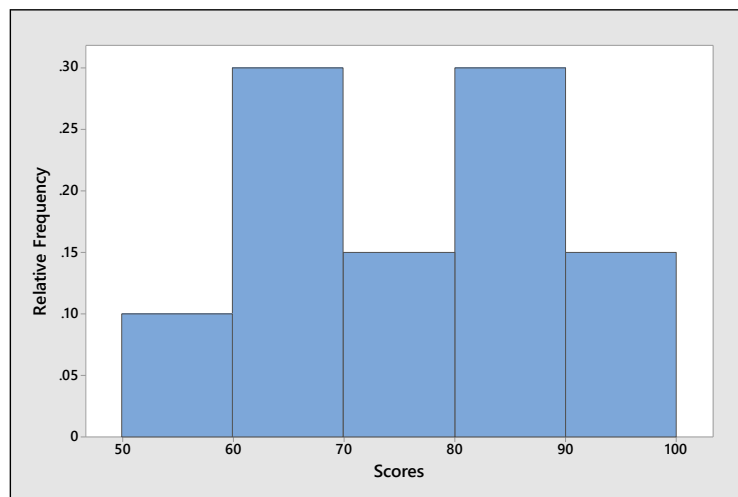
18. The proportion of measurements less than 2 is the same as the proportion of “0”s and “1”s, or $0.25 + 0.45 = .70$.

19. The probability of selecting a “2” in a random selection from these twenty measurements is $6/20 = .30$.

20. There are no outliers in this relatively symmetric, mound-shaped distribution.

1.4.21-23 Answers will vary. The range of the data is $94 - 55 = 39$ and we choose to use 5 classes. Calculate $39/5 = 7.8$ which we choose to round up to 10. Convenient class boundaries are created, starting at 50 and the table and relative frequency histogram are created.

Class Boundaries	Frequency	Relative Frequency
50 to < 60	2	.10
60 to < 70	6	.30
70 to < 80	3	.15
80 to < 90	6	.30
90 to < 100	3	.15



21. The distribution has two peaks at about 65 and 85. Depending on the way in which the student constructs the histogram, these peaks may or may not be clearly seen.

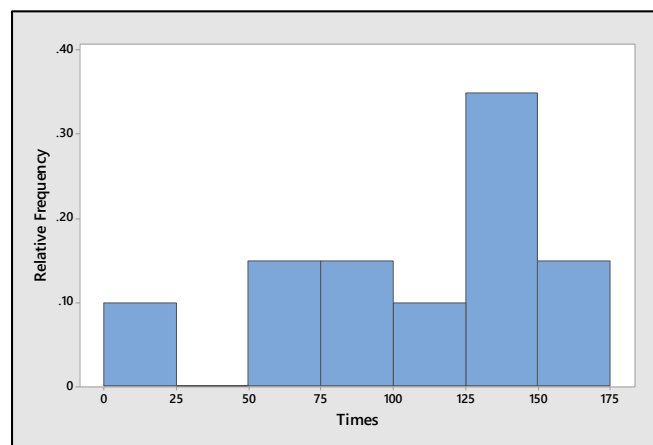
22. The shape is unusual. It might indicate that the students are divided into two groups – those who understand the material and do well on exams, and those who do not have a thorough command of the material.

23. The shapes are roughly the same, but this may not be the case if the student constructs the histogram using different class boundaries.

1.4.24 a There are a few extremely small numbers, indicating that the distribution is probably skewed to the left.

b The range of the data $165 - 8 = 157$. We choose to use seven class intervals of length 25, with subintervals 0 to < 25 , 25 to < 50 , 50 to < 75 , and so on. The tally and relative frequency histogram are shown next.

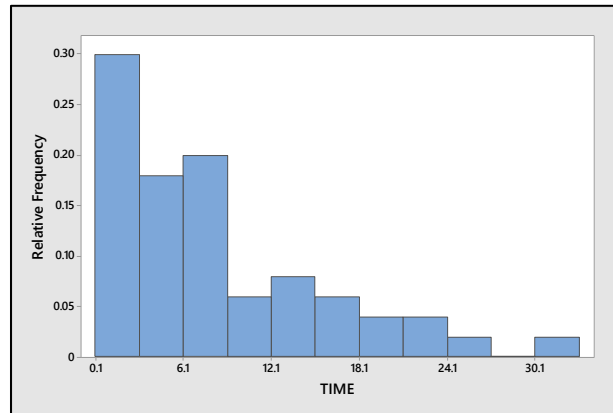
Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	0 to < 25	11	2	$2/20$
2	25 to < 50		0	$0/20$
3	50 to < 75	111	3	$3/20$
4	75 to < 100	111	3	$3/20$
5	100 to < 125	11	2	$2/20$
6	125 to < 150	11111 11	7	$7/20$
7	150 to < 175	111	3	$3/20$



c The distribution is indeed skewed left with two possible outliers: $x = 8$ and $x = 11$.

1.4.25 a The range of the data $32.3 - 0.2 = 32.1$. We choose to use eleven class intervals of length 3 ($32.1/11 = 2.9$, which when rounded to the next largest integer is 3). The subintervals 0.1 to < 3.1 , 3.1 to < 6.1 , 6.1 to < 9.1 , and so on, are convenient and the tally and relative frequency histogram are shown next.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	0.1 to < 3.1	11111 11111 11111	15	$15/50$
2	3.1 to < 6.1	11111 1111	9	$9/50$
3	6.1 to < 9.1	11111 11111	10	$10/50$
4	9.1 to < 12.1	111	3	$3/50$
5	12.1 to < 15.1	1111	4	$4/50$
6	15.1 to < 18.1	111	3	$3/50$
7	18.1 to < 21.1	11	2	$2/50$
8	21.1 to < 24.1	11	2	$2/50$
9	24.1 to < 27.1	1	1	$1/50$
10	27.1 to < 30.1		0	$0/50$
11	30.1 to < 33.1	1	1	$1/50$

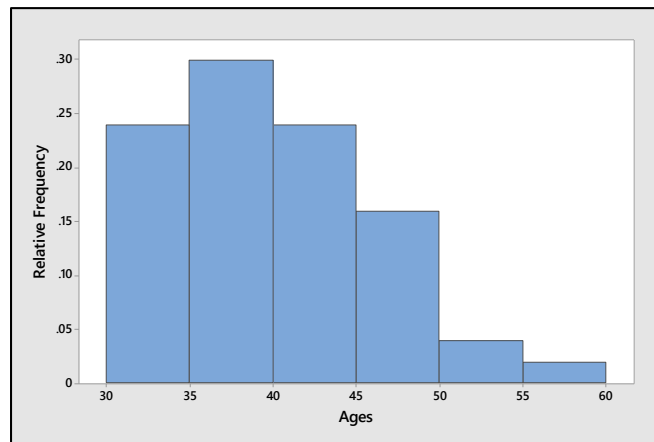


b The data is skewed to the right, with a few unusually large measurements.

c Looking at the data, we see that 36 patients had a disease recurrence within 10 months. Therefore, the fraction of recurrence times less than or equal to 10 is $36/50 = 0.72$.

1.4.26 a We use class intervals of length 5, beginning with the subinterval 30 to < 35. The tally and the relative frequency histogram are shown next.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	30 to < 35	11111 11111 11	12	12/50
2	35 to < 40	11111 11111 11111	15	15/50
3	40 to < 45	11111 11111 11	12	12/50
4	45 to < 50	11111 111	8	8/50
5	50 to < 55	11	2	2/50
6	55 to < 60	1	1	1/50



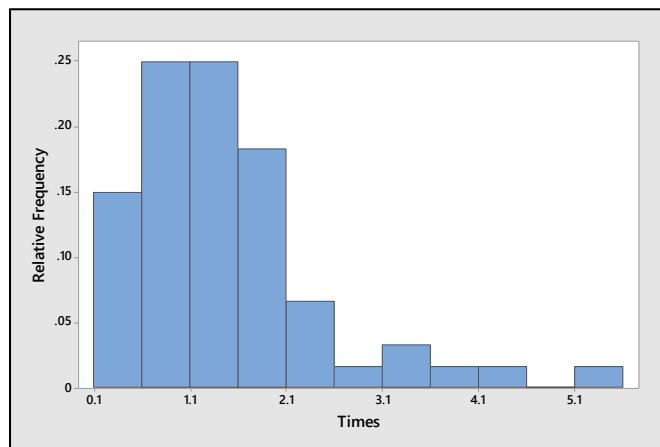
b Use the table or the relative frequency histogram. The proportion of children in the interval 35 to < 45 is $(15 + 12)/50 = .54$.

c The proportion of children aged less than 50 months is $(12 + 15 + 12 + 8)/50 = .94$.

1.4.27 a The data ranges from .2 to 5.2, or 5.0 units. Since the number of class intervals should be between five and twelve, we choose to use eleven class intervals, with each class interval having length 0.50 ($5.0/11 = .45$, which, rounded to the nearest convenient fraction, is .50). We must now select interval boundaries such that no measurement can fall on a boundary point. The subintervals .1 to < .6, .6 to < 1.1, and so on, are convenient and a tally is constructed.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	0.1 to < 0.6	11111 11111	10	.167
2	0.6 to < 1.1	11111 11111 11111	15	.250
3	1.1 to < 1.6	11111 11111 11111	15	.250
4	1.6 to < 2.1	11111 11111	10	.167
5	2.1 to < 2.6	1111	4	.067
6	2.6 to < 3.1	1	1	.017
7	3.1 to < 3.6	11	2	.033
8	3.6 to < 4.1	1	1	.017
9	4.1 to < 4.6	1	1	.017
10	4.6 to < 5.1		0	.000
11	5.1 to < 5.6	1	1	.017

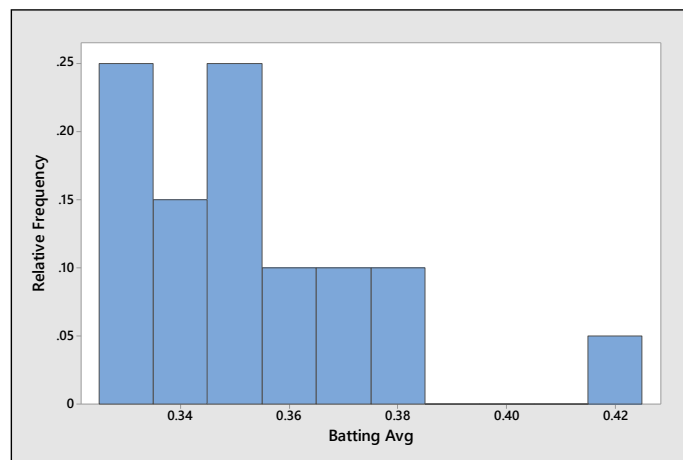
The relative frequency histogram is shown next.



b The distribution is skewed to the right, with several unusually large observations.

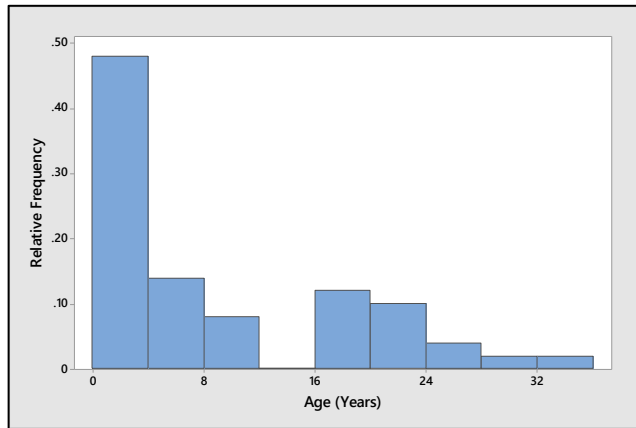
c For some reason, one person had to wait 5.2 minutes. Perhaps the supermarket was understaffed that day, or there may have been an unusually large number of customers in the store.

