

Nonparametric **Statistics**

A Step-by-Step Approach

Second Edition

Gregory W. Corder • Dale I. Foreman



WILEY

NONPARAMETRIC STATISTICS

SECOND EDITION

NONPARAMETRIC STATISTICS
A Step-by-Step Approach

GREGORY W. CORDER

DALE I. FOREMAN

WILEY

Copyright © 2014 by John Wiley & Sons, Inc. All rights reserved.

Published by John Wiley & Sons, Inc., Hoboken, New Jersey.

Published simultaneously in Canada.

No part of this publication may be reproduced, stored in a retrieval system, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, scanning, or otherwise, except as permitted under Section 107 or 108 of the 1976 United States Copyright Act, without either the prior written permission of the Publisher, or authorization through payment of the appropriate per-copy fee to the Copyright Clearance Center, Inc., 222 Rosewood Drive, Danvers, MA 01923, (978) 750-8400, fax (978) 750-4470, or on the web at www.copyright.com. Requests to the Publisher for permission should be addressed to the Permissions Department, John Wiley & Sons, Inc., 111 River Street, Hoboken, NJ 07030, (201) 748-6011, fax (201) 748-6008, or online at <http://www.wiley.com/go/permissions>.

Limit of Liability/Disclaimer of Warranty: While the publisher and author have used their best efforts in preparing this book, they make no representations or warranties with respect to the accuracy or completeness of the contents of this book and specifically disclaim any implied warranties of merchantability or fitness for a particular purpose. No warranty may be created or extended by sales representatives or written sales materials. The advice and strategies contained herein may not be suitable for your situation. You should consult with a professional where appropriate. Neither the publisher nor author shall be liable for any loss of profit or any other commercial damages, including but not limited to special, incidental, consequential, or other damages.

For general information on our other products and services or for technical support, please contact our Customer Care Department within the United States at (800) 762-2974, outside the United States at (317) 572-3993 or fax (317) 572-4002.

Wiley also publishes its books in a variety of electronic formats. Some content that appears in print may not be available in electronic formats. For more information about Wiley products, visit our web site at www.wiley.com.

Library of Congress Cataloging-in-Publication Data:

Corder, Gregory W., 1972– author.

Nonparametric statistics : a step-by-step approach / Gregory W. Corder. Dale I. Foreman – Second edition.

pages cm

Includes bibliographical references and index.

ISBN 978-1-118-84031-3 (hardback)

1. Nonparametric statistics. I. Foreman, Dale I., author II. Title.

QA278.8.C67 2014

519.5'4–dc23

2013042040

Printed in the United States of America

10 9 8 7 6 5 4 3 2 1

CONTENTS

<i>PREFACE</i>	ix
<i>LIST OF VARIABLES</i>	xiii
CHAPTER 1 <i>NONPARAMETRIC STATISTICS: AN INTRODUCTION</i>	1
1.1 Objectives	1
1.2 Introduction	1
1.3 The Nonparametric Statistical Procedures Presented in This Book	3
1.3.1 State the Null and Research Hypotheses	4
1.3.2 Set the Level of Risk (or the Level of Significance) Associated with the Null Hypothesis	4
1.3.3 Choose the Appropriate Test Statistic	5
1.3.4 Compute the Test Statistic	5
1.3.5 Determine the Value Needed for Rejection of the Null Hypothesis Using the Appropriate Table of Critical Values for the Particular Statistic	5
1.3.6 Compare the Obtained Value with the Critical Value	6
1.3.7 Interpret the Results	6
1.3.8 Reporting the Results	6
1.4 Ranking Data	6
1.5 Ranking Data with Tied Values	7
1.6 Counts of Observations	8
1.7 Summary	9
1.8 Practice Questions	9
1.9 Solutions to Practice Questions	10
CHAPTER 2 <i>TESTING DATA FOR NORMALITY</i>	13
2.1 Objectives	13
2.2 Introduction	13
2.3 Describing Data and the Normal Distribution	14
2.4 Computing and Testing Kurtosis and Skewness for Sample Normality	17
2.4.1 Sample Problem for Examining Kurtosis	19
2.4.2 Sample Problem for Examining Skewness	22
2.4.3 Examining Skewness and Kurtosis for Normality Using SPSS	24
2.5 Computing the Kolmogorov–Smirnov One-Sample Test	27
2.5.1 Sample Kolmogorov–Smirnov One-Sample Test	29
2.5.2 Performing the Kolmogorov–Smirnov One-Sample Test Using SPSS	34
2.6 Summary	37
2.7 Practice Questions	37
2.8 Solutions to Practice Questions	38

