

Instructor's Resource Guide

TO ACCOMPANY

Understandable Statistics: Concepts and Methods

Eleventh Edition

Joseph Kupresanin
Cecil College

© 2015 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: 978-130510243-9
ISBN-10: 1-30510243-6

Cengage Learning
200 First Stamford Place, 4th Floor
Stamford, CT 06902
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,
visit **www.cengage.com**.

Purchase any of our products at your local college
store or at our preferred online store
www.cengagebrain.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, posted electronically or distributed, except that portions of the Supplement may be provided to your students IN PRINT FORM ONLY 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. 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.

Table of Contents

PART I: COURSE OVERVIEW AND TEACHING TIPS	2
PART II: FREQUENTLY USED FORMULAS AND TRANSPARENCY MASTERS	23
PART III: SAMPLE CHAPTER TESTS AND ANSWERS	62
PART IV: COMPLETE SOLUTIONS	207

Part I

Course Overview and Teaching Tips

Suggestions for Using the Text

In writing this text, we have followed the premise that a good textbook must be more than just a repository of knowledge. A good textbook should be an agent that interacts with the student to create a working knowledge of the subject. To help achieve this interaction, we have modified the traditional format in order to encourage active student participation.

Each chapter begins with Preview Questions, which indicate the topics addressed in each section of the chapter. Next is a Focus Problem that uses real-world data. The Focus Problems show the students the kinds of questions they can answer once they have mastered the material in the chapter. Consequently, students are asked to solve each chapter's Focus Problem as soon as the concepts required for the solution have been introduced.

Procedure displays, which summarize key strategies for carrying out statistical procedures and methods, and definition boxes are interspersed throughout each chapter. Another special feature of this text is the Guided Exercises built into the reading material. These Guided Exercises, which include complete worked solutions, help the students to focus on key concepts in the newly introduced material. Also, the Section Problems reinforce student understanding and sometimes require the student to look at the concepts from a slightly different perspective than the one presented in the section. Some Section Problems are categorized as Statistical Literacy, Critical Thinking, Basic Computation, or Expand Your Knowledge. The Statistical Literacy problems typically review definitions and statistical symbols used in the text. Critical Thinking problems often ask about statistical formulas and unusual situations that involve common ideas. Basic Computation problems are just that – problems designed to reinforce the raw mechanics of statistical formulas or procedures. Expand Your Knowledge problems appear at the end of each section and present enrichment topics designed to challenge the student with the most advanced concepts in that section.

The Chapter Review Problems are much more comprehensive. They require students to place each problem in the context of all they have learned in the chapter. Data Highlights, found at the end of each chapter, ask students to look at data as presented in newspapers, magazines, and other media and then to apply relevant methods of interpretation. Finally, Linking Concept problems ask students to verbalize their skills and synthesize the material.

We believe that the progression from small-step Guided Exercises to Section Problems, Chapter Review Problems, Data Highlights, and Linking Concepts will enable instructors to use their class time in a very profitable way, going from specific mastery details to more comprehensive decision-making analysis.

Calculators and statistical computer software remove much of the computational burden from statistics. Many basic scientific calculators provide the mean and standard deviation. Calculators that support two-variable statistics provide the coefficients of the least-squares line, the value of the correlation coefficient, and the predicted value of y for a given x . Graphing calculators sort the data, and many provide the least-squares line. Statistical software packages give full support for descriptive statistics and inferential statistics. Students benefit from using these technologies. In many examples and exercises in *Understandable Statistics* we ask students to use calculators to verify answers. For example, in keeping with the use of computer technology and standard practice in research, hypothesis testing is now introduced using P values. The critical region method is still supported but is not given primary emphasis. Illustrations in the text show TI-83Plus/TI-84Plus, MINITAB, SPSS, and Microsoft Excel outputs, so students can see the different types of information available to them through the use of technology.

However, it is not enough to enter data and punch a few buttons to get statistical results. The formulas that produce the statistics contain a great deal of information about the *meaning* of those statistics. The text breaks down formulas into tabular form so that students can see the information in the formula. We find it useful to take class time to discuss formulas. For instance, an essential part of the standard deviation formula is the comparison of each data value with the mean. When we point this out to students, it gives meaning to the standard deviation. When students understand the content of the formulas, the numbers they get from their calculators or computers begin to make sense.

The eleventh edition includes Cumulative Reviews at the end of chapters 3, 6, 9, and 11; these sections tie together the concepts from those chapters to help the student to put those concepts in a larger context. The Technology Notes briefly describe relevant procedures for using the TI-83Plus/TI-84Plus calculator, Microsoft Excel, MINITAB, and SPSS. In addition, the Using Technology sections have been updated to include SPSS material.

For a course in which technologies are strongly incorporated into the curriculum, we provide Technology Guides (for TI-83Plus/TI-84Plus, MINITAB, Microsoft Excel, and SPSS). These guides give specific hints for using the technologies and also provide Lab Activities to help students explore various statistical concepts.

Finally, accompanying the text are several interactive components that help to demonstrate key concepts. Through on-line tutorials and an interactive textbook, the student can manipulate data to see and understand the effects in context.

Alternate Paths through the Text

As with previous editions, the eleventh edition of *Understandable Statistics* is designed to be flexible. In most one-semester courses, it is not possible to cover all the topics. The text provides many topics, so you can tailor a course to fit student needs. The text also aims to be a *readable reference* for topics not specifically included in your course.

Table of Prerequisite Material

Chapter	Prerequisite Sections
1 Getting Started	None
2 Organizing Data	1.1, 1.2
3 Averages and Variation	1.1, 1.2, 2.2
4 Elementary Probability Theory	1.1, 1.2, 2.2, 3.1, 3.2
5 The Binomial Probability Distribution and Related Topics	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, with 4.3 useful but not essential
6 Normal Curves and Sampling Distributions	
With 6.6 omitted	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, 5.1
With 6.6 included	Add 5.2, 5.3
7 Estimation	
With 7.3 and parts of 7.4 omitted	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, 5.1, 6.1, 6.2, 6.3, 6.4, 6.5
With 7.3 and all of 7.4	Add 5.2, 5.3, 6.6
8 Hypothesis Testing	
With 8.3 and parts of 8.4 omitted	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, 5.1, 6.1, 6.2, 6.3, 6.4, 6.5
With 8.3 and all of 8.4	Add 5.2, 5.3, 6.6
9 Correlation and Regression	
9.1 and 9.2	1.1, 1.2, 3.1, 3.2
9.3 and 9.4	Add 4.1, 4.2, 5.1, 6.1, 6.2, 6.3, 6.4, 6.5, 7.1, 8.1, 8.2
10 Chi-Square and <i>F</i> Distributions	
With 10.3 omitted	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, 5.1, 6.1, 6.2, 6.3, 6.4, 6.5, 8.1
With 10.3 included	Add 7.1
11 Nonparametric Statistics	1.1, 1.2, 2.2, 3.1, 3.2, 4.1, 4.2, 5.1, 6.1, 6.2, 6.3, 6.4, 6.5, 8.1, 8.3

Teaching Tips for Each Chapter

CHAPTER 1: GETTING STARTED

Double-Blind Studies (Section 1.3)

The double-blind method of data collection, mentioned at the end of Section 1.3, is an important part of standard research practice. A typical use is in testing new medications. Because the researcher does not know which patients are receiving the experimental drug and which are receiving the established drug (or a placebo), the researcher is prevented from doing things subconsciously that might skew the results.

If, for instance, the researcher communicates a more optimistic attitude to patients in the experimental group, this could influence how they respond to diagnostic questions or actually might influence the course of their illness. And if the researcher wants the new drug to prove effective, this could subconsciously influence how he or she handles information related to each patient's case. All such factors are eliminated in double-blind testing.

The following appears in the physician's dosing instructions package insert for the prescription drug QUIXIN™:

In randomized, double-masked, multicenter controlled clinical trials where patients were dosed for 5 days, QUIXIN™ demonstrated clinical cures in 79% of patients treated for bacterial conjunctivitis on the final study visit day (days 6–10).

Note the phrase *double-masked*. Apparently, this is a synonym for *double-blind*. Since *double-blind* is used widely in the medical literature and in clinical trials, why do you suppose that the company chose to use *double-masked* instead?

Perhaps this will provide some insight: QUIXIN™ is a topical antibacterial solution for the treatment of conjunctivitis; i.e., it is an antibacterial eye drop solution used to treat an inflammation of the conjunctiva, the mucous membrane that lines the inner surface of the eyelid and the exposed surface of the eyeball. Perhaps, since QUIXIN™ is a treatment for eye problems, the manufacturer decided the word *blind* should not appear *anywhere* in the discussion.

Source: Package insert. QUIXIN™ is manufactured by Santen Oy, P.O. Box 33, FIN-33721 Tampere, Finland, and marketed by Santen, Inc., Napa, CA 94558, under license from Daiichi Pharmaceutical Co., Ltd., Tokyo, Japan.

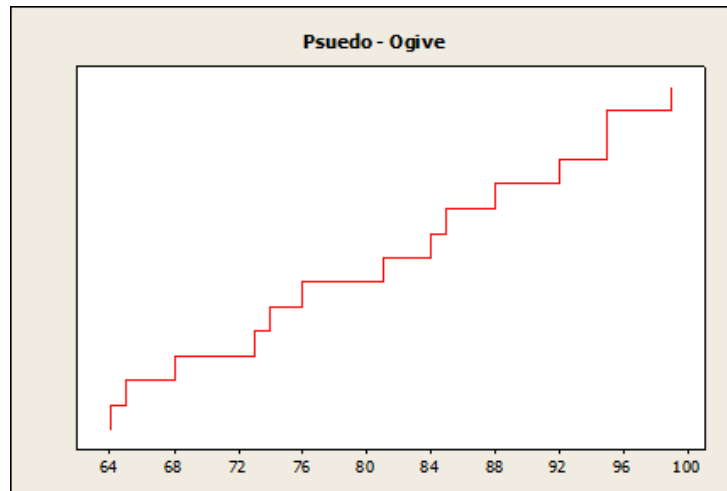
CHAPTER 2: ORGANIZING DATA

Emphasize when to use the various graphs discussed in this chapter: bar graphs when comparing data sets, circle graphs for displaying how data are dispersed into several categories, time-series graphs to display how data change over time, histograms or frequency polygons to display relative frequencies of grouped data, and stem-and-leaf displays for displaying grouped data in a way that does not lose the detail of the original raw data.

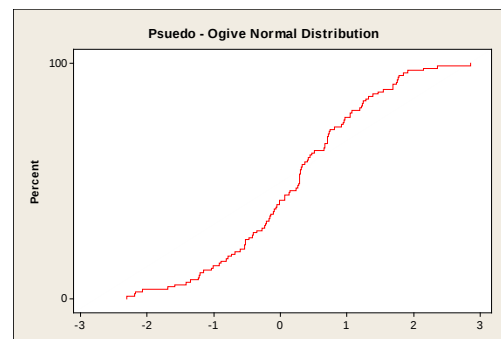
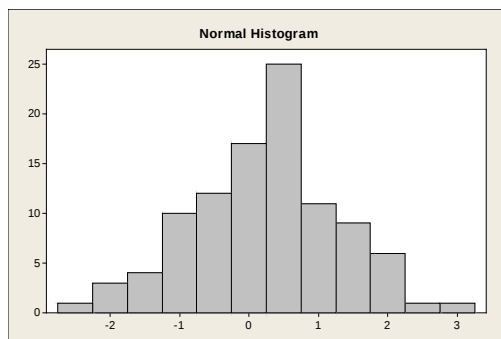
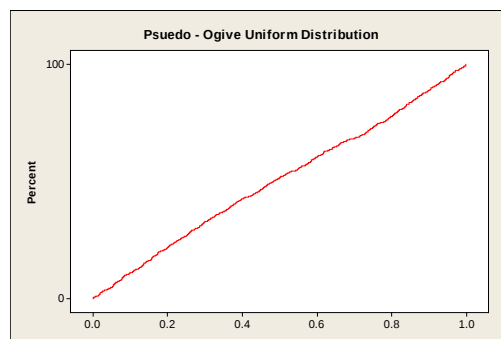
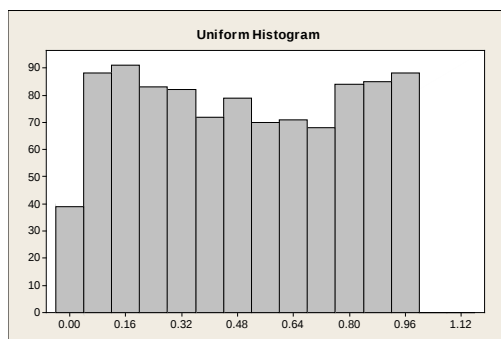
Drawing and Using Ogives (Section 2.1)

The text describes how an ogive, which is a graph displaying a cumulative-frequency distribution, can be constructed easily using a frequency table. However, a graph of the same basic sort can be constructed even more quickly than that. Simply arrange the data values in ascending order and then plot one point for each data value, where the x coordinate is the data value and the y coordinate starts at 1 for the first point and increases by 1 for each successive point. Finally, connect adjacent points with line segments. In the resulting graph, for any x , the corresponding y value will be (roughly) the number of data values less than or equal to x .