1.4.28 a Histograms will vary from student to student. A typical histogram generated by *Minitab* is shown next.

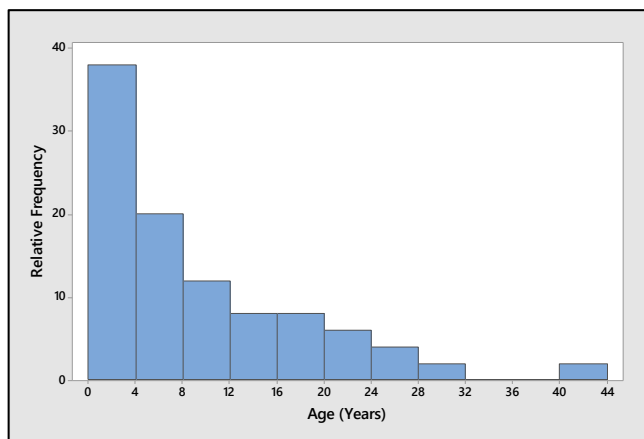


b Since 1 of the 20 players has an average above 0.400, the chance is 1 out of 20 or $1/20 = 0.05$.

1.4.29 a-b Answers will vary from student to student. The students should notice that the distribution is skewed to the right with a few pennies being unusually old. A typical histogram is shown next.

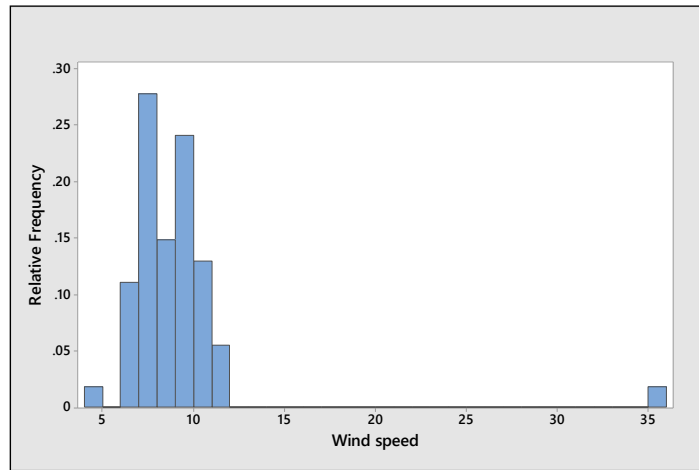


- 1.4.30 a** Answers will vary from student to student. A typical histogram is shown next. It looks very similar to the histogram from Exercise 1.4.29.



- b** There is one outlier, $x = 41$.
- 1.4.31 a** Answers will vary from student to student. The relative frequency histogram below was constructed using classes of length 1.0 starting at $x = 4$. The value $x = 35.1$ is not shown in the table but appears on the graph shown next.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	4.0 to < 5.0	1	1	1/54
2	5.0 to < 6.0	0	0	0/54
3	6.0 to < 7.0	11111 1	6	6/54
4	7.0 to < 8.0	11111 11111 11111	15	15/54
5	8.0 to < 9.0	11111 111	8	8/54
6	9.0 to < 10.0	11111 11111 111	13	13/54
7	10.0 to < 11.0	11111 11	7	7/54
8	11.0 to < 12.0	111	3	3/54



b Since Mt. Washington is a very mountainous area, it is not unusual that the average wind speed would be very high.

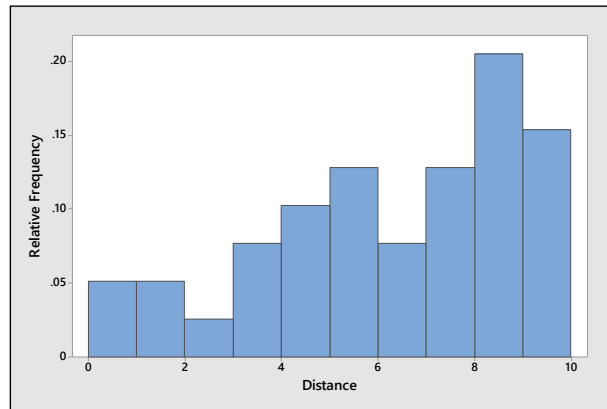
c The value $x = 9.9$ does not lie far from the center of the distribution (excluding $x = 35.1$). It would not be considered unusually high.

1.4.32 a-b The data is somewhat mound-shaped, but it appears to have two local peaks – high points from which the frequencies drop off on either side.

c Since these are student heights, the data can be divided into two groups – heights of males and heights of females. Both groups will have an approximate mound-shape, but the average female height will be lower than the average male height. When the two groups are combined into one data set, it causes a “mixture” of two mound-shaped distributions and produces the two peaks seen in the histogram.

1.4.33 a The relative frequency histogram below was constructed using classes of length 1.0 starting at $x = 0.0$.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	0.0 to < 1.0	11	2	2/39
2	1.0 to < 2.0	11	2	2/39
3	2.0 to < 3.0	1	1	1/39
4	3.0 to < 4.0	111	3	3/39
5	4.0 to < 5.0	111	4	4/39
6	5.0 to < 6.0	11111	5	5/39
7	6.0 to < 7.0	111	3	3/39
8	7.0 to < 8.0	11111	5	5/39
9	8.0 to < 9.0	11111 111	8	8/39
10	9.0 to < 10.0	11111 1	6	6/39

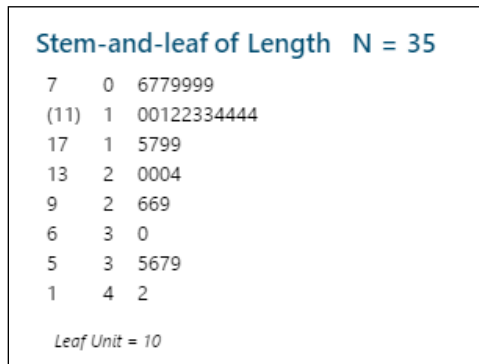


- a** The distribution is skewed to the left, with slightly higher frequency in the first two classes (within two miles of UCR).
- b** As the distance from UCR increases, each successive area increases in size, thus allowing for more Starbucks stores in that region.

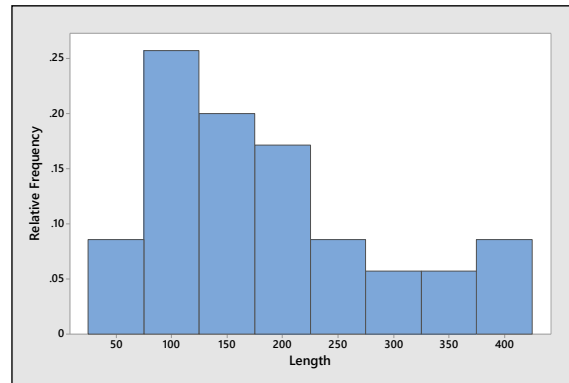
Reviewing What You've Learned

- 1.R.1**
- a** “Ethnic origin” is a *qualitative variable* since a quality (ethnic origin) is measured.
 - b** “Score” is a *quantitative variable* since a numerical quantity (0-100) is measured.
 - c** “Type of establishment” is a *qualitative variable* since a category (Carl’s Jr., McDonald’s or Burger King) is measured.
 - d** “Mercury concentration” is a *quantitative variable* since a numerical quantity is measured.
- 1.R.2** To determine whether a distribution is likely to be skewed, look for the likelihood of observing extremely large or extremely small values of the variable of interest.
- a** The price of an 8-oz can of peas is not likely to contain unusually large or small values.
 - b** Not likely to be skewed.
 - c** If a package is dropped, it is likely that all the shells will be broken. Hence, a few large number of broken shells is possible. The distribution will be skewed.
 - d** If an animal has one tick, he is likely to have more than one. There will be some “0”s with uninfected rabbits, and then a larger number of large values. The distribution will not be symmetric.
- 1.R.3**
- a** The length of time between arrivals at an outpatient clinic is a continuous random variable, since it can be any of the infinite number of positive real values.
 - b** The time required to finish an examination is a continuous random variable as was the random variable described in part **a**.
 - c** Weight is continuous, taking any positive real value.
 - d** Body temperature is continuous, taking any real value.
 - e** Number of people is discrete, taking the values 0, 1, 2, ...
- 1.R.4**
- a** Number of properties is discrete, taking the values 0, 1, 2, ...
 - b** Depth is continuous, taking any non-negative real value.
 - c** Length of time is continuous, taking any non-negative real value.
 - d** Number of aircraft is discrete.

1.R.5 a



b

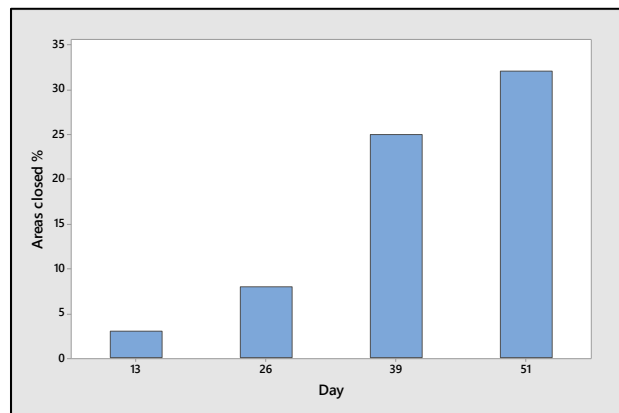
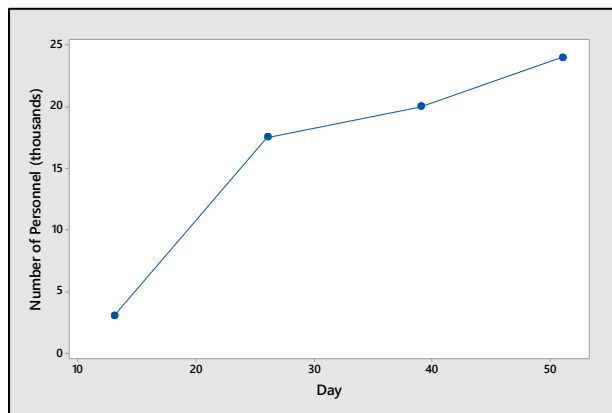


c These data are skewed right.

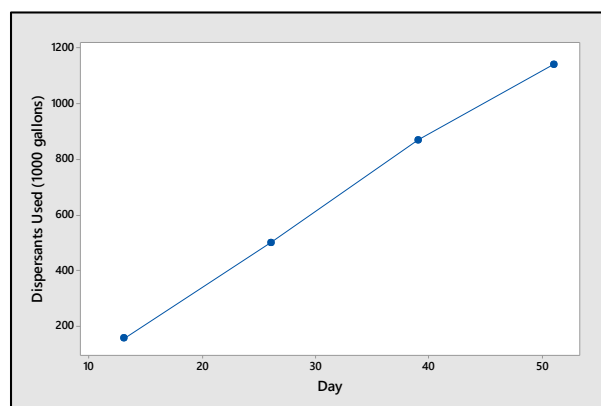
1.R.6 a The five quantitative variables are measured over time two months after the oil spill. Some sort of comparative bar charts (side-by-side or stacked) or a line chart should be used.

b As the time after the spill increases, the values of all five variables increase.

c-d The line chart for *number of personnel* and the bar chart for *fishing areas closed* are shown next.