CHAPTER 3	<i>COMPARING TWO RELATED SAMPLES: THE WILCOXON SIGNED RANK AND THE SIGN TEST</i>	39
3.1	Objectives	39
3.2	Introduction	39
3.3	Computing the Wilcoxon Signed Rank Test Statistic	40
3.3.1	Sample Wilcoxon Signed Rank Test (Small Data Samples)	41
3.3.2	Confidence Interval for the Wilcoxon Signed Rank Test	43
3.3.3	Sample Wilcoxon Signed Rank Test (Large Data Samples)	45
3.4	Computing the Sign Test	49
3.4.1	Sample Sign Test (Small Data Samples)	50
3.4.2	Sample Sign Test (Large Data Samples)	53
3.5	Performing the Wilcoxon Signed Rank Test and the Sign Test Using SPSS	57
3.5.1	Define Your Variables	57
3.5.2	Type in Your Values	57
3.5.3	Analyze Your Data	58
3.5.4	Interpret the Results from the SPSS Output Window	58
3.6	Statistical Power	60
3.7	Examples from the Literature	61
3.8	Summary	61
3.9	Practice Questions	62
3.10	Solutions to Practice Questions	65
CHAPTER 4	<i>COMPARING TWO UNRELATED SAMPLES: THE MANN–WHITNEY U-TEST AND THE KOLMOGOROV–SMIRNOV TWO-SAMPLE TEST</i>	69
4.1	Objectives	69
4.2	Introduction	69
4.3	Computing the Mann–Whitney <i>U</i> -Test Statistic	70
4.3.1	Sample Mann–Whitney <i>U</i> -Test (Small Data Samples)	71
4.3.2	Confidence Interval for the Difference between Two Location Parameters	74
4.3.3	Sample Mann–Whitney <i>U</i> -Test (Large Data Samples)	75
4.4	Computing the Kolmogorov–Smirnov Two-Sample Test Statistic	80
4.4.1	Sample Kolmogorov–Smirnov Two-Sample Test	81
4.5	Performing the Mann–Whitney <i>U</i> -Test and the Kolmogorov–Smirnov Two-Sample Test Using SPSS	84
4.5.1	Define Your Variables	84
4.5.2	Type in Your Values	85
4.5.3	Analyze Your Data	86
4.5.4	Interpret the Results from the SPSS Output Window	86
4.6	Examples from the Literature	88
4.7	Summary	89
4.8	Practice Questions	90
4.9	Solutions to Practice Questions	92
CHAPTER 5	<i>COMPARING MORE THAN TWO RELATED SAMPLES: THE FRIEDMAN TEST</i>	97
5.1	Objectives	97
5.2	Introduction	97