For example, here is the graph for the data set 64, 65, 68, 73, 74, 76, 81, 84, 85, 88, 92, 95, 95, and 99:



This graph is not technically an ogive because the possibility of duplicate data values (such as 95 in this example) means that the graph will not necessarily be a function. But the graph can be used to get a quick fix on the general shape of the cumulative distribution curve. And by implication, the graph can be used to get a quick idea of the shape of the frequency distribution, as illustrated below.



The pseudo-ogive obtained from the example data set suggests a uniform distribution on the interval 63–100 or thereabouts.

CHAPTER 3: AVERAGES AND VARIATION

Students should be instructed in the various ways that sets of numeric data can be represented by a single number. The concepts of this section illustrate for students the need for this kind of representation.

The different ways this can be done are discussed in Section 3.1. The mean, median, and mode vary in appropriateness depending on the situation. In many cases of numeric data, the mean is the most appropriate measure of central tendency. If the mean is larger or smaller than most of the data values, however, then the median may be the number that best represents the data set. The median is most appropriate usually if the data set is annual salaries, costs of houses, or any data set that contains one or a few very large or very small values. The mode would be most appropriate if the population were the votes in an election or Nielsen television ratings, for example. Students should get acquainted with these concepts by calculating the mean, median, and mode for different data sets and then interpreting the meaning of each one and determining which measure of central tendency is the most appropriate.

The range, variance, and standard deviation can be presented to students as other numbers that aid in the representation of a data set in that they measure how data are dispersed. Students will begin to have a better understanding of these measures of dispersion by calculating these numbers for given data sets and interpreting their respective meanings. These concepts of central tendency and dispersion also can be applied to grouped data, and students should become acquainted with interpreting these measures for given realistic situations in which data have been collected.

Chebyshev's theorem is important to discuss with students because it relates to the mean and standard deviation of *any* data set.

Finally, the mean, median, first and third quartiles, and range of a data set can be viewed easily in a box-and-whisker plot.

CHAPTER 4: ELEMENTARY PROBABILITY THEORY

What Is Probability? (Section 4.1)

As the text describes, there are several methods for assigning a probability to an event. Probability based on intuition is often called *subjective* probability. Thus understood, probability is a numerical measure of a person's estimate of the likelihood of some event. Subjective probability is assumed to be reflected in a person's decisions: The higher an event's probability, the more the person would be willing to bet on its occurring.

Probability based on relative frequency is often called *experimental* probability because the relative frequency is calculated from an observed history of experiment outcomes. But we are already using the word *experiment* in a way that is neutral among the different treatments of probability—namely, as the name for the activity that produces various possible outcomes. So when we are talking about probability based on relative frequency, we will call this *observed* probability.

Probability based on equally likely outcomes is often called *theoretical* probability because it is ultimately derived from a theoretical model of the experiment's structure. The experiment may be conducted only once, or not at all, and need not be repeatable.

These three ways of treating probability are compatible and complementary. For a reasonable, well-informed person, the subjective probability of an event should match the theoretical probability, and the theoretical probability, in turn, predicts the observed probability as the experiment is repeated many times.

Also, it should be noted that although in statistics probability is officially a property of *events*, it can be thought of as a property of *statements* as well. The probability of a statement equals the probability of the event that makes the statement true.

Probability and statistics are overlapping fields of study; if they weren't, there would be no need for a chapter on probability in a book on statistics. So the general statement in the text that probability deals with known populations, whereas statistics deals with unknown populations is necessarily a simplification. However, the statement does express an important truth: If we confront an experiment we initially know absolutely nothing about, then we can collect data but we cannot calculate probabilities. In other words, we can only calculate probabilities after we have

formed some idea of, or acquaintance with, the experiment. To find the theoretical probability of an event, we have to know how the experiment is set up. To find the observed probability, we have to have a record of previous outcomes. And as reasonable people, we need some combination of those same two kinds of information to set our subjective probability.

This may seem obvious, but it has important implications for how we understand technical concepts encountered later in the course. There will be times when we would like to make a statement, say, about the mean of a population and then give the probability that this statement is true—i.e., the probability that the event described by the statement occurs (or has occurred). What we discover when we look closely, however, is that often this cannot be done. Often we have to settle for some other conclusion instead. The Teaching Tips for Sections 7.1 and 8.1 describe two instances of this problem.

CHAPTER 5: THE BINOMIAL PROBABILITY DISTRIBUTION AND RELATED TOPICS

Binomial Probabilities (Section 5.2)

Students should be able to show that $pq = p(1 - p)$ has its maximum value at $p = 0.5$. There are at least three ways to demonstrate this: graphically, algebraically, and using calculus.

Graphical method

Recall that $0 \leq p \leq 1$. So, for $q = 1 - p$, $0 \leq q \leq 1$ and $0 \leq pq \leq 1$. Plot $y = pq = p(1 - p)$ using MINITAB, a graphing calculator, or other technology. The graph is a parabola. Observe which value of p maximizes pq . (Many graphing calculators can find the maximum value and where it occurs.)

So pq has a maximum value of 0.25 when $p = 0.5$.

Algebraic method

From the definition of q , it follows that $pq = p(1 - p) = p - p^2 = -p^2 + p + 0$. Recognize that this is a quadratic function of the form $ax^2 + bx + c$, where p is used instead of x , and $a = -1$, $b = 1$, and $c = 0$.

The graph of a quadratic function is a parabola, and the general form of a parabola is $y = a(x - h)^2 + k$. The parabola opens up if $a > 0$ and opens down if $a < 0$ and has a vertex at (h, k) . If the parabola opens up, it has its minimum at $x = h$, and the minimum value of the function is $y = k$. Similarly, if the parabola opens down, it has its maximum value of $y = k$, when $x = h$.

Using the method of completing the square, we can rewrite $y = ax^2 + bx + c$ in the form $y = a(x - h)^2 + k$ to show that $h = -\frac{b}{2a}$ and $k = c - \frac{b^2}{4a}$. When $a = -1$, $b = 1$, and $c = 0$, it follows that $h = 1/2$ and $k = 1/4$. So the value of p that maximizes pq is $p = 1/2$, and then $pq = 1/4$. This confirms the results of the graphical solution.

Calculus-based method

Advanced Placement students probably have had (or are taking) calculus, including tests for local extrema. For a function with continuous first and second derivatives, at an extremum, the first derivative equals 0, and the second derivative is either positive (at a minimum) or negative (at a maximum).

The first derivative of $f(p) = pq = p(1 - p)$ is given by

$$\begin{aligned}f'(p) &= \frac{d}{dp}[p(1 - p)] \\&= \frac{d}{dp}[-p^2 + p] \\&= -2p + 1\end{aligned}$$

Solve $f'(p) = 0$:

$$-2p + 1 = 0$$

$$p = \frac{1}{2}$$

Now find $f''\left(\frac{1}{2}\right)$:

$$\begin{aligned}f''(p) &= \frac{d}{dp}[f'(p)] \\&= \frac{d}{dp}(-2p + 1) \\&= -2\end{aligned}$$

$$\text{So } f''\left(\frac{1}{2}\right) = -2.$$

Since the second derivative is negative when the first derivative equals 0, $f(p)$ has a maximum at $p = \frac{1}{2}$.

This result has implications for confidence intervals for p ; see the Teaching Tips for Chapter 7.

CHAPTER 6: NORMAL CURVES AND SAMPLING DISTRIBUTIONS

Emphasize the differences between discrete and continuous random variables with examples of each.

Emphasize how normal curves can be used to approximate the probabilities of both continuous and discrete random variables, and in the cases when the distribution of a data set can be approximated by a normal curve, such a curve is defined by two quantities: the mean and standard deviation of the data. In such a case, the normal curve is defined by this equation:

$$y = \frac{e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}}{\sigma\sqrt{2\pi}}$$

Review Chebyshev's theorem from Chapter 3. Emphasize that this theorem implies that for *any* data set, at *least* 75% of the data lie within 2 standard deviations on each side of the mean, at *least* 88.9% of the data lie within 3 standard deviations on each side of the mean, and at *least* 93.75% of the data lie within 4 standard deviations on each side of the mean.

In comparison, a data set that has a distribution that is symmetric and bell-shaped, or in particular, an approximate normal distribution, is more restrictive in that

1. Approximately 68% of the data values lie within 1 standard deviation on each side of the mean,
2. Approximately 95% of the data values lie within 2 standard deviations on each side of the mean, and
3. Approximately 99.7% of the data values lie within 3 standard deviations on each side of the mean.

Remind students regularly that a z value equals the number of standard deviations from the mean for data values of *any* distribution approximated by a normal curve.

Emphasize the connection between the area under a normal curve and probability values of the random variable. That is, emphasize that the area under any normal curve equals 1 and that the percentage of area under the curve between given values of the random variable equals the probability that the random variable will be between these values. The values in a z table are areas *and* probability values.

Emphasize the differences between population parameters and sample statistics. Point out that when knowledge of the population is unavailable, then knowledge of a corresponding sample statistic must be used to make inferences about the population.

Emphasize the main two facts derived from the central limit theorem:

1. If x is a random variable with a normal distribution whose mean is μ and standard deviation is σ , then the means of random samples for any fixed-size n taken from the x distribution is a random variable $\bar{x}$ that has a normal distribution with mean μ and standard deviation $\sigma/\sqrt{n}$.
2. If x is a random variable with *any* distribution whose mean is μ and standard deviation is σ , then the mean of random samples of a fixed-size n taken from the x distribution is a random variable $\bar{x}$ that has a distribution that approaches a normal distribution with mean μ and standard deviation $\sigma/\sqrt{n}$ as n increases without limit.

Choosing sample sizes greater than 30 is an important point to emphasize in the situation mentioned in part 2 of the central limit theorem above. This commonly accepted convention ensures that the $\bar{x}$ distribution of part 2 will have an approximate normal distribution regardless of the distribution of the population from which these samples are drawn.

Emphasize that the central limit theorem allows us to infer facts about populations from sample means having normal distributions.

Emphasize the conditions whereby a binomial probability distribution (discussed in Chapter 5) can be approximated by a normal distribution: $np > 5$ and $n(1 - p) > 5$, where n is the number of trials and p is the probability of success in a single trial.

When a normal distribution is used to approximate a discrete random variable (such as the random variable of a binomial probability experiment), the *continuity correction* is an important concept to emphasize to students. A discussion of this important adjustment can be a good opportunity to compare discrete and continuous random variables.

Emphasize that facts about sampling distributions for proportions relating to binomial experiments can be inferred if the same conditions satisfied by a binomial experiment that can be approximated by a normal distribution are satisfied: $np > 5$ and $n(1 - p) > 5$, where n is the number of trials and p is the probability of success in a single trial.

Emphasize the difference in the continuity correction that must be taken into account in the sampling distribution for proportions and the continuity correction for a normal distribution used to approximate the probability distribution of the discrete random variable in a binomial probability experiment. That is, instead of subtracting 0.5 from the left endpoint and adding 0.5 to the right endpoint of an interval involved in a normal distribution approximating a binomial probability distribution, $0.5/n$ must be subtracted from the left endpoint and $0.5/n$ must be added to the right endpoint of such an interval when a normal distribution is used to approximate a sampling distribution for proportions.

CHAPTER 7: ESTIMATION

As the text says, nontrivial probability statements involve variables, not constants. And if the mean of a population is considered a constant, then the event that this mean falls in a certain range with known numerical bounds has either probability 1 or probability 0.