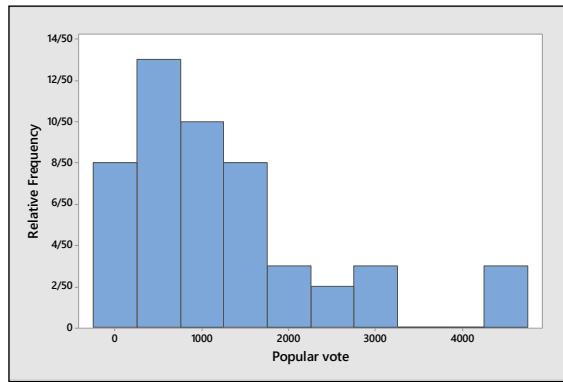


e The line chart for *amount of dispersants* is shown next. There appears to be a straight-line trend.

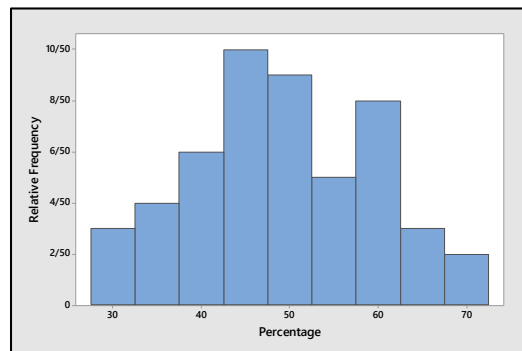


1.R.7 a The popular vote within each state should vary depending on the size of the state. Since there are several very large states (in population) in the United States, the distribution should be skewed to the right.

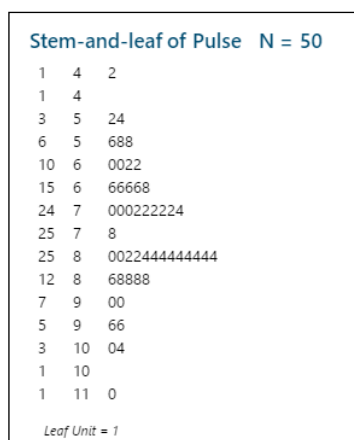
b-c Histograms will vary from student to student but should resemble the histogram generated by *Minitab* in the next figure. The distribution is indeed skewed to the right, with three “outliers” – California, Florida and Texas.



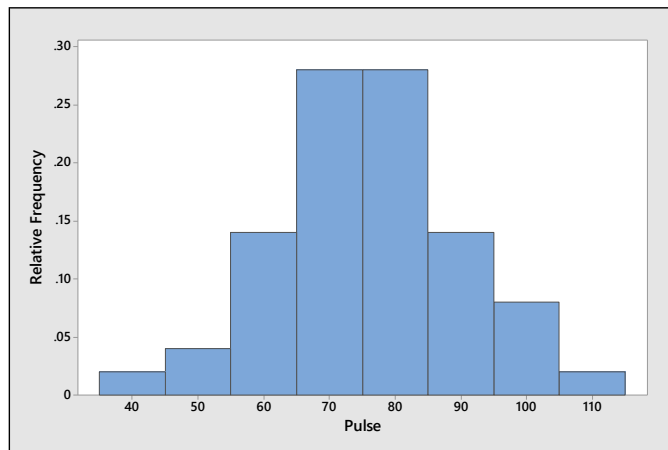
- 1.R.8 a-b** Once the size of the state is removed by calculating the percentage of the popular vote, the unusually large values in the Exercise 7 data set will disappear, and each state will be measured on an equal basis. Student histograms should resemble the histogram shown next. Notice the relatively mound-shape and the lack of any outliers.



- 1.R.9 a-b** Popular vote is skewed to the right while the percentage of popular vote is roughly mound-shaped. While the distribution of popular vote has outliers (California, Florida and Texas), there are no outliers in the distribution of percentage of popular vote. When the stem and leaf plots are turned 90°, the shapes are very similar to the histograms.
- c** Once the size of the state is removed by calculating the percentage of the popular vote, the unusually large values in the set of “popular votes” will disappear, and each state will be measured on an equal basis. The data then distribute themselves in a mound-shape around the average percentage of the popular vote.
- 1.R.10 a** The measurements are obtained by counting the number of beats for 30 seconds, and then multiplying by 2. Thus, the measurements should all be even numbers.
- b** The stem and leaf plot is shown next.

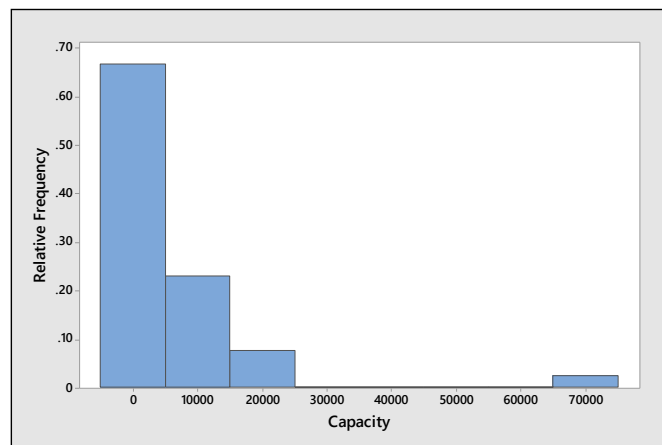


c Answers will vary. A typical histogram, generated by *Minitab*, is shown next.



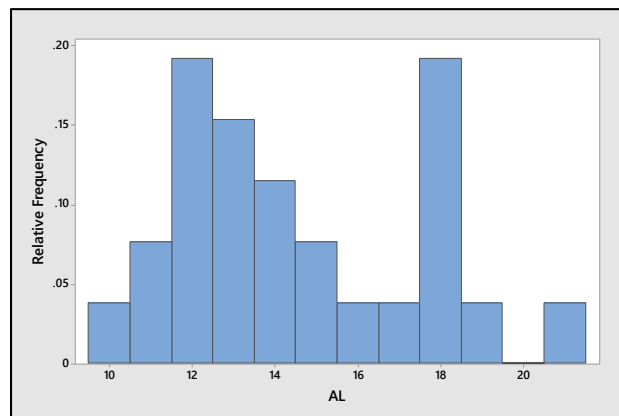
d The distribution of pulse rates is mound-shaped and relatively symmetric around a central location of 75 beats per minute. There are no outliers.

1.R.11 a-b Answers will vary from student to student. A typical histogram is shown next—the distribution is skewed to the right, with an extreme outlier (Texas).

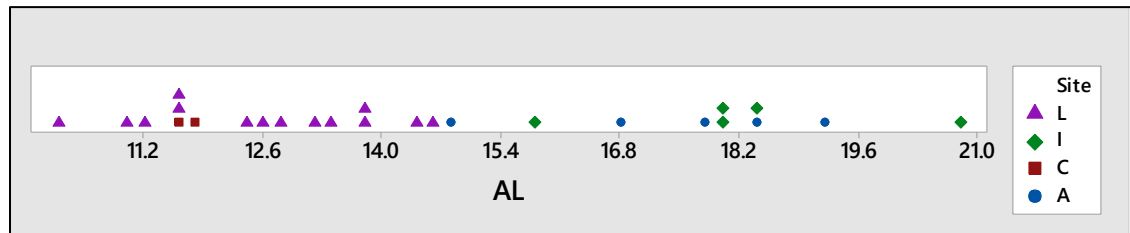


c Answers will vary.

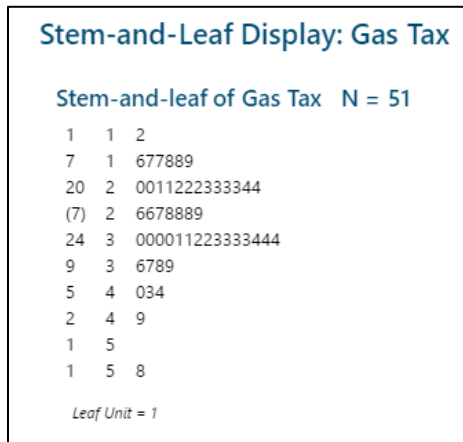
1.R.12 a-b Answers will vary. A typical histogram is shown next. Notice the gaps and the bimodal nature of the histogram, probably due to the fact that the samples were collected at different locations.



c The dotplot is shown as follows. The locations are indeed responsible for the unusual gaps and peaks in the relative frequency histogram.

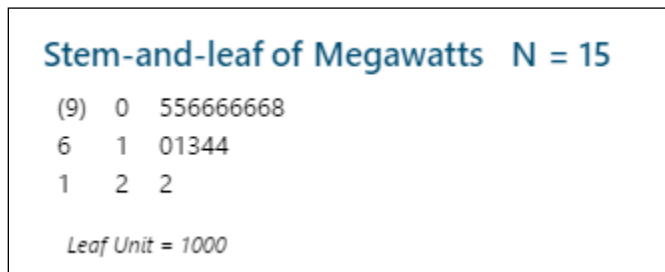


1.R.13 a-b The *Minitab* stem and leaf plot is shown next. The distribution is slightly skewed to the right.



c Pennsylvania (58.20) has an unusually high gas tax.

1.R.14 a-b Answers will vary. The *Minitab* stem and leaf plot is shown next. The distribution is skewed to the right.

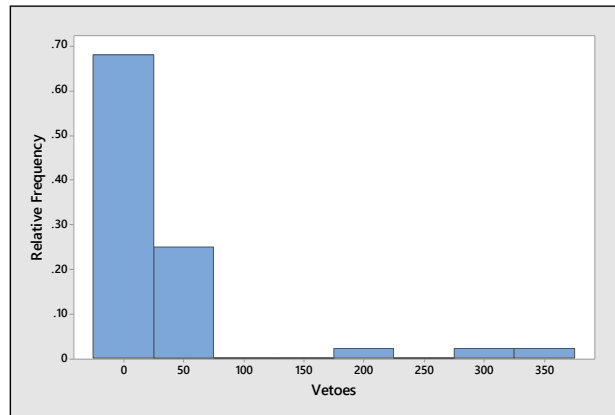


1.R.15 a-b The distribution is approximately mound-shaped, with one unusual measurement, in the class with midpoint at 100.8°. Perhaps the person whose temperature was 100.8 has some sort of illness coming on?

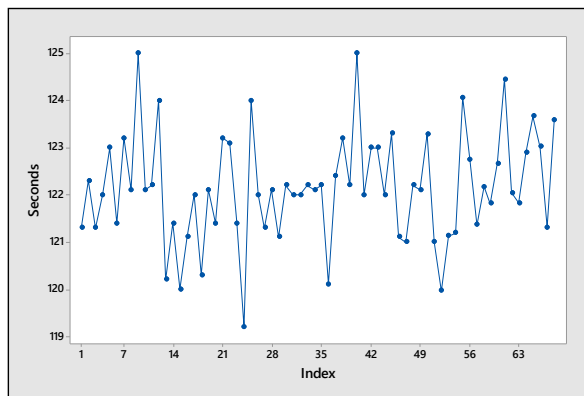
c The value 98.6° is slightly to the right of center.

On Your Own

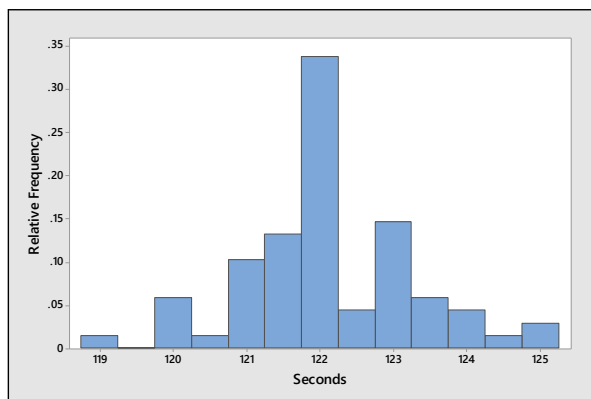
1.R.16 Answers will vary from student to student. The students should notice that the distribution is skewed to the right with a few presidents (Truman, Cleveland, and F.D. Roosevelt) casting an unusually large number of vetoes.