5.3	Computing the Friedman Test Statistic	98
5.3.1	Sample Friedman's Test (Small Data Samples without Ties)	99
5.3.2	Sample Friedman's Test (Small Data Samples with Ties)	101
5.3.3	Performing the Friedman Test Using SPSS	105
5.3.4	Sample Friedman's Test (Large Data Samples without Ties)	108
5.4	Examples from the Literature	111
5.5	Summary	112
5.6	Practice Questions	113
5.7	Solutions to Practice Questions	114
CHAPTER 6	COMPARING MORE THAN TWO UNRELATED SAMPLES: THE KRUSKAL–WALLIS H-TEST	117
6.1	Objectives	117
6.2	Introduction	117
6.3	Computing the Kruskal–Wallis <i>H</i> -Test Statistic	118
6.3.1	Sample Kruskal–Wallis <i>H</i> -Test (Small Data Samples)	119
6.3.2	Performing the Kruskal–Wallis <i>H</i> -Test Using SPSS	124
6.3.3	Sample Kruskal–Wallis <i>H</i> -Test (Large Data Samples)	128
6.4	Examples from the Literature	134
6.5	Summary	134
6.6	Practice Questions	135
6.7	Solutions to Practice Questions	136
CHAPTER 7	COMPARING VARIABLES OF ORDINAL OR DICHOTOMOUS SCALES: SPEARMAN RANK-ORDER, POINT-BISERIAL, AND BISERIAL CORRELATIONS	139
7.1	Objectives	139
7.2	Introduction	139
7.3	The Correlation Coefficient	140
7.4	Computing the Spearman Rank-Order Correlation Coefficient	140
7.4.1	Sample Spearman Rank-Order Correlation (Small Data Samples without Ties)	142
7.4.2	Sample Spearman Rank-Order Correlation (Small Data Samples with Ties)	145
7.4.3	Performing the Spearman Rank-Order Correlation Using SPSS	148
7.5	Computing the Point-Biserial and Biserial Correlation Coefficients	150
7.5.1	Correlation of a Dichotomous Variable and an Interval Scale Variable	150
7.5.2	Correlation of a Dichotomous Variable and a Rank-Order Variable	152
7.5.3	Sample Point-Biserial Correlation (Small Data Samples)	152
7.5.4	Performing the Point-Biserial Correlation Using SPSS	156
7.5.5	Sample Point-Biserial Correlation (Large Data Samples)	159
7.5.6	Sample Biserial Correlation (Small Data Samples)	163
7.5.7	Performing the Biserial Correlation Using SPSS	167
7.6	Examples from the Literature	167
7.7	Summary	167
7.8	Practice Questions	168
7.9	Solutions to Practice Questions	170

CHAPTER 8	<i>TESTS FOR NOMINAL SCALE DATA: CHI-SQUARE AND FISHER EXACT TESTS</i>	172
8.1	Objectives	172
8.2	Introduction	172
8.3	The χ^2 Goodness-of-Fit Test	172
8.3.1	Computing the χ^2 Goodness-of-Fit Test Statistic	173
8.3.2	Sample χ^2 Goodness-of-Fit Test (Category Frequencies Equal)	173
8.3.3	Sample χ^2 Goodness-of-Fit Test (Category Frequencies Not Equal)	176
8.3.4	Performing the χ^2 Goodness-of-Fit Test Using SPSS	180
8.4	The χ^2 Test for Independence	184
8.4.1	Computing the χ^2 Test for Independence	185
8.4.2	Sample χ^2 Test for Independence	186
8.4.3	Performing the χ^2 Test for Independence Using SPSS	190
8.5	The Fisher Exact Test	196
8.5.1	Computing the Fisher Exact Test for 2×2 Tables	197
8.5.2	Sample Fisher Exact Test	197
8.5.3	Performing the Fisher Exact Test Using SPSS	201
8.6	Examples from the Literature	202
8.7	Summary	203
8.8	Practice Questions	204
8.9	Solutions to Practice Questions	206
CHAPTER 9	<i>TEST FOR RANDOMNESS: THE RUNS TEST</i>	210
9.1	Objectives	210
9.2	Introduction	210
9.3	The Runs Test for Randomness	210
9.3.1	Sample Runs Test (Small Data Samples)	212
9.3.2	Performing the Runs Test Using SPSS	213
9.3.3	Sample Runs Test (Large Data Samples)	217
9.3.4	Sample Runs Test Referencing a Custom Value	219
9.3.5	Performing the Runs Test for a Custom Value Using SPSS	221
9.4	Examples from the Literature	225
9.5	Summary	225
9.6	Practice Questions	225
9.7	Solutions to Practice Questions	227
APPENDIX A	<i>SPSS AT A GLANCE</i>	229
APPENDIX B	<i>CRITICAL VALUE TABLES</i>	235
REFERENCES		261
INDEX		265

PREFACE

The social, behavioral, and health sciences have a need for the ability to use nonparametric statistics in research. Many studies in these areas involve data that are classified in the nominal or ordinal scale. At times, interval data from these fields lack parameters for classification as normal. Nonparametric statistical tests are useful tools for analyzing such data.