However, we might instead think of the population mean as itself a variable because, after all, the value of the mean is unknown initially. In other words, we may think of the population from which we are sampling as one of many populations—a population of populations, if you like. One of these populations has been randomly selected for us to work with, and we are trying to figure out which population it is or at least what its mean is.

If we think of our sampling activity in this way, we can then think of the event “The mean lies between a and b ” as having a nontrivial probability of being true. Can we now create a 90% confidence interval and then say that the mean has a 90% probability of being in that interval? It might seem so, but in general, the answer is no. Even though a procedure might have exactly a 90% success rate at creating confidence intervals that contain the mean, a confidence interval created by such a procedure will not, in general, have exactly a 90% chance of containing the mean.

How is this possible? To understand this paradox, let us turn from mean finding to a simpler task: guessing the color of a randomly drawn marble. Suppose that a sack contains some red marbles and some blue marbles. And suppose that we have a friend who will reach in, draw out a marble, and announce its color while we have our backs turned. The friend can be counted on to announce the correct color *exactly 90% of the time* and the wrong color the other 10% of the time. So, if the marble drawn is blue, the friend will say “blue” 9 times out of 10 and “red” 1 time. And likewise for the red marble. This is like creating a 90% confidence interval.

Now the friend reaches in, pulls out a marble, and announces, “blue.” Does this mean that we are 90% sure that the friend is holding a blue marble? *It depends on what we think about the mix of marbles in the bag.* Suppose that we think that the bag contains three red marbles and two blue ones. Then we expect the friend to draw a red marble $3/5$ of the time and announce “blue” 10% of those times, or $3/50$ of all draws. And we expect the friend to draw a blue marble $2/5$ of the time and announce “blue” 90% of those times, or $18/50$ of all draws. This means that the ratio of true “blue” announcements to false ones is 18:3, or 6:1. And thus we attach a probability of $6/7 = 85.7\%$, not 90%, to our friend’s announcement that the marble drawn is blue, even though we believe our friend to be telling the truth 90% of the time. For similar reasons, if the friend says “red,” we will attach a probability of 93.1% to this claim. Simply put, our initial belief that there are more red marbles than blue ones pulls our confidence in a “blue” announcement downward and our confidence in a “red” announcement upward from the 90% level.

Now, if we believe that there is an *equal* number of red and blue marbles in the bag, then, as it turns out, we will attach 90% probability to “blue” announcements and to “red” announcements as well. But *this is a special case*. In general, the probabilities we attach to each of our friend’s statements will be different from the frequency with which we think our friend is telling the truth. Furthermore, if we have *no idea* about the mix of marbles in the bag, then we will be *unable* to set probabilities for our friend’s statements because we will be unable to run the calculation for how often our friend’s “blue” statements are true and our friend’s “red” statements are true. In other words, *we cannot justify simply setting our probability equal to the test’s confidence level by default.*

This story has two morals. First, the probability of a statement is one thing, and the success rate of a procedure that tries to come up with true statements is another. Second, our prior beliefs about the conditions of an experiment are an unavoidable element in our interpretation of any sample data.

Let us apply these lessons to the business of finding confidence intervals for population means. When we create a 90% confidence interval, we will in general *not* be 90% sure that the interval contains the mean. It could *happen* to turn out that we were 90% sure, but this will depend on what ideas we had about the population mean going in. Suppose that we were fairly sure, to start with, that the population mean lay somewhere between 10 and 20, and suppose that we then took a sample that led to the construction of a 90% confidence interval that ran from 31–46. We would *not* conclude, with 90% certainty, that the mean lay between 31 and 46. Instead, we would have a probability lower than that because we previously thought the mean was outside that range. At the same time, we would be much more ready to believe that the mean lay between 31 and 46 than we were before because, after all, a procedure with a 90% success rate produced that prediction. Our exact probability for the “between 31 and 46” statement would depend on our entire initial probability distribution for values of the population mean—something we would have a hard time coming up with if the question were put to us. Thus, under normal circumstances, our exact level of certainty about the confidence interval could not be calculated.

So the general point made in the text holds even if we think of a population mean as a variable. The procedure for

finding a confidence interval of confidence level c does not, in general, produce a statement (about the value of a population mean) that has a probability c of being true.

Confidence Intervals for p (Section 7.3)

The result obtained in the Teaching Tip for Chapter 5 has implications for the confidence interval for p . The most conservative interval estimate of p , the widest possible confidence interval in a given situation, is obtained when

$$E = z_c \sqrt{pq/n} \text{ is calculated using } p = 1/2.$$

CHAPTER 8: HYPOTHESIS TESTING

P -Value Method vs. Critical Region Method

The most popular method of statistical testing is the P -value method. For this reason, the P -value method is emphasized in this book.

The P -value method was used extensively by a famous statistician, R. A. Fisher, and is the most popular method of testing in use today. At the end of Section 8.2, another method of testing called the *critical region method* is presented. It was used extensively by statisticians J. Neyman and E. Pearson. In recent years, the use of this method has been declining.

The critical region method for hypothesis testing is convenient when distribution tables are available for finding critical values. However, most statistical software and research journal articles give P values rather than critical values. Most fields of study that require statistics want students to be able to use P values.

Emphasize that for a fixed, preset level of significance α , both methods are logically equivalent.

What a Hypothesis Test Tells Us (Sections 8.1–8.3)

The procedure for hypothesis testing with significance levels may confuse some students at first, especially because the levels are chosen somewhat arbitrarily. Why, the students may wonder, don't we just calculate the likelihood that the null hypothesis is true? Or is that really what we're doing when we find the P value?

Once again, we run the risk of confusion over the role of probability in our statistical conclusions. The P value is *not* the same thing as the probability, in light of the data, of the null hypothesis. Instead, the P value is the probability that the data would turn out the way they did, assuming that the null hypothesis is true. Just as with confidence intervals, we have to be careful not to think that we are finding the probability of a given statement when we are in fact doing something else.

To illustrate, consider two coins in a sack, one fair and one two-headed. One of these coins is pulled out at random and flipped. It comes up heads. Let us take as our null hypothesis "The flipped coin was the fair one." The probability of the outcome, given the null hypothesis, is $1/2$ because a fair coin will come up heads half the time. This probability is in fact the P value of the outcome. On the other hand, the probability that the null hypothesis is true, given the evidence, is $1/3$ because out of all the heads outcomes one will see in many such trials, $1/3$ are from the fair coin.

Now suppose that instead of containing two coins of known character, the sack contains an unknown mix—some fair coins, some two-headed coins, and possibly some two-tailed coins as well. Then we still can calculate the P value of a heads outcome because the probability of "heads" with a fair coin is still $1/2$. But the probability that the coin is fair, given that we're seeing heads, *cannot be calculated* because we know nothing about the mix of coins in the bag. So the P value of the outcome is one thing, and the probability of the null hypothesis is another.

The lesson now should be familiar. Without some prior idea about the character of an experiment, either based on a theoretical model or based on previous outcomes, we cannot attach a definite probability to a statement about the experimental setup or its outcome.

This is the usual situation in hypothesis testing. We normally lack the information needed to calculate probabilities for the null hypothesis and its alternative. What we do instead is to take the null hypothesis as defining a well-understood scenario from which we *can* calculate the likelihoods of various outcomes—the probabilities of various kinds of sample results, given that the null hypothesis is true. By contrast, the alternative hypothesis includes all sorts of scenarios, in some of which (for instance) two population means are only slightly different, in others of which the two means are far apart, and so on. Unless we have identified the likelihoods of all these possibilities relative to each other and to the null hypothesis, we will not have the background information needed to calculate the probability of the null hypothesis from sample data.

In fact, we will not have the data necessary to calculate the power, $1 - \beta$, of a hypothesis test. Finding the power requires knowing the H_1 distribution. Because we cannot specify the H_1 distribution when we are concerned with things such as diagnosing disease (instead of drawing coins from a sack and the like), we normally cannot determine the probability of the null hypothesis in light of the evidence. Instead, we have to content ourselves with quantifying the risk α of rejecting the hypothesis when it is true.

A Paradox About Hypothesis Tests

The way hypothesis tests work leads to a result that at first seems surprising. It sometimes can happen that, at a given level of significance, a one-tailed test leads to rejection of the null hypothesis, whereas a two-tailed test does not. Apparently, one can be justified in concluding that $\mu > k$ (or $\mu < k$ as the case may be) but not justified in concluding that $\mu \neq k$ —even though the latter conclusion follows from the former! What is going on here?

This paradox dissolves when one remembers that a one-tailed test is used only when one has appropriate information. With the null hypothesis $H_0: \mu = k$, we choose the alternative hypotheses $H_1: \mu > k$ only if we *are already sure* that μ is not less than $H_1: \mu < k$. In effect, this assumption boosts the force of any evidence that μ does not equal k —and if it is not less than or equal to k , it must be greater.

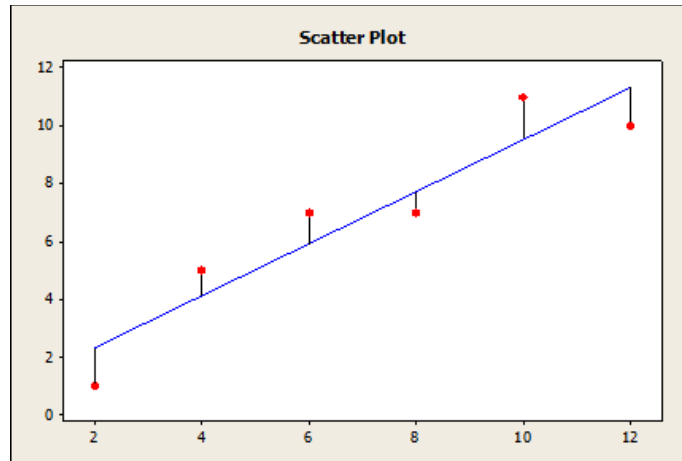
In other words, when a right-tailed test is appropriate, rejecting the null hypothesis means concluding *both* that $\mu > k$ and that $\mu \neq k$. But when there is no justification for a one-tailed test, one must use a two-tailed test and must have somewhat stronger evidence before concluding that $\mu \neq k$.

CHAPTER 9: CORRELATION AND REGRESSION

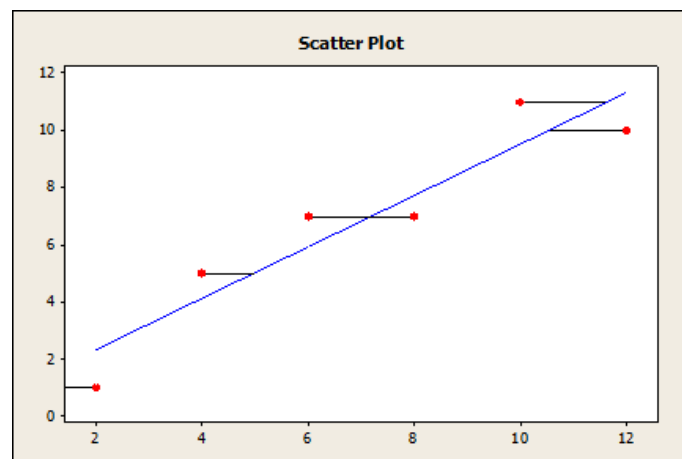
Least-Squares Criteria (Section 9.2)

With some sets of paired data, it will not be obvious which is the explanatory variable and which is the response variable. Here it may be worth mentioning that for linear regression, the choice matters. The results of a linear regression analysis will differ depending on which variable is chosen as the explanatory variable and which is chosen as the response variable. This is not immediately obvious. We might think that with x as the explanatory variable, we could just solve the regression equation $y = a + bx$ for x in terms of y to obtain the regression equation that we would get if we took y as the explanatory variable instead. But this would be a mistake.

The figure below shows the vertical distances from data points to the line of best fit. The line is defined so as to make the sum of the squares of these vertical distances as small as possible.



Now the next figure shows the *horizontal* distances from the data points to the same line. These are the distances whose sum of squares would be minimized if the explanatory and response variables switched roles. With such a switch, the graph would be flipped over, and the horizontal distances would become vertical ones. But the line that minimizes the sum of squares for vertical distances is not, in general, the same line that minimizes the sum of squares for horizontal distances.



So there is more than one way, mathematically, to define the line of best fit for a set of paired data. This raises a question: What is the *proper* way to define the line of best fit?