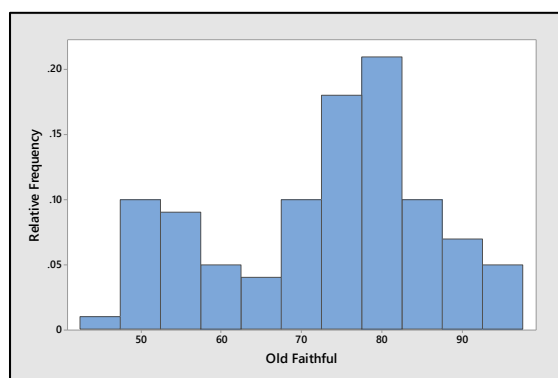
- 1.R.17 a** The line chart is shown next. The year in which a horse raced does not appear to have an effect on his winning time.



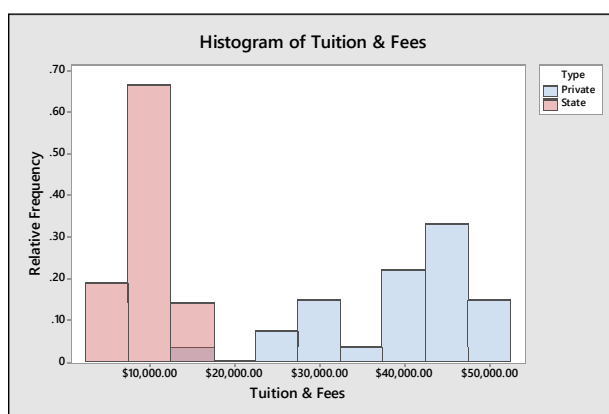
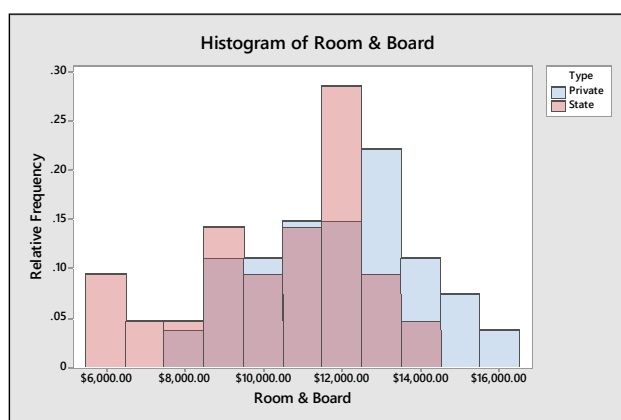
- b** Since the year of the race is not important in describing the data set, the distribution can be described using a relative frequency histogram. The distribution that follows is roughly mound-shaped with an unusually fast ($x = 119.2$) race times the year that *Secretariat* won the derby.



- 1.R.18** Answers will vary from student to student. Students should notice that both distributions are skewed left. The higher peak with a low bar to its left in the laptop group may indicate that students who would generally receive average scores (65-75) are scoring higher than usual. This may or may not be *caused* by the fact that they used laptop computers.
- 1.R.19** Answers will vary. A typical relative frequency histogram is shown next. There is an unusual bimodal feature.



1.R.20 Answers will vary. The student should notice that there is a clear difference in tuition and fees between private and state schools. Both distributions are roughly mound-shaped. The distribution of room and board costs are also roughly mound-shaped, but there is less of a difference between private and state schools.



CASE STUDY: How is Your Blood Pressure?

1. The following variables have been measured on the participants in this study: sex (qualitative); age in years (quantitative discrete); diastolic blood pressure (quantitative continuous, but measured to an integer value) and systolic blood pressure (quantitative continuous, but measured to an integer value). For each person, both systolic and diastolic readings are taken, making the data bivariate.
2. The important variables in this study are diastolic and systolic blood pressure, which can be described singly with histograms in various categories (male vs. female or by age categories). Further, the relationship between systolic and diastolic blood pressure can be displayed together using a scatterplot or a bivariate histogram.
3. Answers will vary from student to student, depending on the choice of class boundaries or the software package which is used. The histograms should look fairly mound-shaped.
4. Answers will vary from student to student.
5. In determining how a student's blood pressure compares to those in a comparable sex and age group, female students (ages 15-20) must compare to the population of females, while male students (ages 15-20) must compare to the population of males. The student should use his or her blood pressure and compare it to the scatterplot generated in part 4.

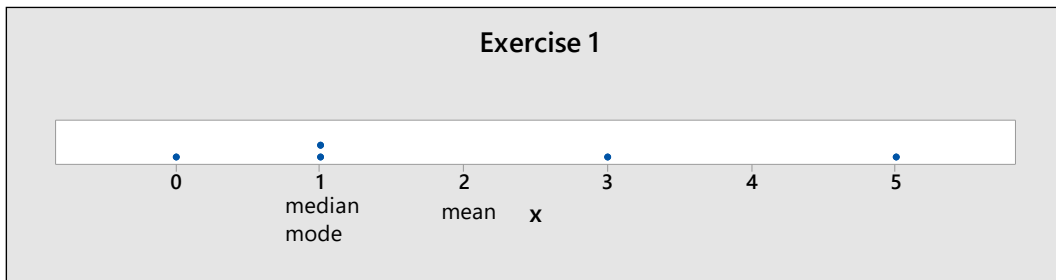
2: Describing Data with Numerical Measures

Section 2.1

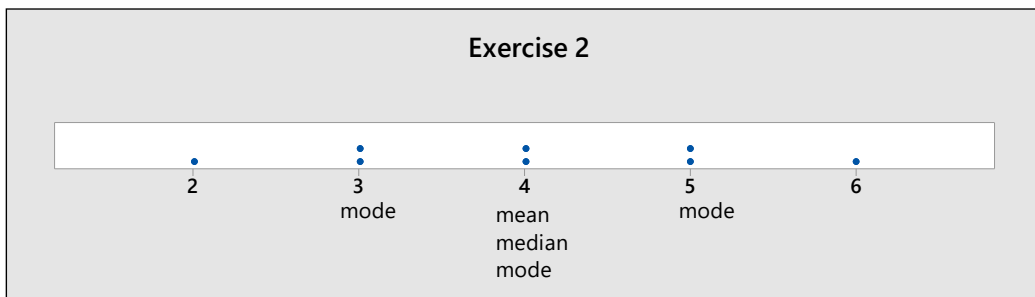
2.1.1 The mean is the sum of the measurements divided by the number of measurements, or

$$\bar{x} = \frac{\sum x_i}{n} = \frac{0+5+1+1+3}{5} = \frac{10}{5} = 2$$

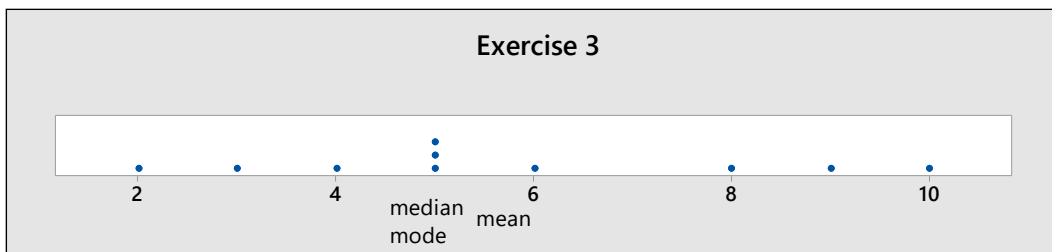
To calculate the median, the observations are first ranked from smallest to largest: 0, 1, 1, 3, 5. Then since $n=5$, the position of the median is $0.5(n+1) = 3$, and the median is the 3rd ranked measurement, or $m = 1$. The mode is the measurement occurring most frequently, or mode = 1. The three measures are located on the dotplot that follows.



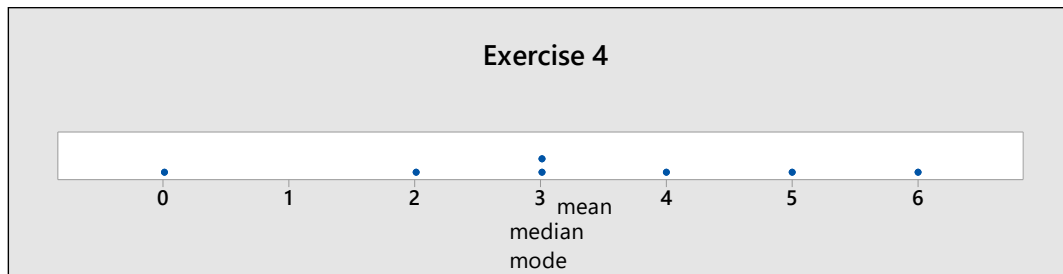
2.1.2 The mean is $\bar{x} = \frac{\sum x_i}{n} = \frac{3+2+\dots+5}{8} = \frac{32}{8} = 4$. To calculate the median, the observations are first ranked from smallest to largest: 2, 3, 3, 4, 4, 5, 5, 6. Since $n=8$ is even, the position of the median is $0.5(n+1) = 4.5$, and the median is the average of the 4th and 5th measurements, or $m = (4+4)/2 = 4$. Since three measurements occur twice, there are three modes: 3, 4, and 5. The dotplot is shown next.



2.1.3 The mean is $\bar{x} = \frac{\sum x_i}{n} = \frac{57}{10} = 5.7$. The ranked observations are: 2, 3, 4, 5, 5, 5, 6, 8, 9, 10. Since $n=10$, the median is halfway between the 5th and 6th ordered observations, or $m = (5+5)/2 = 5$. The measurement $x = 5$ occurs three times. Since this is the highest frequency of occurrence for the data set, the *mode* is 5. The dotplot is shown next.



- 2.1.4** The mean is $\bar{x} = \frac{\sum x_i}{n} = \frac{23}{7} = 3.286$. The ranked observations are: 0, 2, 3, 3, 4, 5, 6. Since $n = 7$, the median is the 4th ordered observation, or $m = 3$. The measurement $x = 3$ occurs twice. Since this is the highest frequency of occurrence for the data set, the *mode* is 3. The dotplot is shown next.



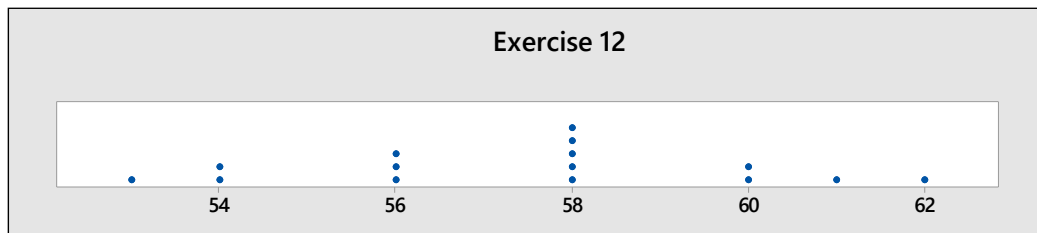
- 2.1.5** Since the mean is less than the median, the data set is skewed left.
- 2.1.6** Since the mean and the median are nearly equal, the data set is roughly symmetric.
- 2.1.7** Since the mean is greater than the median, the data set is skewed right.
- 2.1.8** Since the mean is less than the median, the data set is skewed left.
- 2.1.9** The mean is

$$\bar{x} = \frac{\sum x_i}{n} = \frac{94.3}{24} = 3.929$$

The data set is ordered from smallest to largest below. The position of the median is $.5(n + 1) = 12.5$ and the median is the average of the 12th and 13th measurements, or $(3.9 + 3.9)/2 = 3.9$.