Purpose of This Book

This book is intended to provide a conceptual and procedural approach for nonparametric statistics. It is written so that someone who does not have an extensive mathematical background may work through the process necessary to conduct the given statistical tests presented. In addition, the outcome includes a discussion of the final decision for each statistical test. Each chapter takes the reader through an example from the beginning hypotheses, through the statistical calculations, to the final decision as compared with the hypothesis. The examples are then followed by a detailed, step-by-step analysis using the computer program SPSS®. Finally, research literature is identified which uses the respective nonparametric statistical tests.

Intended Audience

While not limited to such, this book is written for graduate and undergraduate students in social science programs. As stated earlier, it is targeted toward the student who does not have an especially strong mathematical background, but can be used effectively with a mixed group of students that includes students who have both strong and weak mathematical background.

Special Features of This Book

There are currently few books available that provide a practical and applied approach to teaching nonparametric statistics. Many books take a more theoretical approach to the instructional process that can leave students disconnected and frustrated, in need of supplementary material to give them the ability to apply the statistics taught.

It is our hope and expectation that this book provides students with a concrete approach to performing the nonparametric statistical procedures, along with their application and interpretation. We chose these particular nonparametric procedures since they represent a breadth of the typical types of analyses found in social science research. It is our hope that students will confidently learn the content presented with the promise of future successful applications.

In addition, each statistical test includes a section that explains how to use the computer program SPSS. However, the organization of the book provides effective instruction of the nonparametric statistical procedures for those individuals with or without the software. Therefore, instructors (and students) can focus on learning the tests with a calculator, SPSS, or both.

A Note to the Student

We have written this book with you in mind. Each of us has had a great deal of experience working with students just like you. Over the course of that time, it has been our experience that most people outside of the fields of mathematics or hard sciences struggle with and are intimidated by statistics. Moreover, we have found that when statistical procedures are explicitly communicated in a step-by-step manner, almost anyone can use them.

This book begins with a brief introduction (Chapter 1) and is followed with an explanation of how to perform the crucial step of checking your data for normality (Chapter 2). The chapters that follow (Chapters 3–9) highlight several nonparametric statistical procedures. Each of those chapters focuses on a particular type of variable and/or sample condition.

Chapters 3–9 each have a similar organization. They each explain the statistical methods included in their respective chapters. At least one sample problem is included for each test using a step-by-step approach. (In some cases, we provide additional sample problems when procedures differ between large and small samples.) Then, those same sample problems are demonstrated using the statistical software package SPSS. Whether or not your instructor incorporates SPSS, this section will give you the opportunity to learn how to use the program. Toward the end of each chapter, we identify examples of the tests in published research. Finally, we present sample problems with solutions.

As you seek to learn nonparametric statistics, we strongly encourage you to work through the sample problems. Then, using the sample problems as a reference, work through the problems at the end of the chapters and additional data sets provided.

New to the Second Edition

Given an opportunity to write a second edition of this book, we revised and expanded several portions. Our changes are based on feedback from users and reviewers.

We asked several undergraduate and graduate students for feedback on Chapters 1 and 2. Based on their suggestions, we made several minor changes to Chapter 1 with a goal to improve understanding. In Chapter 2, we expanded the section that describes and demonstrates the Kolmogorov–Smirnov (K-S) one-sample test.

After examining current statistics textbooks and emerging research paper, we decided to include two additional tests. We added the sign test to Chapter 3 and the Kolmogorov–Smirnov (K-S) two-sample test to Chapter 4. We also added a discussion on statistical power to Chapter 3 as requested by instructors who had adopted our book for their courses.

Since our book's first edition, SPSS has undergone several version updates. Our new edition of the book also has updated directions and screen captures for images of SPSS. Specifically, these changes reflect SPSS version 21.

We have included web-based tools to support our book's new edition. If you visit the publisher's book support website, you will find a link to a Youtube channel that includes narrated screen casts. The screen casts demonstrate how to use SPSS to perform the tests included in this book. The publisher's book support website also includes a link to a decision tree that helps the user determine an appropriate type of statistical test. The decision tree is organized using Prezi. The branches terminate with links to the screen casts on YouTube.