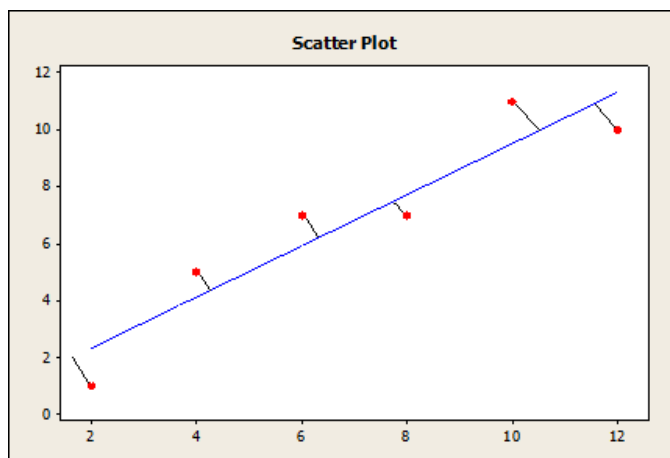
Let us turn this question around. Under what circumstances is a best fit based on *vertical* distances the right way to go? Well, intuitively, the distance from a data point to the line of best fit represents some sort of deviation from the ideal value. We can conceptualize this most easily in terms of measurement error. If we treat the error as a strictly vertical distance, then we are saying that in each data pair, the second value is possibly off, but the first value is exactly correct. In other words, the least-squares method with vertical distances assumes that the first value in each data pair is measured with essentially perfect accuracy, whereas the second is measured only imperfectly.

An illustration shows how these assumptions can be realistic. Suppose that we are measuring the explosive force generated by the ignition of varying amounts of gunpowder. The weight of the gunpowder is the explanatory variable, and the force of the resulting explosion is the response variable. It could easily happen that we were able to measure the weight of gunpowder with great exactitude—down to the thousandth ounce—but that our means of measuring explosion force was quite crude, such as the height to which a wooden block was projected into the air by the explosion. We then would have an experiment with a good deal of error in the response variable measurement

but for, practical purposes, no error in the explanatory variable measurement. This would all be perfectly in accord with the vertical-distance criterion for finding the line of best fit by the least-squares method.

But now consider a different version of the gunpowder experiment. This time we have a highly refined means of measuring explosive force (some sort of electronic device, let us say), and at the same time we have only a very crude means of measuring gunpowder mass (perhaps a rusty pan balance). In this version of the story, the error would be in the measurement of the explanatory variable, and a horizontal least-squares criterion would be called for.

Now, the most common situation is one in which both the explanatory and the response variables contain some error. The preceding discussion suggests that the most appropriate least-squares criterion for goodness of fit for a line through the cluster of data points would be a criterion in which error was represented as a line lying at some slant, as in the figure below.



To apply such a criterion, we would have to figure out how to define distance in two dimensions when the x and y axes have different units of measure. We will not try to solve that puzzle here. Instead, we just summarize what we have learned. There is more than one least-squares criterion for fitting a line to a bivariate data set, and the choice of which criterion to use implies an assumption about which variable(s) is affected by the error (or other deviation) that moves points off the line representing ideal results.

And finally, we now see that the standard use of vertical distances in the least-squares method *implies an assumption that the error is predominantly in the response variable*. This is often a reasonable assumption to make because the explanatory variable is frequently a *control* variable—i.e., a variable under the experimenter's control and that thus generally is capable of being adjusted with a fair amount of precision. The response variable, by contrast, is the one that simply must be measured and that cannot be fine-tuned through an experimenter's adjustment. However, it is worth noting that this is only the typical relationship, not a necessary one (as the second gunpowder scenario shows).

Finally, it is also worth noting that both the vertical and horizontal least-squares criteria will produce a line that passes through the point $(\bar{x}, \bar{y})$. Thus the vertical and horizontal least-squares lines must either coincide (which is atypical but not impossible) or intersect at $(\bar{x}, \bar{y})$. The other property the two lines have in common is the correlation coefficient r . It is easy to see, looking at the formula for r , that the value of r does not depend on which variable is chosen as the explanatory one and which is chosen as the response one.

Variables and the Issue of Cause and Effect (Section 9.2)

The relationship between two measured variables x and y may not, in physical terms, be one of cause and effect, respectively. It often is, of course, but instead it may happen that y is the cause and x is the effect. Note that in an example where x = cricket chirps per second and y = air temperature, y is obviously the cause and x the effect. In other situations, x and y will be two effects of a common, possibly unknown, cause. For example, x might be a patient's blood sugar level, and y might be the patient's body temperature. Both these variables could be caused by an illness, which might be quantified in terms of a count of bacterial activity. The point to remember is that although the x -causes- y scenario is typical, strictly speaking, the designations *explanatory variable* and *response variable* should be understood not in terms of a causal relationship but in terms of which quantity is known initially and which is inferred.

CHAPTER 10: CHI-SQUARE AND F DISTRIBUTIONS

Emphasize that both the χ^2 distribution and the F distribution are not symmetric and have only nonnegative values.

Emphasize that the applications of the χ^2 distribution include the test for independence of two factors, goodness of fit of a present distribution to a given distribution, and whether a variance (or standard deviation) has changed or varies from a known population variance (or standard deviation). The χ^2 distribution is also used to find a confidence interval for a variance (or standard deviation).

Emphasize that applications of the F distribution include the test of whether the variances (or equivalently, standard deviations) of two independent, normal distributions are equal. A second application of the F distribution is the one-way ANOVA test, which determines whether a significant difference exists between any of several sample means of groups taken from populations that are each assumed to be normally distributed, independent of one another, and in which the groups come from distributions with approximately the same standard deviation. A third application of the F distribution is a two-way ANOVA test: a test of whether differences exist in the population means of varying levels of two factors where each level of each factor is assumed to be from a normal distribution and where all levels of both factors are assumed to have equal variances.

CHAPTER 11: NONPARAMETRIC STATISTICS

Review the classifications of data discussed in Chapter 1: ratio, interval, ordinal, and nominal.

Emphasize that the methods of nonparametric statistics are quite general and are applied when no assumptions are made about the population distributions from which samples are drawn, such as that the distributions are normal or binomial, for example.

Emphasize that the sign test is used when comparing sample distributions from two populations that are not independent, such as when a sample is measured twice, as in a "before-and-after" study. Emphasize that the sign test requires that the number of positive and negative signs between the samples number at least 12. Point out that since the proportion of plus signs to total number of plus and minus signs of the sampling distribution for x follows a normal distribution, the critical values for the sign test are based on z values from a normal distribution.

Emphasize that the rank-sum test for testing the difference between sample means can be used when it is not known whether the populations the samples come from are normally distributed or when assumptions about equal population variances are not satisfied. An important point to emphasize is that the rank-sum test requires that the sample size of each sample be at least 11. Emphasize that since the sampling distribution for the sum of ranks R follows a normal distribution, the critical values and sample statistics of the test are z values from a normal distribution.

Emphasize that the Spearman rank correlation is used to compare ranked data from two sources.

Emphasize that for the Spearman rank correlation coefficient r_s , $-1 \leq r_s \leq 1$, and discuss the meanings of $r_s = 1$, $r_s = -1$, $r_s = 0$, r_s close to 1, and r_s close to -1 .

Compare the similarity of r_s to the correlation coefficient r from Chapter 10.

The runs test for randomness is a very useful nonparametric test.

H_0 : The symbols are randomly mixed in the sequence.

H_1 : The symbols are not randomly mixed in the sequence.

In Section 11.4 we restrict our attention to a sequence of two symbols. Any sequence of numbers can be converted to a sequence of symbols A for above the median and B for below the median. Remind students that Table 9 in Appendix II provides critical values for $\alpha = 0.05$ only. There are tables available for other levels of significance. However, in this text, we restrict α to 0.05 when n_1 and n_2 are both less than or equal to 20.

Hints for Distance Education Courses

Distance education uses various media, each of which can be used in a one-way or interactive mode. Here is a representative list:

		One-Way	Interactive
Medium:	Audio	MP3 files	Phone
	Audiovisual	DVDs, MP4	Videoconferencing
	Data	Computer-resident tutorials, Web tutorials	Social Media, E-mail, chat, discussion boards
	Print	Texts, workbooks	Mailed-in assignments, mailed-back instructor comments, fax exchanges

Sometimes the modes are given as asynchronous (students working on their schedules) versus synchronous (students and instructors working at the same time), but synchronous scheduling normally makes sense only when this enables some element of interactivity in the instruction.

Naturally, the media and modes may be mixed and matched. A course might, for instance, use a one-way video feed with interactive audio plus discussion lists.

THINGS TO KEEP IN MIND

In many online courses, printed material is a foundational part of the instruction. In an online course, the textbook or e-book is *at least* as important as in a traditional course because it is the one resource that requires no special equipment to use and whose use is not made more difficult by the distance separating student and instructor.

Because students generally obtain all course materials at once, before instruction begins, midcourse adjustments of course content generally are not practicable. Plan the course carefully up front so that everything is in place when instruction begins.

In online courses, students can be assumed to have ready access to computers while working on their own. This creates the opportunity for technology-based assignments that in a traditional course might be feasible at best as optional work (e.g., assignments using SPSS, MINITAB, or Microsoft Excel; see the corresponding guides that accompany *Understandable Statistics*). However, any time students have to spend learning how to use unfamiliar software will add to their overall workload and possibly to their frustration level. Remember this when choosing technology-based work to incorporate.

Remember that even (and perhaps especially) in online education, students take a course because they want to interact with a human being rather than just read a book. The goal of distance instruction is to make that possible for students who cannot enroll in a traditional course. Lectures should not turn into slide shows with voice commentary, even though these may be technologically easier to transmit than, say, real-time video. Keep the human element uppermost.

All students should be self-motivated, but in real life, nearly all students benefit from a little friendly supervision and encouragement. This goes double for distance education. Make an extra effort to check in with students one-on-one, to ask how things are going, and to remind them of things they may be forgetting or neglecting.

CHALLENGES IN DISTANCE EDUCATION

Technology malfunctions often plague distance courses. To prevent this from happening in yours:

- Don't take on too much at once. As the student sites multiply, so do the technical difficulties. Try the methodology with one or two remote sites before expanding.
- Plan all technology use well in advance and thoroughly test all equipment before the course starts.
- Have redundant and backup means for conducting class sessions. If, for instance, a two-way videoconferencing link goes down, plan for continuing the lecture by speakerphone, with students referring to predistributed printed materials as needed.
- Allow enough slack time in lectures for extra logistical tasks and occasional technical difficulties.
- If possible, do a precourse dry run with at least some of the students that so they can get familiar with the equipment and procedures and alert you to any difficulties they run into.
- When it is feasible, have a facilitator at each student site. This person's main job is to make sure that the technology at the student end works smoothly. If the facilitator can assist with course administration and answer student questions about course material, so much the better.

In a distance course, establishing rapport with students and making them comfortable can be difficult.

- An informal lecture style, often effective in traditional classrooms, can be even more effective in a distance course. Be cheerful and use humor. (In cross-cultural contexts, though, remember that what is funny to you may fall flat with your audience.)
- Remember that your voice will not reach the students with the same clarity it would in a traditional classroom. Speak clearly, not too fast, and avoid overly long sentences. Pause regularly.
- Early in the course, work in some concrete, real-world examples and applications to help the students relax, roll up their sleeves, and forget about the distance learning aspect of the course.
- If the course is interactive, via teleconferencing or real-time typed communication, get students into "send" mode as soon as possible. Ask them questions. Call on individuals if you judge that this is appropriate.
- A student site assistant with a friendly manner also can help students to settle into the course quickly.

The distance learning format will make it hard for you to gauge how well students are responding to your instruction. In a traditional course, student incomprehension or frustration often registers in facial expression, tone of voice, or muttered comments—all of which are, depending on the instructional format, either difficult or impossible to pick up on in a distance course. Have some way for students to give feedback on how well the course is going for them. Possibilities:

- Quickly written surveys ("On a scale of 1 to 5, please rate . . .") every few weeks.
- Periodic "How are things going?" phone calls or messages from you.
- A student site assistant can act as your "eyes and ears" for this aspect of the instruction, and students may be more comfortable voicing frustrations to him or her than to you.

Cheating is a problem in any course but especially so in distance courses. Once again, an on-site facilitator is an asset. Another means of forestalling cheating is to have open-book exams, which takes away the advantage of sneaking a peek at the text.