3.2	3.5	3.7	3.9	4.2	4.4
3.3	3.5	3.9	4.0	4.3	4.4
3.4	3.6	3.9	4.0	4.3	4.5
3.5	3.6	3.9	4.2	4.3	4.8

- 2.1.10** Since the mean and the median are nearly equal, the data set is roughly symmetric.
- 2.1.11** The mean is $\bar{x} = \frac{\sum x_i}{n} = \frac{862}{15} = 57.467$. The ranked observations are: 53, 54, 54, 56, 56, 56, 58, 58, 58, 58, 58, 60, 60, 61, 62. Since $n = 15$, the median is the $.5(15 + 1) = 8^{\text{th}}$ ordered observation, or $m = 58$. The measurement $x = 58$ occurs five times. Since this is the highest frequency of occurrence for the data set, the *mode* is 58.
- 2.1.12** Since the mean and the median are nearly equal, the data set is roughly symmetric. The dotplot is shown next, confirming our result.



- 2.1.13** Similar to previous exercises. The mean is

$$\bar{x} = \frac{\sum x_i}{n} = \frac{0.99 + 1.92 + \cdots + 0.66}{14} = \frac{12.55}{14} = 0.896$$

To calculate the median, rank the observations from smallest to largest. The position of the median is $0.5(n + 1) = 7.5$, and the median is the average of the 7th and 8th ranked measurement or

$m = (0.67 + 0.69)/2 = 0.68$. The mode is .60, which occurs twice. Since the mean is slightly larger than the median, the distribution is slightly skewed to the right.

2.1.14 Similar to previous exercises. The mean is

$$\bar{x} = \frac{\sum x_i}{n} = \frac{125 + 45 + \cdots + 80}{15} = \frac{1305}{15} = 87$$

To calculate the median, rank the observations from smallest to largest. The position of the median is $0.5(n+1) = 8$, and the median is the 8th ranked measurement or $m = 105$. There are two modes—110 and 115—which both occur twice. Since the mean is smaller than the median, the distribution is skewed to the left.

2.1.15 Similar to previous exercises. The mean is

$$\bar{x} = \frac{\sum x_i}{n} = \frac{2150}{10} = 215$$

The ranked observations are: 175, 185, 190, 190, 200, 225, 230, 240, 250, 265. The position of the median is $0.5(n+1) = 5.5$ and the median is the average of the 5th and 6th observation or

$$\frac{200 + 225}{2} = 212.5$$

and the mode is 190 which occurs twice. The mean is slightly larger than the median, and the data may be slightly skewed right.

2.1.16 **a** $\bar{x} = \frac{\sum x_i}{n} = \frac{13904}{4} = 3476$ **b** $\bar{x} = \frac{\sum x_i}{n} = \frac{12674}{4} = 3168.50$

c The average premium cost in several different cities is not as important to the consumer as the cost for your particular characteristics, including the city in which you live.

2.1.17 **a** Although there may be a few households who own more than one DVR, the majority should own either 0 or 1. The distribution should be slightly skewed to the right.

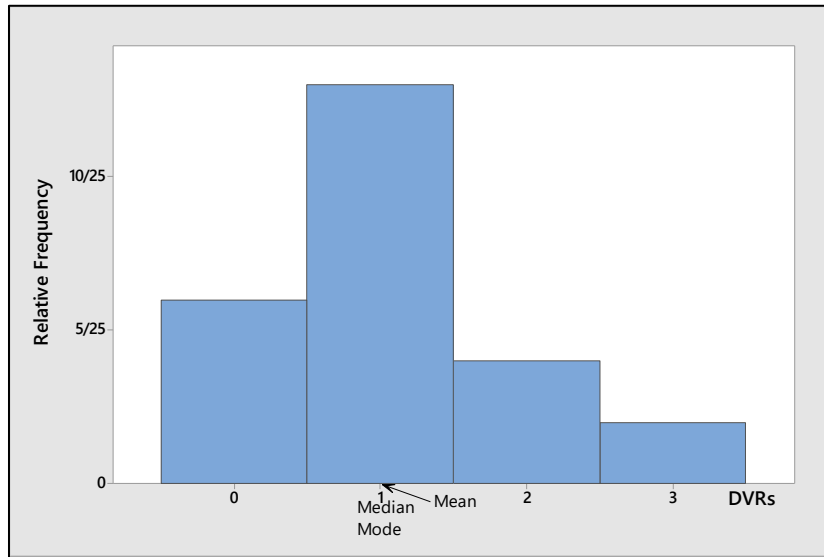
b Since most households will have only one DVR, we guess that the mode is 1.

c The mean is

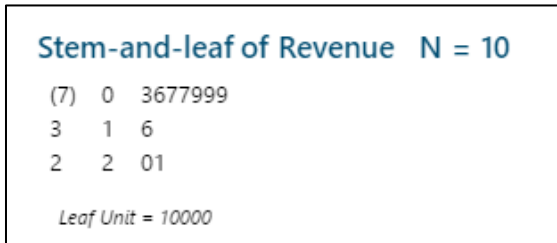
$$\bar{x} = \frac{\sum x_i}{n} = \frac{1+0+\cdots+1}{25} = \frac{27}{25} = 1.08$$

To calculate the median, the observations are first ranked from smallest to largest: There are six 0s, thirteen 1s, four 2s, and two 3s. Then since $n = 25$, the position of the median is $0.5(n+1) = 13$, which is the 13th ranked measurement, or $m = 1$. The mode is the measurement occurring most frequently, or mode = 1.

d The relative frequency histogram is shown next, with the three measures superimposed. Notice that the mean falls slightly to the right of the median and mode, indicating that the measurements are slightly skewed to the right.



2.1.18 a The stem and leaf plot that follows was generated by *Minitab*. It is skewed to the right.



b The mean is

$$\bar{x} = \frac{\sum x_i}{n} = \frac{166380 + 79919 + \cdots + 205004}{10} = \frac{1129104}{10} = 112,910.40$$

To calculate the median, rank the observations from smallest to largest.

37,949	69,495	71,890	79,919	93,662
94,571	94,595	166,380	205,004	215,639

Then since $n = 10$, the position of the median is $0.5(n+1) = 5.5$, the average of the 5th and 6th ranked measurements or $m = (93662 + 94571)/2 = 94,116.5$.

c Since the mean is strongly affected by outliers, the median would be a better measure of center for this data set.

2.1.19 It is obvious that any one family cannot have 2.5 children, since the number of children per family is a quantitative discrete variable. The researcher is referring to the average number of children per family calculated for all families in the United States during the 1930s. The average does not necessarily have to be integer-valued.

2.1.20 The distribution of sports salaries will be skewed to the right, because of the very high salaries of some sports figures. Hence, the median salary would be a better measure of center than the mean.

2.1.21 a The mean is

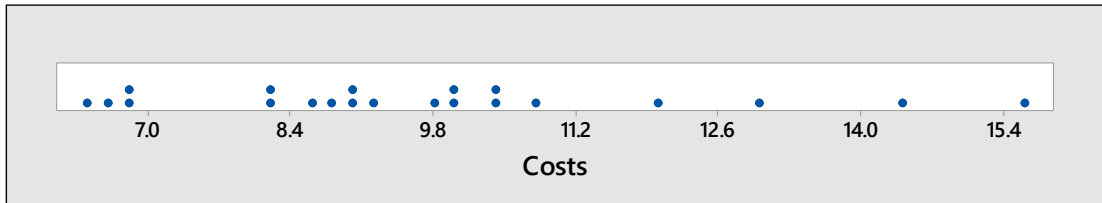
$$\bar{x} = \frac{\sum x_i}{n} = \frac{10.4 + 10.7 + \cdots + 8.8}{21} = \frac{203.5}{21} = 9.690$$

To calculate the median, rank the observations from smallest to largest.

6.4 6.5 6.8 6.8 8.2 8.2
 8.6 8.8 8.9 9.0 9.2
 9.7 10.0 10.0 10.4 10.4
 10.7 12.0 13.0 14.3 15.6

Then since $n = 21$, the position of the median is $0.5(n+1) = 11$, the 11th ranked measurement or $m = 9.2$. There are four modes—6.8, 8.2, 10.0 and 10.4—all of which occur twice.

b-c Since the mean is larger than the median, the data set is skewed right. The dotplot is shown next.



2.1.22 a Similar to previous exercises.

$$\bar{x} = \frac{\sum x_i}{n} = \frac{743.87}{16} = 46.492$$

b The ranked observations are:

39.99 39.99 39.99 39.99
 39.99 41.46 42.46 43.99
 46.49 47.19 47.49 47.90
 51.97 54.99 59.99 59.99

The position of the median is $0.5(n+1) = 8.5$, and the median is the average of the 8th and 9th observations or $m = (43.99 + 46.49)/2 = 45.24$. The mode is 39.99, which occurs 5 times.

c Answers will vary. Ease of purchase, speed of delivery, delivery costs, and return policy might be important.

Section 2.2

2.2.1 Calculate $\bar{x} = \frac{\sum x_i}{n} = \frac{12}{5} = 2.4$. Then create a table of differences, $(x_i - \bar{x})$ and their squares, $(x_i - \bar{x})^2$.

x_i	$x_i - \bar{x}$	$(x_i - \bar{x})^2$
2	-0.4	0.16
1	-1.4	1.96
1	-1.4	1.96
3	0.6	0.36
5	2.6	6.76
Total	0	11.20

Then, using the definition formula,

$$s^2 = \frac{\sum (x_i - \bar{x})^2}{n-1} = \frac{(2-2.4)^2 + \dots + (5-2.4)^2}{4} = \frac{11.20}{4} = 2.8$$

To use the computing formula, calculate $\sum x_i^2 = 2^2 + 1^2 + \dots + 5^2 = 40$. Then

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{40 - \frac{(12)^2}{5}}{4} = \frac{11.2}{4} = 2.8$$

The sample standard deviation is the positive square root of the variance or

$$s = \sqrt{s^2} = \sqrt{2.8} = 1.673$$

2.2.2 Calculate $\bar{x} = \frac{\sum x_i}{n} = \frac{17}{8} = 2.125$. Then create a table of differences, $(x_i - \bar{x})$ and their squares, $(x_i - \bar{x})^2$.

x_i	$x_i - \bar{x}$	$(x_i - \bar{x})^2$
4	1.875	3.515625
1	-1.125	1.265625
3	.875	.765625
1	-1.125	1.265625
3	.875	.765625
1	-1.125	1.265625
2	-.125	.015625
2	-.125	.015625
Total	0	8.875

Then, using the definition formula,

$$s^2 = \frac{\sum (x_i - \bar{x})^2}{n-1} = \frac{(4-2.125)^2 + \cdots + (2-2.125)^2}{7} = \frac{8.875}{7} = 1.2679$$

To use the computing formula, calculate $\sum x_i^2 = 4^2 + 1^2 + \cdots + 2^2 = 45$. Then

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{45 - \frac{(17)^2}{8}}{7} = \frac{8.875}{7} = 1.2679 \text{ and } s = \sqrt{s^2} = \sqrt{1.2679} = 1.126.$$

2.2.3 Calculate $\bar{x} = \frac{\sum x_i}{n} = \frac{31}{8} = 3.875$. Then create a table of differences, $(x_i - \bar{x})$ and their squares, $(x_i - \bar{x})^2$.

x_i	$x_i - \bar{x}$	$(x_i - \bar{x})^2$	x_i	$x_i - \bar{x}$	$(x_i - \bar{x})^2$
3	-.875	.765625	4	.125	.015625
1	-2.875	8.265625	4	.125	.015625
5	1.125	1.265625	3	-.875	.765625
6	2.125	4.515625	5	1.125	1.265625

Then, using the definition formula,

$$s^2 = \frac{\sum (x_i - \bar{x})^2}{n-1} = \frac{(3-3.875)^2 + \cdots + (5-3.875)^2}{7} = \frac{16.875}{7} = 2.4107$$

To use the computing formula, calculate $\sum x_i^2 = 3^2 + 1^2 + \cdots + 5^2 = 137$. Then

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{137 - \frac{(31)^2}{8}}{7} = \frac{16.875}{7} = 2.4107 \text{ and } s = \sqrt{s^2} = \sqrt{2.4107} = 1.553.$$

2.2.4 Use the computing formula, with $\sum x_i^2 = 374.65$ and $\sum x_i = 94.3$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{94.3}{24} = 3.929$,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{374.65 - \frac{(94.3)^2}{24}}{23} = \frac{4.129533}{23} = .179547$$

$$\text{and } s = \sqrt{s^2} = \sqrt{.179547} = .4237.$$

The range is calculated as $R = 4.8 - 3.2 = 1.6$ and $R/s = 1.6/.4237 = 3.78$, so that the range is approximately 4 standard deviations.