Gregory W. Corder
Dale I. Foreman

LIST OF VARIABLES

English Symbols

C	number of columns in a contingency table; number of categories
C_F	tie correction for the Friedman test
C_H	tie correction for the Kruskal–Wallis H -test
$D, \tilde{D}$	divergence between values from cumulative frequency distributions
D_i	difference between a ranked pair
df	degrees of freedom
f_e	expected frequency
f_o	observed frequency
$\hat{f}_r$	empirical frequency value
F_r	Friedman test statistic
g	number of tied groups in a variable
h	correction for continuity
H	Kruskal–Wallis test statistic
H_A	alternate hypothesis
H_O	null hypothesis
k	number of groups
K	kurtosis
M	midpoint of a sample
n	sample size
N	total number of values in a contingency table
p	probability
P_i	a category's proportion with respect to all categories
r_b	biserial correlation coefficient for a sample
r_{pb}	point-biserial correlation coefficient for a sample
r_s	Spearman rank-order correlation coefficient for a sample
R	number of runs; number of rows in a contingency table
R_i	sum of the ranks from a particular sample
s	standard deviation of a sample
S_k	skewness
SE	standard error
t	t statistic
t_i	number of tied values in a tie group
T	Wilcoxon signed rank test statistic
U	Mann–Whitney test statistic
$\bar{x}$	sample mean
y	height of the unit normal curve ordinate at the point dividing two proportions

xiv LIST OF VARIABLES

z	the number of standard deviations away from the mean
Z	Kolmogorov–Smirnov test statistic

Greek Symbols

α	alpha, probability of making a type I error
α_B	adjusted level of risk using the Bonferroni procedure
β	beta, probability of making a type II error
θ	theta, median of a population
μ	mu, mean value for a population
ρ	rho, correlation coefficient for a population
σ	sigma, standard deviation of a population
Σ	sigma, summation
χ^2	chi-square test statistic

NONPARAMETRIC STATISTICS: AN INTRODUCTION

1.1 OBJECTIVES

In this chapter, you will learn the following items:

- The difference between parametric and nonparametric statistics.
- How to rank data.
- How to determine counts of observations.

1.2 INTRODUCTION

If you are using this book, it is possible that you have taken some type of introductory statistics class in the past. Most likely, your class began with a discussion about probability and later focused on particular methods of dealing with populations and samples. Correlations, z -scores, and t -tests were just some of the tools you might have used to describe populations and/or make inferences about a population using a simple random sample.

Many of the tests in a traditional, introductory statistics text are based on samples that follow certain assumptions called parameters. Such tests are called *parametric tests*. Specifically, parametric assumptions include samples that

- are randomly drawn from a normally distributed population,
- consist of independent observations, except for paired values,
- consist of values on an interval or ratio measurement scale,
- have respective populations of approximately equal variances,
- are adequately large,* and
- approximately resemble a normal distribution.

*The minimum sample size for using a parametric statistical test varies among texts. For example, Pett (1997) and Salkind (2004) noted that most researchers suggest $n > 30$. Warner (2008) encouraged considering $n > 20$ as a minimum and $n > 10$ per group as an absolute minimum.

If any of your samples breaks one of these rules, you violate the assumptions of a parametric test. You do have some options, however.

You might change the nature of your study so that your data meet the needed parameters. For instance, if you are using an ordinal or nominal measurement scale, you might redesign your study to use an interval or ratio scale. (See Box 1.1 for a description of measurement scales.) Also, you might seek additional participants to enlarge your sample sizes. Unfortunately, there are times when one or neither of these changes is appropriate or even possible.

BOX 1.1

MEASUREMENT SCALES.

We can measure and convey variables in several ways. *Nominal* data, also called categorical data, are represented by counting the number of times a particular event or condition occurs. For example, you might categorize the political alignment of a group of voters. Group members could either be labeled democratic, republican, independent, undecided, or other. No single person should fall into more than one category.

A *dichotomous* variable is a special classification of nominal data; it is simply a measure of two conditions. A dichotomous variable is either discrete or continuous. A *discrete dichotomous* variable has no particular order and might include such examples as gender (male vs. female) or a coin toss (heads vs. tails). A *continuous dichotomous* variable has some type of order to the two conditions and might include measurements such as pass/fail or young/old.

Ordinal scale data describe values that occur in some order of rank. However, distance between any two ordinal values holds no particular meaning. For example, imagine lining up a group of people according to height. It would be very unlikely that the individual heights would increase evenly. Another example of an ordinal scale is a Likert-type scale. This scale asks the respondent to make a judgment using a scale of three, five, or seven items. The range of such a scale might use a 1 to represent *strongly disagree* while a 5 might represent *strongly agree*. This type of scale can be considered an ordinal measurement since any two respondents will vary in their interpretation of scale values.