Good student–instructor interaction takes conscious effort and planning in a distance course. Provide students with a variety of ways to contact you:

- Social Media is the handiest way for most students to stay in touch.
- Phone. When students are most likely to be free in the evenings, set the number up for your home address, and schedule evening office hours during which students can count on reaching you.
- When students can make occasional in-person visits to your office, provide for this as well.

ADVANTAGES IN DISTANCE EDUCATION

Compared with traditional courses, more of the information shared in a distance course is, or can be, preserved for later review. Students can review videotaped lectures, instructor–student exchanges via e-mail can be reread, class discussions are on a reviewable discussion list, and so on.

To the extent that students are on their own, working out of texts or watching prerecorded video, course material can be modularized and customized to suit the needs of individual students. Where a traditional course would necessarily be offered as a 4-unit lecture series, the counterpart distance course could be broken into four 1-unit modules, with students free to take only the modules they need. This is especially beneficial when the course is aimed at students who already have some professional experience with statistics and need to fill in gaps rather than comprehensive instruction.

STUDENT INTERACTION

Surprisingly, some instructors have found that students interact more with one another in a well-designed distance course than in a traditional course, even when the students are physically separated from one another. Part of the reason may be a higher level of motivation among distance learners. But another reason is that the same technologies that facilitate student–instructor communication—such things as e-mail and discussion lists—also facilitate student–student communication. In some cases, distance learners actually have done better than traditional learners taking the very same course. Better student interaction was thought to be the main reason.

One implication of this greater student–student interaction is that while group projects involving statistical evaluations of real-world data might seem more difficult to set up in a distance course, they are actually no harder, and the students learn just as much. The Web has many real-world data sources, such as the U.S. Department of Commerce (**home.doc.gov**), which has links to the U.S. Census Bureau (**www.census.gov**); the Bureau of Economic Analysis (**www.bea.gov**); and other agencies that compile publicly available data.

Suggested References

The American Statistical Association

Contact Information

1429 Duke Street
Alexandria, VA 22314-3415
Phone: (703) 684-1221 or toll-free: (888) 231-3473
Fax: (703) 684-2037
www.amstat.org

ASA Publications

Stats: The Magazine for Students of Statistics
CHANCE magazine
The American Statistician
AmStat News

BOOKS

- Huff, Darryll, and Irving Geis. *How to Lie with Statistics*. New York: WW Norton & Co, 1954.
Classic text on the use and misuse of statistics.
- Moore, David S. *Statistics: Concepts and Controversies*, 5th ed. New York: W.H. Freeman Inc., 2005.
Does not go deeply into computational aspects of statistical methods. Good resource for emphasizing concepts and applications.
- Tufte, Edward R. *The Visual Display of Quantitative Information*, 2nd ed. Cheshire, CT: Graphics Press, 2001. A beautiful book, the first of three by Tufte on the use of graphical images to summarize and interpret numerical data. The books are virtual works of art in their own right.
- Tanur, Judith M., et al. *Statistics: A Guide to the Unknown*, 3rd ed. Pacific Grove, CA: Duxbury Press, 1988.
Another excellent source of illustrations.

REFERENCES FOR DISTANCE LEARNING

- Bolland, Thomas W. (1994). "Successful Customers of Statistics at a Distant Learning Site." In *Proceedings of the Quality and Productivity Section, American Statistical Association*, pp. 300–304.
- Lawrence, Betty, and Leonard M. Gaines. (1997). "An Evaluation of the Effectiveness of an Activity-Based Approach to Statistics for Distance Learners." In *Proceedings of the Section on Statistical Education, American Statistical Association*, pp. 271–272.
- Wegman, Edward J., and Jeffrey L. Solka (1999). "Implications for Distance Learning: Methodologies for Statistical Education." In *Proceedings of the Section on Statistical Education, the Section on Teaching Statistics in the Health Sciences, and the Section on Statistical Consulting, American Statistical Association*, pp. 13–16.
- "Distance Learning: Principles for Effective Design, Delivery, and Evaluation." University of Idaho web site: **www.uidaho.edu/eo/distgla.html**.

Part II

Frequently Used Formulas and Transparency Masters

Frequently Used Formulas

n = sample size N = population size f = frequency

Chapter 2

Class width = $\frac{\text{high} - \text{low}}{\text{number classes}}$ (increase to next integer)

Class midpoint = $\frac{\text{upper limit} + \text{lower limit}}{2}$

Lower boundary = lower boundary of previous class + class width

Chapter 3

Sample mean $\bar{x} = \frac{\sum x}{n}$

Population mean $\mu = \frac{\sum x}{N}$

Weighted average = $\frac{\sum xw}{\sum w}$

Range = largest data value – smallest data value

Sample standard deviation $s = \sqrt{\frac{\sum (x - \bar{x})^2}{n - 1}}$

Computation formula $s = \sqrt{\frac{\sum x^2 - (\sum x)^2 / n}{n - 1}}$

Population standard deviation $\sigma = \sqrt{\frac{\sum (x - \mu)^2}{N}}$

Sample variance s^2

Population variance σ^2

Sample coefficient of variation $CV = \frac{s}{\bar{x}} \cdot 100$

Sample mean for grouped data $\bar{x} = \frac{\sum xf}{n}$

Sample standard deviation for grouped data $s = \sqrt{\frac{\sum (x - \bar{x})^2 f}{n - 1}} = \sqrt{\frac{\sum x^2 f - (\sum xf)^2 / n}{n - 1}}$

Chapter 4

Probability of the complement of event A $P(A^c) = 1 - P(A)$

Multiplication rule for independent events $P(A \text{ and } B) = P(A) \cdot P(B)$

General multiplication rules $P(A \text{ and } B) = P(A) \cdot P(B|A)$
 $P(A \text{ and } B) = P(B) \cdot P(A|B)$

Addition rule for mutually exclusive events $P(A \text{ or } B) = P(A) + P(B)$

General addition rule $P(A \text{ or } B) = P(A) + P(B) - P(A \text{ and } B)$

Permutation rule $P_{n,r} = \frac{n!}{(n-r)!}$

Combination rule $C_{n,r} = \frac{n!}{r!(n-r)!}$

Chapter 5

Mean of a discrete probability distribution $\mu = \sum xP(x)$

Standard deviation of a discrete probability distribution $\sigma = \sqrt{\sum (x - \mu)^2 P(x)}$

Given $L = a + bx$ $\mu_L = a + b\mu$
 $\sigma_L = |b|\sigma$

Given $W = ax_1 + bx_2$ (x_1 and x_2 independent) $\mu_W = a\mu_1 + b\mu_2$
 $\sigma_W = \sqrt{a^2\sigma_1^2 + b^2\sigma_2^2}$

For Binomial Distributions

r = number of successes; p = probability of success; $q = 1 - p$

Binomial probability distribution $P(r) = C_{n,r} p^r q^{n-r}$

Mean $\mu = np$

Standard deviation $\sigma = \sqrt{npq}$

Geometric Probability Distribution

n = number of trial on which first success occurs $P(n) = p(1-p)^{n-1}$

Poisson Probability Distribution

r = number of successes

λ = mean number of successes over given interval $P(\lambda) = \frac{e^{-\lambda} \lambda^r}{r!}$

Chapter 6

Raw score $x = z\sigma + \mu$

Standard score $z = \frac{x - \mu}{\sigma}$

Mean of $\bar{x}$ distribution $\mu_{\bar{x}} = \mu$

Standard deviation of $\bar{x}$ distribution $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$

Standard score for $\bar{x}$ $z = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}}$

Mean of $\hat{p}$ distribution $\mu_{\hat{p}} = p$

Standard deviation of $\hat{p}$ distribution $\sigma_{\hat{p}} = \sqrt{\frac{pq}{n}}; q = 1 - p$

Chapter 7

Confidence Interval

For μ

$$\bar{x} - E < \mu < \bar{x} + E$$

where $E = z_c \frac{\sigma}{\sqrt{n}}$ when σ is known

$$E = t_c \frac{s}{\sqrt{n}} \text{ when } \sigma \text{ is unknown}$$

with $d.f. = n - 1$

For p ($np > 5$ and $n(1 - p) > 5$)

$$\hat{p} - E < p < \hat{p} + E$$

where $E = z_c \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$

$$\hat{p} = \frac{r}{n}$$

For $\mu_1 - \mu_2$ (independent samples)

$$(\bar{x}_1 - \bar{x}_2) - E < \mu_1 - \mu_2 < (\bar{x}_1 - \bar{x}_2) + E$$

where $E = z_c \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$ when σ_1 and σ_2 are known

$$E = t_c \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} \text{ when } \sigma_1 \text{ or } \sigma_2 \text{ is unknown}$$

with $d.f.$ = smaller of $n_1 - 1$ and $n_2 - 1$

(Note: Software uses Satterthwaite's approximation for degrees of freedom $d.f.$)

For difference of proportions $p_1 - p_2$

$$(\hat{p}_1 - \hat{p}_2) - E < p_1 - p_2 < (\hat{p}_1 - \hat{p}_2) + E$$

$$\text{where } E = z_c \sqrt{\frac{\hat{p}_1 \hat{q}_1}{n_1} + \frac{\hat{p}_2 \hat{q}_2}{n_2}}$$

$$\hat{p}_1 = r_1 / n_1; \hat{p}_2 = r_2 / n_2; \hat{q}_1 = 1 - \hat{p}_1; \hat{q}_2 = 1 - \hat{p}_2$$

Sample Size for Estimating

$$\text{Means } n = \left(\frac{z_c \sigma}{E} \right)^2$$

$$\text{Proportions } n = p(1-p) \left(\frac{z_c}{E} \right)^2 \text{ with preliminary estimate for } p$$

$$n = \frac{1}{4} \left(\frac{z_c}{E} \right)^2 \text{ without preliminary estimate for } p$$

Chapter 8

Sample Test Statistics for Tests of Hypotheses

$$\text{For } \mu (\sigma \text{ known}) \quad z = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}}$$

$$\text{For } \mu (\sigma \text{ unknown}) \quad t = \frac{\bar{x} - \mu}{s / \sqrt{n}}; d.f. = n - 1$$

$$\text{For } p (np > 5 \text{ and } nq > 5) \quad z = \frac{\hat{p} - p}{\sqrt{pq/n}}, \text{ where } q = 1 - p; \hat{p} = r/n$$

$$\text{For paired differences } d \quad t = \frac{\bar{d} - \mu_{\bar{d}}}{s_d / \sqrt{n}}; d.f. = n - 1$$

$$\text{Difference of means } (\sigma_1 \text{ and } \sigma_2 \text{ known}) \quad z = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

Difference of means (σ_1 and σ_2 unknown) $z = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$; $d.f. = \text{smaller of } n_1 - 1 \text{ and } n_2 - 1$

(Note: Software uses Satterthwaite's approximation for degrees of freedom $d.f.$)

Difference of proportions $z = \frac{\hat{p}_1 - \hat{p}_2}{\sqrt{\frac{\bar{p}\bar{q}}{n_1} + \frac{\bar{p}\bar{q}}{n_2}}}$, where $\bar{p} = \frac{r_1 + r_2}{n_1 + n_2}$ and $\bar{q} = 1 - \bar{p}$; $\hat{p}_1 = r_1 / n_1$; $\hat{p}_2 = r_2 / n_2$

Chapter 9

Regression and Correlation

Pearson product moment correlation coefficient: $r = \frac{n \sum xy - (\sum x)(\sum y)}{\sqrt{n \sum x^2 - (\sum x)^2} \sqrt{n \sum y^2 - (\sum y)^2}}$

Least-squares line $\hat{y} = a + bx$, where $b = \frac{n \sum xy - (\sum x)(\sum y)}{n \sum x^2 - (\sum x)^2}$ and $a = \bar{y} - b\bar{x}$

Coefficient of determination $= r^2$

Sample test statistic for r $t = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}}$ with $d.f. = n - 2$

Standard error of estimate $S_e = \sqrt{\frac{\sum y^2 - a \sum y - b \sum xy}{n - 2}}$

Confidence interval for y $\hat{y} - E < y < \hat{y} + E$, where $E = t_c S_e \sqrt{1 + \frac{1}{n} + \frac{(x - \bar{x})^2}{n \sum x^2 - (\sum x)^2}}$ with $d.f. = n - 2$