2.2.5 Use the computing formula, with $\sum x_i^2 = 49,634$ and $\sum x_i = 862$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{862}{15} = 57.467$,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{49634 - \frac{(862)^2}{15}}{14} = \frac{97.73333}{14} = 6.98095$$

$$\text{and } s = \sqrt{s^2} = \sqrt{6.98095} = 2.642.$$

The range is calculated as $R = 62 - 53 = 9$ and $R/s = 9/2.642 = 3.41$, so that the range is approximately 3.5 standard deviations.

2.2.6 Use the computing formula, with $\sum x_i^2 = 130$ and $\sum x_i = 32$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{32}{10} = 3.2$,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{130 - \frac{(32)^2}{10}}{9} = \frac{27.60}{9} = 3.066667 \text{ and } s = \sqrt{s^2} = \sqrt{3.066667} = 1.75119.$$

The range, $R = 6 - 1 = 5$, is $5/1.75119 = 2.855$ standard deviations.

2.2.7 The range is $R = 1.92 - .53 = 1.39$. Use the computing formula, with $\sum x_i^2 = 13.3253$ and $\sum x_i = 12.55$. Then,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{13.3253 - \frac{(12.55)^2}{14}}{13} = \frac{2.07512}{13} = .15962 \text{ and } s = \sqrt{s^2} = \sqrt{.15962} = .3995.$$

2.2.8 The range is $R = 125 - 25 = 100$. Use the computing formula, with $\sum x_i^2 = 131,125$ and $\sum x_i = 1305$. Then,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{131,125 - \frac{(1305)^2}{15}}{14} = \frac{17,590}{14} = 1256.42857 \text{ and } s = \sqrt{s^2} = \sqrt{1256.42857} = 35.446.$$

2.2.9 The range is $R = 265 - 175 = 90$. Use the computing formula, with $\sum x_i^2 = 470,900$ and $\sum x_i = 2150$. Then,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{470,900 - \frac{(2150)^2}{10}}{9} = \frac{8650}{9} = 961.11111 \text{ and } s = \sqrt{s^2} = \sqrt{961.11111} = 31.002.$$

2.2.10 a The range is $R = 2.39 - 1.28 = 1.11$.

b Calculate $\sum x_i^2 = 1.28^2 + 2.39^2 + \cdots + 1.51^2 = 15.415$. Then

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{15.415 - \frac{(8.56)^2}{5}}{4} = \frac{.76028}{4} = .19007$$

and $s = \sqrt{s^2} = \sqrt{.19007} = .436$

c The range, $R = 1.11$, is $1.11/.436 = 2.5$ standard deviations.

2.2.11 a The range is $R = 459.21 - 233.97 = 225.24$. **b** $\bar{x} = \frac{\sum x_i}{n} = \frac{3777.30}{12} = 314.775$

c Calculate $\sum x_i^2 = 243.92^2 + 233.97^2 + \dots + 286.41^2 = 1,267,488.3340$. Then

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{1267488.334 - \frac{(3777.30)^2}{12}}{11} = 7135.338773$$

and $s = \sqrt{s^2} = \sqrt{7135.338773} = 84.471$.

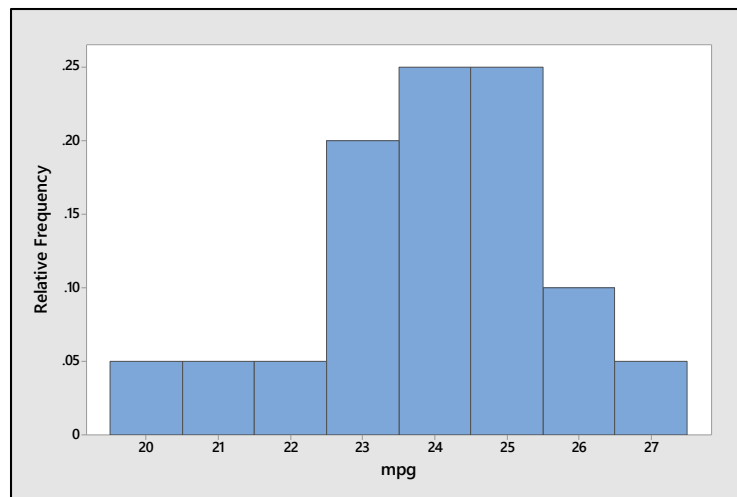
2.2.12 a Calculate $n = 10$, $\sum x_i = 68.5$, $\sum x_i^2 = 478.375$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{68.5}{10} = 6.85$ and

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{478.375 - \frac{(68.5)^2}{10}}{9}} = 1.008$$

b The most frequently recorded measurement is the mode or $x = 7$ hours of sleep.

2.2.13 a Max = 27, Min = 20.2 and the range is $R = 27 - 20.2 = 6.8$.

b Answers will vary. A typical histogram is shown next. The distribution is slightly skewed to the left.



c Calculate $n = 20$, $\sum x_i = 479.2$, $\sum x_i^2 = 11532.82$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{479.2}{20} = 23.96$

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{11532.82 - \frac{(479.2)^2}{20}}{19}} = \sqrt{2.694} = 1.641$$

2.2.14 a The range is $R = 71 - 40 = 31$.

b Calculate $n = 10$, $\sum x_i = 592$, $\sum x_i^2 = 36014$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{592}{10} = 59.2$

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{36014 - \frac{(592)^2}{10}}{9}} = \sqrt{107.5111} = 10.369$$

c Calculate $R/s = 31/10.369 = 3.00$, so that the range is approximately 3 standard deviations.

2.2.15 They are probably referring to the average number of times that men and women go camping per year.

Section 2.3

2.3.1 The range of the data is $R = 6 - 1 = 5$ and the range approximation with $n = 10$ is $s \approx \frac{R}{3} = 1.67$. The standard deviation of the sample is

$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{130 - \frac{(32)^2}{10}}{9}} = \sqrt{3.0667} = 1.751$$

which is very close to the estimate.

2.3.2 The range of the data is $R = 30 - 22 = 8$ and the range approximation with $n = 28$ is $s \approx \frac{R}{4} = 2$. The standard deviation of the sample is

$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{19285 - \frac{(733)^2}{28}}{27}} = \sqrt{3.55952} = 1.887$$

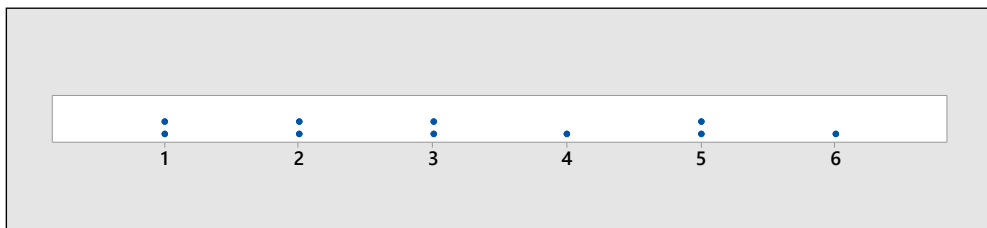
which is very close to the estimate.

2.3.3 The range of the data is $R = 7.1 - 4.0 = 3.1$ and the range approximation with $n = 15$ is $s \approx \frac{R}{3} = 1.033$. The standard deviation of the sample is

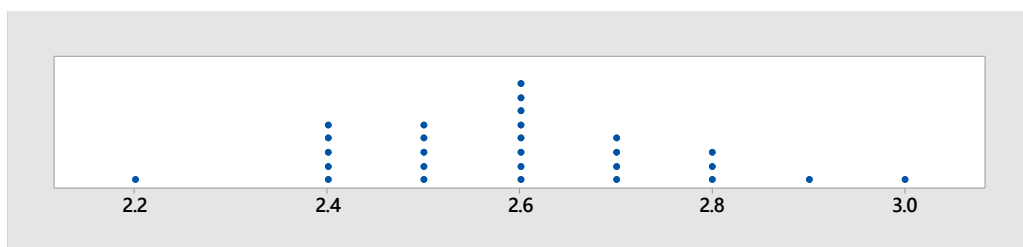
$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{531.17 - \frac{(88.1)^2}{15}}{14}} = \sqrt{.98067} = .990$$

which is very close to the estimate.

2.3.4 From the dotplot on the next page, you can see that the data set is not mound-shaped. Hence you can use Tchebysheff's Theorem, but not the Empirical Rule to describe the data.

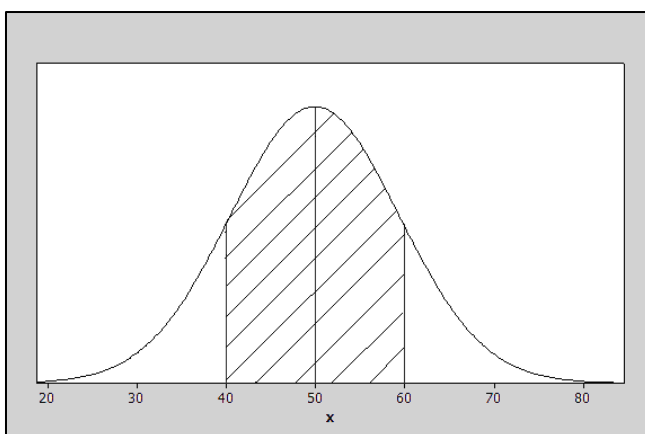


2.3.5 From the dotplot below, you can see that the data set is relatively mound-shaped. Hence you can use both Tchebysheff's Theorem and the Empirical Rule to describe the data.



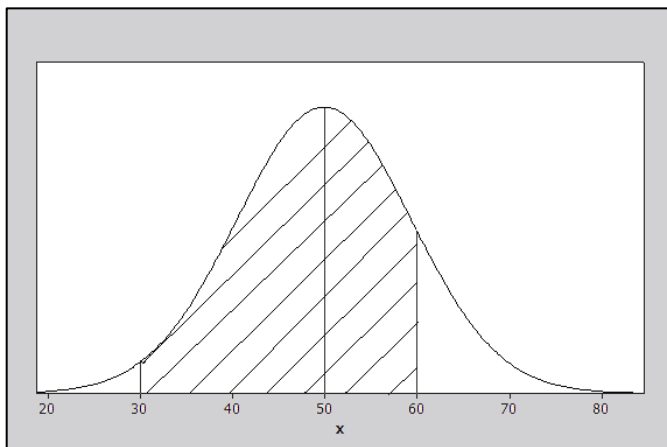
2.3.6-11 Since the distribution is relatively mound-shaped, use the Empirical Rule.