An *interval* scale is a measure in which the relative distances between any two sequential values are the same. To borrow an example from the physical sciences, we consider the Celsius scale for measuring temperature. An increase from -8 to -7°C degrees is identical to an increase from 55 to 56°C .

A *ratio* scale is slightly different from an interval scale. Unlike an interval scale, a ratio scale has an absolute zero value. In such a case, the zero value indicates a measurement limit or a complete absence of a particular condition. To borrow another example from the physical sciences, it would be appropriate to measure light intensity with a ratio scale. Total darkness is a complete absence of light and would receive a value of zero.

On a general note, we have presented a classification of measurement scales similar to those used in many introductory statistics texts. To the best of our knowledge, this hierarchy of scales was first made popular by Stevens (1946). While Stevens has received agreement (Stake, 1960; Townsend & Ashby, 1984) and criticism (Anderson, 1961; Gaito, 1980; Velleman & Wilkinson, 1993), we believe the scale classification we present suits the nature and organization of this book. We direct anyone seeking additional information on this subject to the preceding citations.

If your samples do not resemble a normal distribution, you might have learned a strategy that modifies your data for use with a parametric test. First, if you can justify your reasons, you might remove extreme values from your samples called outliers. For example, imagine that you test a group of children and you wish to generalize the findings to typical children in a normal state of mind. After you collect the test results, most children earn scores around 80% with some scoring above and below the average. Suppose, however, that one child scored a 5%. If you find that this child speaks no English because he arrived in your country just yesterday, it would be reasonable to exclude his score from your analysis. Unfortunately, outlier removal is rarely this straightforward and deserves a much more lengthy discussion than we offer here.* Second, you might utilize a parametric test by applying a mathematical transformation to the sample values. For example, you might square every value in a sample. However, some researchers argue that transformations are a form of data tampering or can distort the results. In addition, transformations do not always work, such as circumstances when data sets have particularly long tails. Third, there are more complicated methods for analyzing data that are beyond the scope of most introductory statistics texts. In such a case, you would be referred to a statistician.

Fortunately, there is a family of statistical tests that do not demand all the parameters, or rules, that we listed earlier. They are called *nonparametric tests*, and this book will focus on several such tests.

1.3 THE NONPARAMETRIC STATISTICAL PROCEDURES PRESENTED IN THIS BOOK

This book describes several popular nonparametric statistical procedures used in research today. Table 1.1 identifies an overview of the types of tests presented in this book and their parametric counterparts.

TABLE 1.1

Type of analysis	Nonparametric test	Parametric equivalent
Comparing two related samples	Wilcoxon signed ranks test and sign test	<i>t</i> -Test for dependent samples
Comparing two unrelated samples	Mann–Whitney <i>U</i> -test and Kolmogorov–Smirnov two-sample test	<i>t</i> -Test for independent samples
Comparing three or more related samples	Friedman test	Repeated measures, analysis of variance (ANOVA)
Comparing three or more unrelated samples	Kruskal–Wallis <i>H</i> -test	One-way ANOVA

(Continued)

*Malthouse (2001) and Osborne and Overbay (2004) presented discussions about the removal of outliers.

TABLE 1.1 (Continued)

Type of analysis	Nonparametric test	Parametric equivalent
Comparing categorical data	Chi square (χ^2) tests and Fisher exact test	None
Comparing two rank-ordered variables	Spearman rank-order correlation	Pearson product–moment correlation
Comparing two variables when one variable is discrete dichotomous	Point-biserial correlation	Pearson product–moment correlation
Comparing two variables when one variable is continuous dichotomous	Biserial correlation	Pearson product–moment correlation
Examining a sample for randomness	Runs test	None

When demonstrating each nonparametric procedure, we will use a particular step-by-step method.

1.3.1 State the Null and Research Hypotheses

First, we state the hypotheses for performing the test. The two types of hypotheses are null and alternate. The *null hypothesis* (H_0) is a statement that indicates no difference exists between conditions, groups, or variables. The *alternate hypothesis* (H_A), also called a research hypothesis, is the statement that predicts a difference or relationship between conditions, groups, or variables.