Sample test statistic for slope b $t = \frac{b}{S_e} \sqrt{\sum x^2 - \frac{1}{n}(\sum x)^2}$ with $d.f. = n - 2$

Confidence interval for β $b - E < \beta < b + E$, where $E = \frac{t_c S_e}{\sqrt{\sum x^2 - \frac{1}{n}(\sum x)^2}}$ with $d.f. = n - 2$

Chapter 10

$\chi^2 = \sum \frac{(O - E)^2}{E}$ where $E = \frac{(\text{row total})(\text{column total})}{\text{sample size}}$

Tests of independence $d.f. = (R - 1)(C - 1)$

Tests of homogeneity $d.f. = (R - 1)(C - 1)$

Goodness of fit $d.f. = (\text{number of entries}) - 1$

Confidence interval for σ^2 ; $d.f. = n - 1$ $\frac{(n-1)s^2}{\chi_U^2} < \sigma^2 < \frac{(n-1)s^2}{\chi_L^2}$

Sample test statistic for σ^2 $\chi^2 = \frac{(n-1)s^2}{\sigma^2}$ with $d.f. = n - 1$

Testing Two Variances

Sample test statistic $F = \frac{s_1^2}{s_2^2}$, where $s_1^2 \geq s_2^2$; $d.f._N = n_1 - 1$; $d.f._D = n_2 - 1$

ANOVA

k = number of groups; N = total sample size

$$SS_{TOT} = \sum x_{TOT}^2 - \frac{(\sum x_{TOT})^2}{N}; SS_{BET} = \sum_{\text{all groups}} \left[\frac{(\sum x_i)^2}{n_i} \right] - \frac{(\sum x_{TOT})^2}{N}; SS_W = \sum_{\text{all groups}} \left[\sum x_i^2 - \frac{(\sum x_i)^2}{n_i} \right]$$

$$SS_{TOT} = SS_{BET} + SS_W; MS_{BET} = \frac{SS_{BET}}{d.f._{BET}} \text{ where } d.f._{BET} = k - 1; MS_W = \frac{SS_W}{d.f._W} \text{ where } d.f._W = N - k;$$

$$F = \frac{MS_{BET}}{MS_W} \text{ where } d.f. \text{ numerator} = d.f._{BET} = k - 1; d.f. \text{ denominator} = d.f._W = N - k$$

Two-Way ANOVA

r = number of rows; c = number of columns

Row factor F : $\frac{MS \text{ row factor}}{MS \text{ error}}$; column factor F : $\frac{MS \text{ column factor}}{MS \text{ error}}$; interaction F : $\frac{MS \text{ interaction}}{MS \text{ error}}$

with degrees of freedom for

Row factor = $r - 1$; interaction = $(r - 1)(c - 1)$; column factor = $c - 1$; error = $rc(n - 1)$

Chapter 11

Sample test statistic for x = proportion of plus signs to all signs ($n \geq 12$) $z = \frac{x - 0.5}{\sqrt{0.25/n}}$

Sample test statistic for R = sum of ranks

$$z = \frac{R - \mu_R}{\sigma_R}, \quad \text{where } \mu_R = \frac{n_1(n_1 + n_2 + 1)}{2} \text{ and } \sigma_R = \sqrt{\frac{n_1 n_2 (n_1 + n_2 + 1)}{12}}$$

Spearman rank correlation coefficient $r_s = 1 - \frac{6 \sum d^2}{n(n^2 - 1)}$, where $d = x - y$

Sample test statistic for runs test R = number of runs in sequence

Transparency Masters

TABLE A Areas of a Standard Normal Distribution (Alternate Version of Appendix II Table 5)

The table entries represent the area under the standard normal curve from 0 to the specified value of z.										
z	.00	.01	.02	.03	.04	.05	.06	.07	.08	.09
0.0	.0000	.0040	.0080	.0120	.0160	.0199	.0239	.0279	.0319	.0359
0.1	.0398	.0438	.0478	.0517	.0557	.0596	.0636	.0675	.0714	.0753
0.2	.0793	.0832	.0871	.0910	.0948	.0987	.1026	.1064	.1103	.1141
0.3	.1179	.1217	.1255	.1293	.1331	.1368	.1406	.1443	.1480	.1517
0.4	.1554	.1591	.1628	.1664	.1700	.1736	.1772	.1808	.1844	.1879
0.5	.1915	.1950	.1985	.2019	.2054	.2088	.2123	.2157	.2190	.2224
0.6	.2257	.2291	.2324	.2357	.2389	.2422	.2454	.2486	.2517	.2549
0.7	.2580	.2611	.2642	.2673	.2704	.2734	.2764	.2794	.2823	.2852
0.8	.2881	.2910	.2939	.2967	.2995	.3023	.3051	.3078	.3106	.3133
0.9	.3159	.3186	.3212	.3238	.3264	.3289	.3315	.3340	.3365	.3389
1.0	.3413	.3438	.3461	.3485	.3508	.3531	.3554	.3577	.3599	.3621
1.1	.3643	.3665	.3686	.3708	.3729	.3749	.3770	.3790	.3810	.3830
1.2	.3849	.3869	.3888	.3907	.3925	.3944	.3962	.3980	.3997	.4015
1.3	.4032	.4049	.4066	.4082	.4099	.4115	.4131	.4147	.4162	.4177
1.4	.4192	.4207	.4222	.4236	.4251	.4265	.4279	.4292	.4306	.4319
1.5	.4332	.4345	.4357	.4370	.4382	.4394	.4406	.4418	.4429	.4441
1.6	.4452	.4463	.4474	.4484	.4495	.4505	.4515	.4525	.4535	.4545
1.7	.4554	.4564	.4573	.4582	.4591	.4599	.4608	.4616	.4625	.4633
1.8	.4641	.4649	.4656	.4664	.4671	.4678	.4686	.4693	.4699	.4706
1.9	.4713	.4719	.4726	.4732	.4738	.4744	.4750	.4756	.4761	.4767
2.0	.4772	.4778	.4783	.4788	.4793	.4798	.4803	.4808	.4812	.4817
2.1	.4821	.4826	.4830	.4834	.4838	.4842	.4846	.4850	.4854	.4857
2.2	.4861	.4864	.4868	.4871	.4875	.4878	.4881	.4884	.4887	.4890
2.3	.4893	.4896	.4898	.4901	.4904	.4906	.4909	.4911	.4913	.4916
2.4	.4918	.4920	.4922	.4925	.4927	.4929	.4931	.4932	.4934	.4936
2.5	.4938	.4940	.4941	.4943	.4945	.4946	.4948	.4949	.4951	.4952
2.6	.4953	.4955	.4956	.4957	.4959	.4960	.4961	.4962	.4963	.4964
2.7	.4965	.4966	.4967	.4968	.4969	.4970	.4971	.4972	.4973	.4974
2.8	.4974	.4975	.4976	.4977	.4977	.4978	.4979	.4979	.4980	.4981
2.9	.4981	.4982	.4982	.4983	.4984	.4984	.4985	.4985	.4986	.4986
3.0	.4987	.4987	.4987	.4988	.4988	.4989	.4989	.4989	.4990	.4990
3.1	.4990	.4991	.4991	.4991	.4992	.4992	.4992	.4992	.4993	.4993
3.2	.4993	.4993	.4994	.4994	.4994	.4994	.4994	.4995	.4995	.4995
3.3	.4995	.4995	.4995	.4996	.4996	.4996	.4996	.4996	.4996	.4997
3.4	.4997	.4997	.4997	.4997	.4997	.4997	.4997	.4997	.4997	.4998
3.5	.4998	.4998	.4998	.4998	.4998	.4998	.4998	.4998	.4998	.4998
3.6	.4998	.4998	.4998	.4999	.4999	.4999	.4999	.4999	.4999	.4999

For values of z greater than or equal to 3.70, use 0.4999 to approximate the shaded area under the standard normal curve.

TABLE 1 Random Numbers

92630	78240	19267	95457	53497	23894	37708	79862	76471	66418
79445	78735	71549	44843	26404	67318	00701	34986	66751	99723
59654	71966	27386	50004	05358	94031	29281	18544	52429	06080
31524	49587	76612	39789	13537	48086	59483	60680	84675	53014
06348	76938	90379	51392	55887	71015	09209	79157	24440	30244
28703	51709	94456	48396	73780	06436	86641	69239	57662	80181
68108	89266	94730	95761	75023	48464	65544	96583	18911	16391
99938	90704	93621	66330	33393	95261	95349	51769	91616	33238
91543	73196	34449	63513	83834	99411	58826	40456	69268	48562
42103	02781	73920	56297	72678	12249	25270	36678	21313	75767
17138	27584	25296	28387	51350	61664	37893	05363	44143	42677
28297	14280	54524	21618	95320	38174	60579	08089	94999	78460
09331	56712	51333	06289	75345	08811	82711	57392	25252	30333
31295	04204	93712	51287	05754	79396	87399	51773	33075	97061
36146	15560	27592	42089	99281	59640	15221	96079	09961	05371
29553	18432	13630	05529	02791	81017	49027	79031	50912	09399
23501	22642	63081	08191	89420	67800	55137	54707	32945	64522
57888	85846	67967	07835	11314	01545	48535	17142	08552	67457
55336	71264	88472	04334	63919	36394	11196	92470	70543	29776
10087	10072	55980	64688	68239	20461	89381	93809	00796	95945
34101	81277	66090	88872	37818	72142	67140	50785	21380	16703
53362	44940	60430	22834	14130	96593	23298	56203	92671	15925
82975	66158	84731	19436	55790	69229	28661	13675	99318	76873
54827	84673	22898	08094	14326	87038	42892	21127	30712	48489
25464	59098	27436	89421	80754	89924	19097	67737	80368	08795

TABLE 1 *continued*

67609	60214	41475	84950	40133	02546	09570	45682	50165	15609
44921	70924	61295	51137	47596	86735	35561	76649	18217	63446
33170	30972	98130	95828	49786	13301	36081	80761	33985	68621
84687	85445	06208	17654	51333	02878	35010	67578	61574	20749
71886	56450	36567	09395	96951	35507	17555	35212	69106	01679
00475	02224	74722	14721	40215	21351	08596	45625	83981	63748
25993	38881	68361	59560	41274	69742	40703	37993	03435	18873
92882	53178	99195	93803	56985	53089	15305	50522	55900	43026
25138	26810	07093	15677	60688	04410	24505	37890	67186	62829
84631	71882	12991	83028	82484	90339	91950	74579	03539	90122
34003	92326	12793	61453	48121	74271	28363	66561	75220	35908
53775	45749	05734	86169	42762	70175	97310	73894	88606	19994
59316	97885	72807	54966	60859	11932	35265	71601	55577	67715
20479	66557	50705	26999	09854	52591	14063	30214	19890	19292
86180	84931	25455	26044	02227	52015	21820	50599	51671	65411
21451	68001	72710	40261	61281	13172	63819	48970	51732	54113
98062	68375	80089	24135	72355	95428	11808	29740	81644	86610
01788	64429	14430	94575	75153	94576	61393	96192	03227	32258
62465	04841	43272	68702	01274	05437	22953	18946	99053	41690
94324	31089	84159	92933	99989	89500	91586	02802	69471	68274
05797	43984	21575	09908	70221	19791	51578	36432	33494	79888
10395	14289	52185	09721	25789	38562	54794	04897	59012	89251
35177	56986	25549	59730	64718	52630	31100	62384	49483	11409
25633	89619	75882	98256	02126	72099	57183	55887	09320	72363
16464	48280	94254	45777	45150	68865	11382	11782	22695	41988

Source: Reprinted from *A Million Random Digits with 100,000 Normal Deviates* by the Rand Corporation (New York: The Free Press, 1955). Copyright 1955 by the Rand Corporation. Used by permission.