6. The interval from 40 to 60 represents $\mu \pm \sigma = 50 \pm 10$. Since the distribution is relatively mound-shaped, the proportion of measurements between 40 and 60 is .68 (or 68%) according to the Empirical Rule and is shown next.



7. Refer to the figure above. Again, using the Empirical Rule, the interval $\mu \pm 2\sigma = 50 \pm 2(10)$ or between 30 and 70 contains approximately .95 (or 95%) of the measurements.

8. Refer to the figure that follows.



Since approximately 68% of the measurements are between 40 and 60, the symmetry of the distribution implies that 34% of the measurements are between 50 and 60. Similarly, since 95% of the measurements are between 30 and 70, approximately 47.5% are between 30 and 50. Thus, the proportion of measurements between 30 and 60 is $0.34 + 0.475 = 0.815$.

9. From the figure in Exercise 6, the proportion of the measurements between 50 and 60 is 0.34 and the proportion of the measurements which are greater than 50 is 0.50. Therefore, the proportion that is greater than 60 must be $0.5 - 0.34 = 0.16$.

10. Refer to Exercise 9. The proportion of the measurements that are greater than 60 is .16. Since the total area under the curve is 1 or 100%, the proportion less than 60 must be $1 - .16 = .84$.

11. From the figure in Exercise 6, the proportion of the measurements between 40 and 50 is 0.34 and the proportion of the measurements which are greater than 50 is 0.50. Therefore, the proportion that is 40 or more must be $0.5 + 0.34 = 0.84$.

2.3.12-14 Since nothing is known about the shape of the data distribution, you must use Tchebysheff's Theorem to describe the data.

12. The interval from 60 to 90 represents $\mu \pm 3\sigma$ which will contain at least 8/9 of the measurements.

13. The interval from 65 to 85 represents $\mu \pm 2\sigma$ which will contain at least 3/4 of the measurements.

14. The interval from 62.5 to 87.5 represents $\mu \pm 2.5\sigma$ which will contain at least $1 - \frac{1}{k^2} = 1 - \frac{1}{2.5^2} = .84$ or 84% of the measurements.

2.3.15 The range of the data is $R = 1.1 - 0.5 = 0.6$ and the approximate value of s is

$$s \approx \frac{R}{3} = 0.2$$

Calculate $\sum x_i = 7.6$ and $\sum x_i^2 = 6.02$. The sample variance is

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{6.02 - \frac{(7.6)^2}{10}}{9} = \frac{0.244}{9} = .027111$$

and the standard deviation of the sample is $s = \sqrt{.027111} = 0.165$ which is very close to the range approximation.

2.3.16 The range of the data is $R = 1.92 - .53 = 1.39$ and the approximate value of s is

$$s \approx \frac{R}{3} = .4633$$

Calculate $\sum x_i = 12.55$ and $\sum x_i^2 = 13.3253$. The sample variance is

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{13.3253 - \frac{(12.55)^2}{14}}{13} = .159625$$

and the standard deviation of the sample is $s = \sqrt{.159625} = 0.3995$ which is very close to the range approximation.

2.3.17 The range is $R = 125 - 25 = 100$ and the approximate value of s is

$$s \approx \frac{R}{3} = 33.33$$

Use the computing formula, with $\sum x_i^2 = 131,125$ and $\sum x_i = 1305$. Then,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{131,125 - \frac{(1305)^2}{15}}{14} = \frac{17,590}{14} = 1256.42857 \text{ and } s = \sqrt{s^2} = \sqrt{1256.42857} = 35.446.$$

which is very close to the range approximation.

2.3.18 The range of the data is $R = 2.39 - 1.28 = 1.11$ and the approximate value of s is

$$s \approx \frac{R}{2.5} = 0.444$$

Calculate $\sum x_i = 8.56$ and $\sum x_i^2 = 15.4150$. The sample variance is

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{15.4150 - \frac{(8.56)^2}{5}}{4} = .19007$$

and the standard deviation of the sample is $s = \sqrt{.19007} = 0.436$ which is very close to the range approximation.

- 2.3.19 a** The stem and leaf plot generated by *Minitab* shows that the data is roughly mound-shaped. Note however the gap in the center of the distribution and the two measurements in the upper tail.

Stem-and-leaf of Weight N = 27		
1	7	5
2	8	3
6	8	7999
8	9	23
13	9	66789
13	10	
(3)	10	688
11	11	2244
7	11	788
4	12	4
3	12	8
2	13	
2	13	8
1	14	1

Leaf Unit = 0.01

- b** Calculate $\sum x_i = 28.41$ and $\sum x_i^2 = 30.6071$, the sample mean is $\bar{x} = \frac{\sum x_i}{n} = \frac{28.41}{27} = 1.052$ and the standard deviation of the sample is

$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{30.6071 - \frac{(28.41)^2}{27}}{26}} = 0.166$$

- c** The following table gives the actual percentage of measurements falling in the intervals $\bar{x} \pm ks$ for $k = 1, 2, 3$.

k	$\bar{x} \pm ks$	Interval	Number in Interval	Percentage
1	1.052 ± 0.166	0.866 to 1.218	21	78%
2	1.052 ± 0.332	0.720 to 1.384	26	96%
3	1.052 ± 0.498	0.554 to 1.550	27	100%

- d** The percentages in part **c** do not agree too closely with those given by the Empirical Rule, especially in the one standard deviation range. This is caused by the lack of mounding (indicated by the gap) in the center of the distribution.

- e** The lack of any one-pound packages is probably a marketing technique intentionally used by the supermarket. People who buy slightly less than one-pound would be drawn by the slightly lower price, while those who need exactly one-pound of meat for their recipe might tend to opt for the larger package, increasing the store's profit.

- 2.3.20** According to the Empirical Rule, if a distribution of measurements is approximately mound-shaped,

- a** approximately 68% or 0.68 of the measurements fall in the interval $\mu \pm \sigma = 12 \pm 2.3$ or 9.7 to 14.3
- b** approximately 95% or 0.95 of the measurements fall in the interval $\mu \pm 2\sigma = 12 \pm 4.6$ or 7.4 to 16.6
- c** approximately 99.7% or 0.997 of the measurements fall in the interval $\mu \pm 3\sigma = 12 \pm 6.9$ or 5.1 to 18.9

Therefore, approximately 0.3% or 0.003 will fall outside this interval.

- 2.3.21 a** The stem and leaf plots are shown next. The second set has a slightly higher location and spread.

Stem-and-leaf of Titanium content			Method = 1	N = 10
1	10	0		
3	11	00		
4	12	0		
(4)	13	0000		
2	14	0		
1	15	0		
Leaf Unit = 0.0001				

Stem-and-leaf of Titanium content			Method = 2	N = 10
1	11	0		
3	12	00		
5	13	00		
5	14	0		
4	15	00		
2	16	0		
1	17	0		
Leaf Unit = 0.0001				

- b Method 1:** Calculate $\sum x_i = 0.125$ and $\sum x_i^2 = 0.001583$. Then $\bar{x} = \frac{\sum x_i}{n} = 0.0125$ and

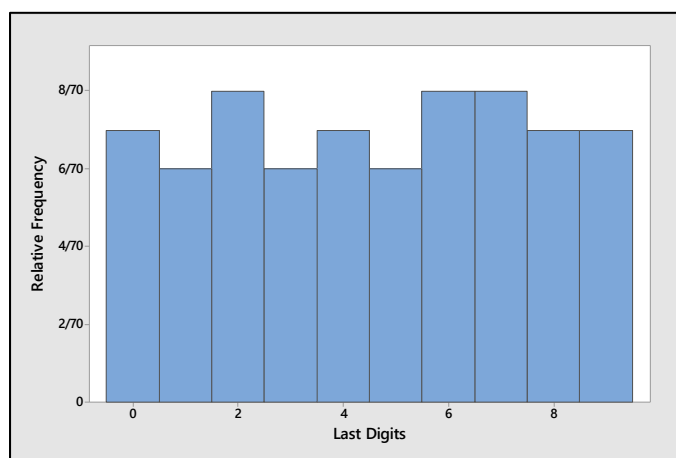
$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{0.001583 - \frac{(0.125)^2}{10}}{9}} = 0.00151$$

- Method 2:** Calculate $\sum x_i = 0.138$ and $\sum x_i^2 = 0.001938$. Then $\bar{x} = \frac{\sum x_i}{n} = 0.0138$ and

$$s = \sqrt{s^2} = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{0.001938 - \frac{(0.138)^2}{10}}{9}} = 0.00193$$

The results confirm the conclusions of part **a**.

- 2.3.22 a** The relative frequency histogram is shown next.



The distribution is relatively “flat” or “uniformly distributed” between 0 and 9. Hence, the center of the distribution should be approximately halfway between 0 and 9 or $(0+9)/2 = 4.5$.

b The range of the data is $R = 9 - 0 = 9$. Using the range approximation, $s \approx R/4 = 9/4 = 2.25$.

c Using the data entry method the students should find $\bar{x} = 4.586$ and $s = 2.892$, which are fairly close to our approximations.

2.3.23 a Similar to previous exercises. The intervals, counts and percentages are shown in the table.

k	$\bar{x} \pm ks$	Interval	Number in Interval	Percentage
1	4.586 ± 2.892	1.694 to 7.478	43	61%
2	4.586 ± 5.784	-1.198 to 10.370	70	100%
3	4.586 ± 8.676	-4.090 to 13.262	70	100%

b The percentages in part **a** do not agree with those given by the Empirical Rule. This is because the shape of the distribution is not mound-shaped, but flat.

2.3.24 a Although most of the animals will die at around 32 days, there may be a few animals that survive a very long time, even with the infection. The distribution will probably be skewed right.

b Using Tchebysheff’s Theorem, at least $3/4$ of the measurements should be in the interval $\mu \pm \sigma \Rightarrow 32 \pm 72$ or 0 to 104 days.

2.3.25 a The value of x is $\mu - \sigma = 32 - 36 = -4$.

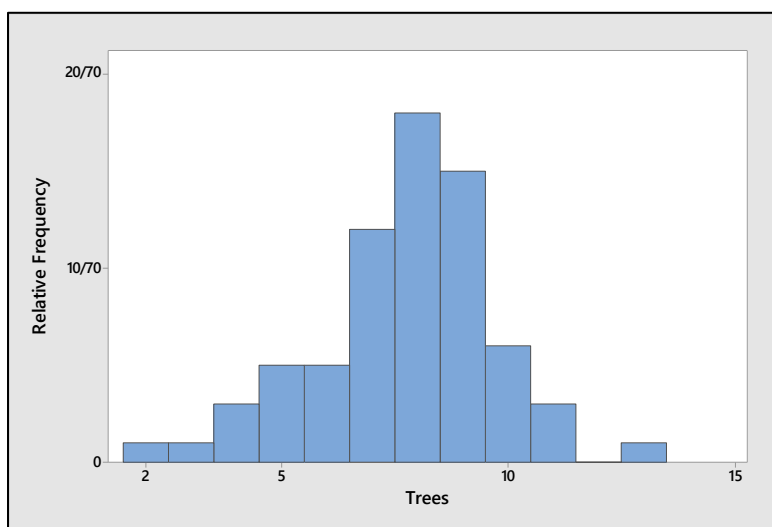
b The interval $\mu \pm \sigma$ is 32 ± 36 or -4 to 68. This interval should contain approximately 68% of the survival times, so that 16% will be longer than 68 days and 16% will be less than -4 days.

c The latter is clearly impossible. Therefore, the approximate values given by the Empirical Rule are not accurate, indicating that the distribution cannot be mound-shaped.

2.3.26 a We choose to use 12 classes of length 1.0. The tally and the relative frequency histogram follow.

Class i	Class Boundaries	Tally	f_i	Relative frequency, f_i/n
1	2 to < 3	1	1	1/70
2	3 to < 4	1	1	1/70
3	4 to < 5	111	3	3/70
4	5 to < 6	11111	5	5/70
5	6 to < 7	11111	5	5/70
6	7 to < 8	11111 11111 11	12	12/70
7	8 to < 9	11111 11111 11111 111	18	18/70
8	9 to < 10	11111 11111 11111	15	15/70
9	10 to < 11	11111 1	6	6/70

10	11 to < 12	111	3	3/70
11	12 to < 13		0	0
12	13 to < 14	1	1	1/70



b Calculate $n = 70$, $\sum x_i = 541$ and $\sum x_i^2 = 4453$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{541}{70} = 7.729$ is an estimate of μ .

c The sample standard deviation is

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{4453 - \frac{(541)^2}{70}}{69}} = \sqrt{3.9398} = 1.985$$

The three intervals, $\bar{x} \pm ks$ for $k = 1, 2, 3$ are calculated below. The table shows the actual percentage of measurements falling in a particular interval as well as the percentage predicted by Tchebysheff's Theorem and the Empirical Rule. Note that the Empirical Rule should be fairly accurate, as indicated by the mound-shape of the histogram in part **a**.

k	$\bar{x} \pm ks$	Interval	Fraction in Interval	Tchebysheff	Empirical Rule
1	7.729 ± 1.985	5.744 to 9.774	$50/70 = 0.71$	at least 0	≈ 0.68
2	7.729 ± 3.970	3.759 to 11.699	$67/70 = 0.96$	at least 0.75	≈ 0.95
3	7.729 ± 5.955	1.774 to 13.684	$70/70 = 1.00$	at least 0.89	≈ 0.997

2.3.27 a-b Calculate $R = 93 - 51 = 42$ so that $s \approx R/4 = 42/4 = 10.5$.

c Calculate $n = 30$, $\sum x_i = 2145$ and $\sum x_i^2 = 158,345$. Then

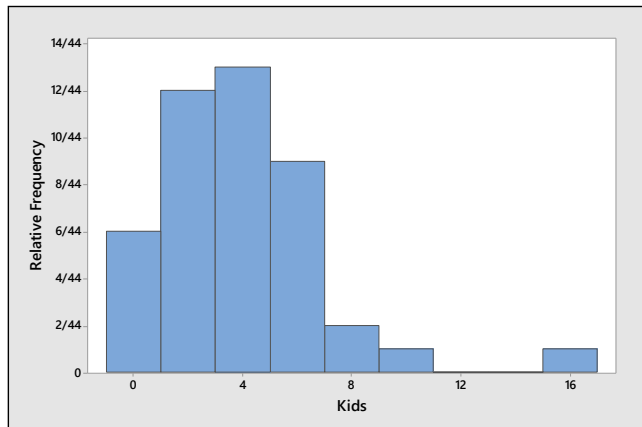
$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{158,345 - \frac{(2145)^2}{30}}{29} = 171.6379 \text{ and } s = \sqrt{171.6379} = 13.101$$

which is fairly close to the approximate value of s from part **b**.

d The two intervals are calculated below. The proportions agree with Tchebysheff's Theorem but are not too close to the percentages given by the Empirical Rule. (This is because the distribution is not quite mound-shaped.)

k	$\bar{x} \pm ks$	Interval	Fraction in Interval	Tchebysheff	Empirical Rule
2	71.5 ± 26.20	45.3 to 97.7	$30/30 = 1.00$	at least 0.75	≈ 0.95
3	71.5 ± 39.30	32.2 to 110.80	$30/30 = 1.00$	at least 0.89	≈ 0.997

2.3.28 a Answers will vary. A typical histogram is shown next. The distribution is skewed to the right.



b Calculate $n = 44$, $\sum x_i = 158$ and $\sum x_i^2 = 926$. Then

$$\bar{x} = \frac{\sum x_i}{n} = \frac{158}{44} = 3.591,$$

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{926 - \frac{(158)^2}{44}}{43} = 8.34038$$

and $s = \sqrt{8.34038} = 2.888$

c The three intervals, $\bar{x} \pm ks$ for $k = 1, 2, 3$ are calculated below. The table shows the actual percentage of measurements falling in a particular interval as well as the percentage predicted by Tchebysheff's Theorem and the Empirical Rule. Note that the Empirical Rule is not very accurate for the first interval, since the histogram in part **a** is skewed.

k	$\bar{x} \pm ks$	Interval	Fraction in Interval	Tchebysheff	Empirical Rule
1	3.591 ± 2.888	.703 to 6.419	$34/44 = .77$	at least 0	≈ 0.68
2	3.591 ± 5.776	-2.785 to 9.367	$42/44 = .95$	at least 0.75	≈ 0.95
3	3.591 ± 8.664	-5.073 to 12.255	$43/44 = .98$	at least 0.89	≈ 0.997

2.3.29 a Answers will vary. A typical stem and leaf plot is generated by *Minitab*.



b Calculate $n = 16$, $\sum x_i = 386$ and $\sum x_i^2 = 9476$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{386}{16} = 24.125$,

$$s^2 = \frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1} = \frac{9476 - \frac{(386)^2}{16}}{15} = 10.91667$$

and $s = \sqrt{s^2} = \sqrt{10.91667} = 3.304$.

c Calculate $\bar{x} \pm 2s \Rightarrow 24.125 \pm 6.608$ or 17.517 to 30.733. From the original data set, 16 of the 16 measurements, or 100% fall in this interval.

2.3.30 The following information is available:

$$n = 400, \bar{x} = 600, s^2 = 4900$$

The standard deviation of these scores is then 70, and the results of Tchebysheff's Theorem follow:

k	$\bar{x} \pm ks$	Interval	Tchebysheff
1	600 ± 70	530 to 670	at least 0
2	600 ± 140	460 to 740	at least 0.75
3	600 ± 210	390 to 810	at least 0.89

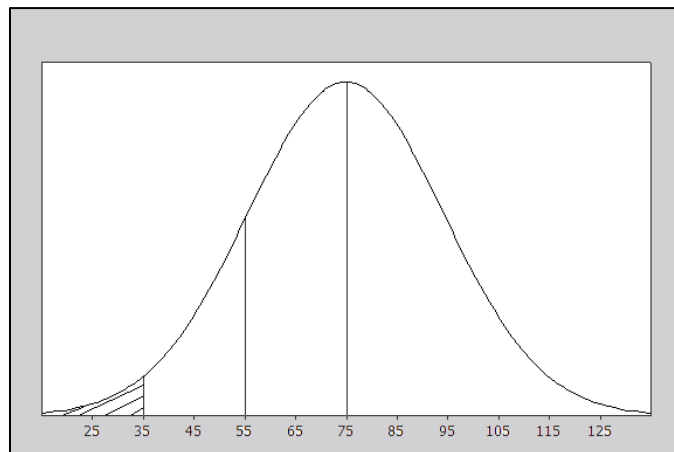
If the distribution of scores is mound-shaped, we use the Empirical Rule, and conclude that approximately 68% of the scores would lie in the interval 530 to 670 (which is $\bar{x} \pm s$). Approximately 95% of the scores would lie in the interval 460 to 740.

2.3.31 The results of the Empirical Rule follow:

k	$\bar{x} \pm ks$	Interval	Empirical Rule
1	420 ± 5	415 to 425	approximately 0.68
2	420 ± 10	410 to 430	approximately 0.95
3	420 ± 15	405 to 435	approximately 0.997

Notice that we are assuming that attendance follows a mound-shaped distribution and hence that the Empirical Rule is appropriate.

2.3.32 a It is known that duration times are approximately normal, with mean 75 and standard deviation 20. In order to determine the probability that a commercial lasts less than 35 seconds, we must determine the fraction of the curve which lies within the shaded area in the figure that follows. Using the Empirical Rule, the fraction of the area between 35 and 75 is half of 0.95 or 0.475. Hence, the fraction below 35 would be $0.5 - 0.475 = 0.025$.



b The fraction of the curve area that lies above the 55 second mark may again be determined by using the Empirical Rule. Refer to the figure in part **a**. The fraction between 55 and 75 is 0.34 and the fraction above 75 is 0.5. Hence, the probability that a commercial lasts longer than 55 seconds is $0.5 + 0.34 = 0.84$.

Section 2.4

2.4.1 The ordered data are: 0, 1, 3, 4, 4, 5, 6, 6, 6, 7, 7, 8

Calculate $n=12$, $\sum x_i = 57$ and $\sum x_i^2 = 337$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{57}{12} = 4.75$ and the sample standard deviation is

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{337 - \frac{(57)^2}{12}}{11}} = \sqrt{6.022727} = 2.454$$

For the smallest observation, $x=0$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{0-4.75}{2.454} = -1.94$ and for the largest observation,

$x=8$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{8-4.75}{2.454} = 1.32$. Since neither z -score exceeds 2 in absolute value, none of the observations are unusually small or large.

2.4.2 The ordered data are: 1.0, 1.7, 2.0, 2.1, 2.3, 2.8, 2.9, 4.4, 5.1, 6.5, 7.8

Calculate $n=11$, $\sum x_i = 38.6$ and $\sum x_i^2 = 182.3$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{38.6}{11} = 3.509$ and the sample standard deviation is

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{182.30 - \frac{(38.6)^2}{11}}{10}} = \sqrt{4.68491} = 2.164$$

For the smallest observation, $x=1.0$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{1.0-3.509}{2.164} = -1.16$ and for the largest observation,

$x=7.8$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{7.8-3.509}{2.164} = 1.98$. Since neither z -score exceeds 2 in absolute value, none of the observations are unusually small or large.

2.4.3 The ordered data are: 2, 3, 4, 5, 6, 6, 6, 7, 8, 9, 9, 10, 25

Calculate $n=13$, $\sum x_i = 100$ and $\sum x_i^2 = 1162$. Then $\bar{x} = \frac{\sum x_i}{n} = \frac{100}{13} = 7.692$ and the sample standard deviation is

$$s = \sqrt{\frac{\sum x_i^2 - \frac{(\sum x_i)^2}{n}}{n-1}} = \sqrt{\frac{1162 - \frac{(100)^2}{13}}{12}} = \sqrt{32.73077} = 5.721$$

For the smallest observation, $x=2$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{2-7.692}{5.721} = -.99$ and for the largest observation,

$x=25$, $z\text{-score} = \frac{x-\bar{x}}{s} = \frac{25-7.692}{5.721} = 3.03$. The value $x=25$ is an outlier.

2.4.4 The ordered data are: 2, 3, 4, 6, 7, 8, 10. With $n=7$, the median is in position $0.5(n+1)=4$, the lower quartile is in position $0.25(n+1)=2$ and the upper quartile is in position $0.75(n+1)=6$. Hence, $m=6$, $Q_1=3$ and $Q_3=8$.

2.4.5 The ordered data are: 0, 1, 2, 3, 5, 5, 6, 6, 7. With $n=9$, the median is in position $0.5(n+1)=5$, the lower quartile is in position $0.25(n+1)=2.5$ (halfway between the 2nd and 3rd ordered observations) and the upper