TABLE 2 Binomial Coefficients $C_{n,r}$

$n \setminus r$	0	1	2	3	4	5	6	7	8	9	10
1	1	1									
2	1	2	1								
3	1	3	3	1							
4	1	4	6	4	1						
5	1	5	10	10	5	1					
6	1	6	15	20	15	6	1				
7	1	7	21	35	35	21	7	1			
8	1	8	28	56	70	56	28	8	1		
9	1	9	36	84	126	126	84	36	9	1	
10	1	10	45	120	210	252	210	120	45	10	1
11	1	11	55	165	330	462	462	330	165	55	11
12	1	12	66	220	495	792	924	792	495	220	66
13	1	13	78	286	715	1,287	1,716	1,716	1,287	715	286
14	1	14	91	364	1,001	2,002	3,003	3,432	3,003	2,002	1,001
15	1	15	105	455	1,365	3,003	5,005	6,435	6,435	5,005	3,003
16	1	16	120	560	1,820	4,368	8,008	11,440	12,870	11,440	8,008
17	1	17	136	680	2,380	6,188	12,376	19,448	24,310	24,310	19,448
18	1	18	153	816	3,060	8,568	18,564	31,824	43,758	48,620	43,758
19	1	19	171	969	3,876	11,628	27,132	50,388	75,582	92,378	92,378
20	1	20	190	1,140	3,845	15,504	38,760	77,520	125,970	167,960	184,756

TABLE 3 Binomial Probability Distribution $C_{n,r} p^r q^{n-r}$

This table shows the probability of r successes in n independent trials, each with probability of success p .																					
n	r	p																			
		.01	.05	.10	.15	.20	.25	.30	.35	.40	.45	.50	.55	.60	.65	.70	.75	.80	.85	.90	.95
2	0	.980	.902	.810	.723	.640	.563	.490	.423	.360	.303	.250	.203	.160	.123	.090	.063	.040	.023	.010	.002
	1	.020	.095	.180	.255	.320	.375	.420	.455	.480	.495	.500	.495	.480	.455	.420	.375	.320	.255	.180	.095
	2	.000	.002	.010	.023	.040	.063	.090	.123	.160	.203	.250	.303	.360	.423	.490	.563	.640	.723	.810	.902
3	0	.970	.857	.729	.614	.512	.422	.343	.275	.216	.166	.125	.091	.064	.043	.027	.016	.008	.003	.001	.000
	1	.029	.135	.243	.325	.384	.422	.441	.444	.432	.408	.375	.334	.288	.239	.189	.141	.096	.057	.027	.007
	2	.000	.007	.028	.057	.096	.141	.189	.239	.288	.334	.375	.408	.432	.444	.441	.422	.384	.325	.243	.135
	3	.000	.000	.001	.003	.008	.016	.027	.043	.064	.091	.125	.166	.216	.275	.343	.422	.512	.614	.729	.857
4	0	.961	.815	.656	.522	.410	.316	.240	.179	.130	.092	.062	.041	.026	.015	.008	.004	.002	.001	.000	.000
	1	.039	.171	.292	.368	.410	.422	.412	.384	.346	.300	.250	.200	.154	.112	.076	.047	.026	.011	.004	.000
	2	.001	.014	.049	.098	.154	.211	.265	.311	.346	.368	.375	.368	.346	.311	.265	.211	.154	.098	.049	.014
	3	.000	.000	.004	.011	.026	.047	.076	.112	.154	.200	.250	.300	.346	.384	.412	.422	.410	.368	.292	.171
	4	.000	.000	.000	.001	.002	.004	.008	.015	.026	.041	.062	.092	.130	.179	.240	.316	.410	.522	.656	.815
5	0	.951	.774	.590	.444	.328	.237	.168	.116	.078	.050	.031	.019	.010	.005	.002	.001	.000	.000	.000	.000
	1	.048	.204	.328	.392	.410	.396	.360	.312	.259	.206	.156	.113	.077	.049	.028	.015	.006	.002	.000	.000
	2	.001	.021	.073	.138	.205	.264	.309	.336	.346	.337	.312	.276	.230	.181	.132	.088	.051	.024	.008	.001
	3	.000	.001	.008	.024	.051	.088	.132	.181	.230	.276	.312	.337	.346	.336	.309	.264	.205	.138	.073	.021
	4	.000	.000	.000	.002	.006	.015	.028	.049	.077	.113	.156	.206	.259	.312	.360	.396	.410	.392	.328	.204
	5	.000	.000	.000	.000	.000	.001	.002	.005	.010	.019	.031	.050	.078	.116	.168	.237	.328	.444	.590	.774
6	0	.941	.735	.531	.377	.262	.178	.118	.075	.047	.028	.016	.008	.004	.002	.001	.000	.000	.000	.000	.000
	1	.057	.232	.354	.399	.393	.356	.303	.244	.187	.136	.094	.061	.037	.020	.010	.004	.002	.000	.000	.000
	2	.001	.031	.098	.176	.246	.297	.324	.328	.311	.278	.234	.186	.138	.095	.060	.033	.015	.006	.001	.000
	3	.000	.002	.015	.042	.082	.132	.185	.236	.276	.303	.312	.303	.276	.236	.185	.132	.082	.042	.015	.002
	4	.000	.000	.001	.006	.015	.033	.060	.095	.138	.186	.234	.278	.311	.328	.324	.297	.246	.176	.098	.031
	5	.000	.000	.000	.000	.002	.004	.010	.020	.037	.061	.094	.136	.187	.244	.303	.356	.393	.399	.354	.232
	6	.000	.000	.000	.000	.000	.000	.001	.002	.004	.008	.016	.028	.047	.075	.118	.178	.262	.377	.531	.735
7	0	.932	.698	.478	.321	.210	.133	.082	.049	.028	.015	.008	.004	.002	.001	.000	.000	.000	.000	.000	.000
	1	.066	.257	.372	.396	.367	.311	.247	.185	.131	.087	.055	.032	.017	.008	.004	.001	.000	.000	.000	.000
	2	.002	.041	.124	.210	.275	.311	.318	.299	.261	.214	.164	.117	.077	.047	.025	.012	.004	.001	.000	.000
	3	.000	.004	.023	.062	.115	.173	.227	.268	.290	.292	.273	.239	.194	.144	.097	.058	.029	.011	.003	.000
	4	.000	.000	.003	.011	.029	.058	.097	.144	.194	.239	.273	.292	.290	.268	.227	.173	.115	.062	.023	.004
	5	.000	.000	.000	.001	.004	.012	.025	.047	.077	.117	.164	.214	.261	.299	.318	.311	.275	.210	.124	.041
	6	.000	.000	.000	.000	.000	.001	.004	.008	.017	.032	.055	.087	.131	.185	.247	.311	.367	.396	.372	.257
	7	.000	.000	.000	.000	.000	.000	.000	.001	.002	.004	.008	.015	.028	.049	.082	.133	.210	.321	.478	.698

TABLE 3 *continued*

		p																			
n	r	.01	.05	.10	.15	.20	.25	.30	.35	.40	.45	.50	.55	.60	.65	.70	.75	.80	.85	.90	.95
8	0	.923	.663	.430	.272	.168	.100	.058	.032	.017	.008	.004	.002	.001	.000	.000	.000	.000	.000	.000	.000
	1	.075	.279	.383	.385	.336	.267	.198	.137	.090	.055	.031	.016	.008	.003	.001	.000	.000	.000	.000	.000
	2	.003	.051	.149	.238	.294	.311	.296	.259	.209	.157	.109	.070	.041	.022	.010	.004	.001	.000	.000	.000
	3	.000	.005	.033	.084	.147	.208	.254	.279	.279	.257	.219	.172	.124	.081	.047	.023	.009	.003	.000	.000
	4	.000	.000	.005	.018	.046	.087	.136	.188	.232	.263	.273	.263	.232	.188	.136	.087	.046	.018	.005	.000
	5	.000	.000	.000	.003	.009	.023	.047	.081	.124	.172	.219	.257	.279	.279	.254	.208	.147	.084	.033	.005
	6	.000	.000	.000	.000	.001	.004	.010	.022	.041	.070	.109	.157	.209	.259	.296	.311	.294	.238	.149	.051
	7	.000	.000	.000	.000	.000	.000	.001	.003	.008	.016	.031	.055	.090	.137	.198	.267	.336	.385	.383	.279
8	.000	.000	.000	.000	.000	.000	.000	.000	.001	.002	.004	.008	.017	.032	.058	.100	.168	.272	.430	.663	
9	0	.914	.630	.387	.232	.134	.075	.040	.021	.010	.005	.002	.001	.000	.000	.000	.000	.000	.000	.000	.000
	1	.083	.299	.387	.368	.302	.225	.156	.100	.060	.034	.018	.008	.004	.001	.000	.000	.000	.000	.000	.000
	2	.003	.063	.172	.260	.302	.300	.267	.216	.161	.111	.070	.041	.021	.010	.004	.001	.000	.000	.000	.000
	3	.000	.008	.045	.107	.176	.234	.267	.272	.251	.212	.164	.116	.074	.042	.021	.009	.003	.001	.000	.000
	4	.000	.001	.007	.028	.066	.117	.172	.219	.251	.260	.246	.213	.167	.118	.074	.039	.017	.005	.001	.000
	5	.000	.000	.001	.005	.017	.039	.074	.118	.167	.213	.246	.260	.251	.219	.172	.117	.066	.028	.007	.001
	6	.000	.000	.000	.001	.003	.009	.021	.042	.074	.116	.164	.212	.251	.272	.267	.234	.176	.107	.045	.008
	7	.000	.000	.000	.000	.000	.001	.004	.010	.021	.041	.070	.111	.161	.216	.267	.300	.302	.260	.172	.063
	8	.000	.000	.000	.000	.000	.000	.000	.001	.004	.008	.018	.034	.060	.100	.156	.225	.302	.368	.387	.299
9	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.002	.005	.010	.021	.040	.075	.134	.232	.387	.630	
10	0	.904	.599	.349	.197	.107	.056	.028	.014	.006	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.091	.315	.387	.347	.268	.188	.121	.072	.040	.021	.010	.004	.002	.000	.000	.000	.000	.000	.000	.000
	2	.004	.075	.194	.276	.302	.282	.233	.176	.121	.076	.044	.023	.011	.004	.001	.000	.000	.000	.000	.000
	3	.000	.010	.057	.130	.201	.250	.267	.252	.215	.166	.117	.075	.042	.021	.009	.003	.001	.000	.000	.000
	4	.000	.001	.011	.040	.088	.146	.200	.238	.251	.238	.205	.160	.111	.069	.037	.016	.006	.001	.000	.000
	5	.000	.000	.001	.008	.026	.058	.103	.154	.201	.234	.246	.234	.201	.154	.103	.058	.026	.008	.001	.000
	6	.000	.000	.000	.001	.006	.016	.037	.069	.111	.160	.205	.238	.251	.238	.200	.146	.088	.040	.011	.001
	7	.000	.000	.000	.000	.001	.003	.009	.021	.042	.075	.117	.166	.215	.252	.267	.250	.201	.130	.057	.010
	8	.000	.000	.000	.000	.000	.000	.001	.004	.011	.023	.044	.076	.121	.176	.233	.282	.302	.276	.194	.075
	9	.000	.000	.000	.000	.000	.000	.000	.000	.002	.004	.010	.021	.040	.072	.121	.188	.268	.347	.387	.315
10	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.006	.014	.028	.056	.107	.197	.349	.599	

TABLE 3 *continued*

<i>n</i>	<i>r</i>	<i>p</i>																			
		.01	.05	.10	.15	.20	.25	.30	.35	.40	.45	.50	.55	.60	.65	.70	.75	.80	.85	.90	.95
11	0	.895	.569	.314	.167	.086	.042	.020	.009	.004	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.099	.329	.384	.325	.236	.155	.093	.052	.027	.013	.005	.002	.001	.000	.000	.000	.000	.000	.000	.000
	2	.005	.087	.213	.287	.295	.258	.200	.140	.089	.051	.027	.013	.005	.002	.001	.000	.000	.000	.000	.000
	3	.000	.014	.071	.152	.221	.258	.257	.225	.177	.126	.081	.046	.023	.010	.004	.001	.000	.000	.000	.000
	4	.000	.001	.016	.054	.111	.172	.220	.243	.236	.206	.161	.113	.070	.038	.017	.006	.002	.000	.000	.000
	5	.000	.000	.002	.013	.039	.080	.132	.183	.221	.236	.226	.193	.147	.099	.057	.027	.010	.002	.000	.000
	6	.000	.000	.000	.002	.010	.027	.057	.099	.147	.193	.226	.236	.221	.183	.132	.080	.039	.013	.002	.000
	7	.000	.000	.000	.000	.002	.006	.017	.038	.070	.113	.161	.206	.236	.243	.220	.172	.111	.054	.016	.001
	8	.000	.000	.000	.000	.000	.001	.004	.010	.023	.046	.081	.126	.177	.225	.257	.258	.221	.152	.071	.014
	9	.000	.000	.000	.000	.000	.000	.001	.002	.005	.013	.027	.051	.089	.140	.200	.258	.295	.287	.213	.087
	10	.000	.000	.000	.000	.000	.000	.000	.000	.001	.002	.005	.013	.027	.052	.093	.155	.236	.325	.684	.329
11	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.004	.009	.020	.042	.086	.167	.314	.569	
12	0	.886	.540	.282	.142	.069	.032	.014	.006	.002	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.107	.341	.377	.301	.206	.127	.071	.037	.017	.008	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000
	2	.006	.099	.230	.292	.283	.232	.168	.109	.064	.034	.016	.007	.002	.001	.000	.000	.000	.000	.000	.000
	3	.000	.017	.085	.172	.236	.258	.240	.195	.142	.092	.054	.028	.012	.005	.001	.000	.000	.000	.000	.000
	4	.000	.002	.021	.068	.133	.194	.231	.237	.213	.170	.121	.076	.042	.020	.008	.002	.001	.000	.000	.000
	5	.000	.000	.004	.019	.053	.103	.158	.204	.227	.223	.193	.149	.101	.059	.029	.011	.003	.001	.000	.000
	6	.000	.000	.000	.004	.016	.040	.079	.128	.177	.212	.226	.212	.177	.128	.079	.040	.016	.004	.000	.000
	7	.000	.000	.000	.001	.003	.011	.029	.059	.101	.149	.193	.223	.227	.204	.158	.103	.053	.019	.004	.000
	8	.000	.000	.000	.000	.001	.002	.008	.020	.042	.076	.121	.170	.216	.237	.231	.194	.133	.068	.021	.002
	9	.000	.000	.000	.000	.000	.000	.001	.005	.012	.028	.054	.092	.142	.195	.240	.258	.236	.172	.085	.017
	10	.000	.000	.000	.000	.000	.000	.000	.001	.002	.007	.016	.034	.064	.109	.168	.232	.283	.292	.230	.099
	11	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.008	.017	.037	.071	.127	.206	.301	.377	.341
	12	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.002	.006	.014	.032	.069	.142	.282	.540
15	0	.860	.463	.206	.087	.035	.013	.005	.002	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.130	.366	.343	.231	.132	.067	.031	.013	.005	.002	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	2	.009	.135	.267	.286	.231	.156	.092	.048	.022	.009	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000
	3	.000	.031	.129	.218	.250	.225	.170	.111	.063	.032	.014	.005	.002	.000	.000	.000	.000	.000	.000	.000
	4	.005	.005	.043	.116	.188	.225	.219	.179	.127	.078	.042	.019	.007	.002	.001	.000	.000	.000	.000	.000
	5	.000	.001	.010	.045	.103	.165	.206	.212	.186	.140	.092	.051	.024	.010	.003	.001	.000	.000	.000	.000
	6	.000	.000	.002	.013	.043	.092	.147	.191	.207	.191	.153	.105	.061	.030	.012	.003	.001	.000	.000	.000
	7	.000	.000	.000	.003	.014	.039	.081	.132	.177	.201	.196	.165	.118	.071	.035	.013	.003	.001	.000	.000
	8	.000	.000	.000	.001	.003	.013	.035	.071	.118	.165	.196	.201	.177	.132	.081	.039	.014	.003	.000	.000
	9	.000	.000	.000	.000	.001	.003	.012	.030	.061	.105	.153	.191	.207	.191	.147	.092	.043	.013	.002	.000
	10	.000	.000	.000	.000	.000	.001	.003	.010	.024	.051	.092	.140	.186	.212	.206	.165	.103	.045	.010	.001
	11	.000	.000	.000	.000	.000	.000	.001	.002	.007	.019	.042	.078	.127	.179	.219	.225	.188	.116	.043	.005
	12	.000	.000	.000	.000	.000	.000	.000	.000	.002	.005	.014	.032	.063	.111	.170	.225	.250	.218	.129	.031
	13	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.009	.022	.048	.092	.156	.231	.586	.267	.135
	14	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.002	.005	.013	.031	.067	.132	.231	.343	.366
	15	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.002	.005	.013	.035	.087	.206	.463

TABLE 3 *continued*

		p																			
n	r	.01	.05	.10	.15	.20	.25	.30	.35	.40	.45	.50	.55	.60	.65	.70	.75	.80	.85	.90	.95
16	0	.851	.440	.185	.074	.028	.010	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.138	.371	.329	.210	.113	.053	.023	.009	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	2	.010	.146	.275	.277	.211	.134	.073	.035	.015	.006	.002	.001	.000	.000	.000	.000	.000	.000	.000	.000
	3	.000	.036	.142	.229	.246	.208	.146	.089	.047	.022	.009	.003	.001	.000	.000	.000	.000	.000	.000	.000
	4	.000	.006	.051	.131	.200	.225	.204	.155	.101	.057	.028	.011	.004	.001	.000	.000	.000	.000	.000	.000
	5	.000	.001	.014	.056	.120	.180	.210	.201	.162	.112	.067	.034	.014	.005	.001	.000	.000	.000	.000	.000
	6	.000	.000	.003	.018	.055	.110	.165	.198	.198	.168	.122	.075	.039	.017	.006	.001	.000	.000	.000	.000
	7	.000	.000	.000	.005	.020	.052	.101	.152	.189	.197	.175	.132	.084	.044	.019	.006	.001	.000	.000	.000
	8	.000	.000	.000	.001	.006	.020	.049	.092	.142	.181	.196	.181	.142	.092	.049	.020	.006	.001	.000	.000
	9	.000	.000	.000	.000	.001	.006	.019	.044	.084	.132	.175	.197	.189	.152	.101	.052	.020	.005	.000	.000
	10	.000	.000	.000	.000	.000	.001	.006	.017	.039	.075	.122	.168	.198	.198	.165	.110	.055	.018	.003	.000
	11	.000	.000	.000	.000	.000	.000	.001	.005	.014	.034	.067	.112	.162	.201	.210	.180	.120	.056	.014	.001
	12	.000	.000	.000	.000	.000	.000	.000	.001	.004	.011	.028	.057	.101	.155	.204	.225	.200	.131	.051	.006
	13	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.009	.022	.047	.089	.146	.208	.246	.229	.142	.036
	14	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.002	.006	.015	.035	.073	.134	.211	.277	.275	.146
	15	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.009	.023	.053	.113	.210	.329	.371
16	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.010	.028	.074	.185	.440	
20	0	.818	.358	.122	.039	.012	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	1	.165	.377	.270	.137	.058	.021	.007	.002	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	2	.016	.189	.285	.229	.137	.067	.028	.010	.003	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	3	.001	.060	.190	.243	.205	.134	.072	.032	.012	.004	.001	.000	.000	.000	.000	.000	.000	.000	.000	.000
	4	.000	.013	.090	.182	.218	.190	.130	.074	.035	.014	.005	.001	.000	.000	.000	.000	.000	.000	.000	.000
	5	.000	.002	.032	.103	.175	.202	.179	.127	.075	.036	.015	.005	.001	.000	.000	.000	.000	.000	.000	.000
	6	.000	.000	.009	.045	.109	.169	.192	.171	.124	.075	.036	.015	.005	.001	.000	.000	.000	.000	.000	.000
	7	.000	.000	.002	.016	.055	.112	.164	.184	.166	.122	.074	.037	.015	.005	.001	.000	.000	.000	.000	.000
	8	.000	.000	.000	.005	.022	.061	.114	.161	.180	.162	.120	.073	.035	.014	.004	.001	.000	.000	.000	.000
	9	.000	.000	.000	.001	.007	.027	.065	.116	.160	.177	.160	.119	.071	.034	.012	.003	.000	.000	.000	.000
	10	.000	.000	.000	.000	.002	.010	.031	.069	.117	.159	.176	.159	.117	.069	.031	.010	.002	.000	.000	.000
	11	.000	.000	.000	.000	.000	.003	.012	.034	.071	.119	.160	.177	.160	.116	.065	.027	.007	.001	.000	.000
	12	.000	.000	.000	.000	.000	.001	.004	.014	.035	.073	.120	.162	.180	.161	.114	.061	.022	.005	.000	.000
	13	.000	.000	.000	.000	.000	.000	.001	.005	.015	.037	.074	.122	.166	.184	.164	.112	.055	.016	.002	.000
	14	.000	.000	.000	.000	.000	.000	.000	.001	.005	.015	.037	.075	.124	.171	.192	.169	.109	.045	.009	.000
	15	.000	.000	.000	.000	.000	.000	.000	.000	.001	.005	.015	.036	.075	.127	.179	.202	.175	.103	.032	.002
	16	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.005	.014	.035	.074	.130	.190	.218	.182	.090	.013
	17	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.004	.012	.032	.072	.134	.205	.243	.190	.060
	18	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.010	.028	.067	.137	.229	.285	.189
	19	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.002	.007	.021	.058	.137	.270	.377
20	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.001	.003	.012	.039	.122	.358	

TABLE 4 Poisson Probability Distribution

For a given value of λ , entry indicates the probability of obtaining a specified value of r .										
λ										
r	.1	.2	.3	.4	.5	.6	.7	.8	.9	1.0
0	.9048	.8187	.7408	.6703	.6065	.5488	.4966	.4493	.4066	.3679
1	.0905	.1637	.2222	.2681	.3033	.3293	.3476	.3595	.3659	.3679
2	.0045	.0164	.0333	.0536	.0758	.0988	.1217	.1438	.1647	.1839
3	.0002	.0011	.0033	.0072	.0126	.0198	.0284	.0383	.0494	.0613
4	.0000	.0001	.0003	.0007	.0016	.0030	.0050	.0077	.0111	.0153
5	.0000	.0000	.0000	.0001	.0002	.0004	.0007	.0012	.0020	.0031
6	.0000	.0000	.0000	.0000	.0000	.0000	.0001	.0002	.0003	.0005
7	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0001

λ										
r	1.1	1.2	1.3	1.4	1.5	1.6	1.7	1.8	1.9	2.0
0	.3329	.3012	.2725	.2466	.2231	.2019	.1827	.1653	.1496	.1353
1	.3662	.3614	.3543	.3452	.3347	.3230	.3106	.2975	.2842	.2707
2	.2014	.2169	.2303	.2417	.2510	.2584	.2640	.2678	.2700	.2707
3	.0738	.0867	.0998	.1128	.1255	.1378	.1496	.1607	.1710	.1804
4	.0203	.0260	.0324	.0395	.0471	.0551	.0636	.0723	.0812	.0902
5	.0045	.0062	.0084	.0111	.0141	.0176	.0216	.0260	.0309	.0361
6	.0008	.0012	.0018	.0026	.0035	.0047	.0061	.0078	.0098	.0120
7	.0001	.0002	.0003	.0005	.0008	.0011	.0015	.0020	.0027	.0034
8	.0000	.0000	.0001	.0001	.0001	.0002	.0003	.0005	.0006	.0009
9	.0000	.0000	.0000	.0000	.0000	.0000	.0001	.0001	.0001	.0002

λ										
r	2.1	2.2	2.3	2.4	2.5	2.6	2.7	2.8	2.9	3.0
0	.1225	.1108	.1003	.0907	.0821	.0743	.0672	.0608	.0550	.0498
1	.2572	.2438	.2306	.2177	.2052	.1931	.1815	.1703	.1596	.1494
2	.2700	.2681	.2652	.2613	.2565	.2510	.2450	.2384	.2314	.2240
3	.1890	.1966	.2033	.2090	.2138	.2176	.2205	.2225	.2237	.2240
4	.0992	.1082	.1169	.1254	.1336	.1414	.1488	.1557	.1622	.1680
5	.0417	.0476	.0538	.0602	.0668	.0735	.0804	.0872	.0940	.1008
6	.0146	.0174	.0206	.0241	.0278	.0319	.0362	.0407	.0455	.0504
7	.0044	.0055	.0068	.0083	.0099	.0118	.0139	.0163	.0188	.0216
8	.0011	.0015	.0019	.0025	.0031	.0038	.0047	.0057	.0068	.0081
9	.0003	.0004	.0005	.0007	.0009	.0011	.0014	.0018	.0022	.0027
10	.0001	.0001	.0001	.0002	.0002	.0003	.0004	.0005	.0006	.0008
11	.0000	.0000	.0000	.0000	.0000	.0001	.0001	.0001	.0002	.0002
12	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0001