The alternate hypothesis may be directional or nondirectional, depending on the context of the research. A directional, or one-tailed, hypothesis predicts a statistically significant change in a particular direction. For example, a treatment that predicts an improvement would be directional. A nondirectional, or two-tailed, hypothesis predicts a statistically significant change, but in no particular direction. For example, a researcher may compare two new conditions and predict a difference between them. However, he or she would not predict which condition would show the largest result.

1.3.2 Set the Level of Risk (or the Level of Significance) Associated with the Null Hypothesis

When we perform a particular statistical test, there is always a chance that our result is due to chance instead of any real difference. For example, we might find that two samples are significantly different. Imagine, however, that no real difference exists. Our results would have led us to reject the null hypothesis when it was actually true. In this situation, we made a type I error. Therefore, statistical tests assume some level of risk that we call alpha, or α .

There is also a chance that our statistical results would lead us to not reject the null hypothesis. However, if a real difference actually does exist, then we made a type II error. We use the Greek letter beta, β , to represent a type II error. See Table 1.2 for a summary of type I and type II errors.

TABLE 1.2

	We do not reject the null hypothesis	We reject the null hypothesis
The null hypothesis is actually true	No error	Type-I error, α
The null hypothesis is actually false	Type-II error, β	No error

After the hypotheses are stated, we choose the level of risk (or the level of significance) associated with the null hypothesis. We use the commonly accepted value of $\alpha = 0.05$. By using this value, there is a 95% chance that our statistical findings are real and not due to chance.

1.3.3 Choose the Appropriate Test Statistic

We choose a particular type of test statistic based on characteristics of the data. For example, the number of samples or groups should be considered. Some tests are appropriate for two samples, while other tests are appropriate for three or more samples.

Measurement scale also plays an important role in choosing an appropriate test statistic. We might select one set of tests for nominal data and a different set for ordinal variables. A common ordinal measure used in social and behavioral science research is the Likert scale. Nanna and Sawilowsky (1998) suggested that nonparametric tests are more appropriate for analyses involving Likert scales.

1.3.4 Compute the Test Statistic

The test statistic, or obtained value, is a computed value based on the particular test you need. Moreover, the method for determining the obtained value is described in each chapter and varies from test to test. For small samples, we use a procedure specific to a particular statistical test. For large samples, we approximate our data to a normal distribution and calculate a z -score for our data.

1.3.5 Determine the Value Needed for Rejection of the Null Hypothesis Using the Appropriate Table of Critical Values for the Particular Statistic

For small samples, we reference a table of critical values located in Appendix B. Each table provides a critical value to which we compare a computed test statistic. Finding a critical value using a table may require you to use such data characteristics

as the degrees of freedom, number of samples, and/or number of groups. In addition, you may need the desired level of risk, or alpha (α).

For large samples, we determine a critical region based on the level of risk (or the level of significance) associated with the null hypothesis, α . We will determine if the computed z -score falls within a critical region of the distribution.

1.3.6 Compare the Obtained Value with the Critical Value

Comparing the obtained value with the critical value allows us to identify a difference or relationship based on a particular level of risk. Once this is accomplished, we can state whether we must reject or must not reject the null hypothesis. While this type of phrasing may seem unusual, the standard practice in research is to state results in terms of the null hypothesis.

Some of the critical value tables are limited to particular sample or group size(s). When a sample size exceeds a table's range of value(s), we approximate our data to a normal distribution. In such cases, we use Table B.1 in Appendix B to establish a critical region of z -scores. Then, we calculate a z -score for our data and compare it with a critical region of z -scores. For example, if we use a two-tailed test with $\alpha = 0.05$, we do not reject the null hypothesis if the z -score is between -1.96 and $+1.96$. In other words, we do not reject if the null hypothesis if $-1.96 \leq z \leq 1.96$.

1.3.7 Interpret the Results

We can now give meaning to the numbers and values from our analysis based on our context. If sample differences were observed, we can comment on the strength of those differences. We can compare the observed results with the expected results. We might examine a relationship between two variables for its relative strength or search a series of events for patterns.

1.3.8 Reporting the Results

Communicating results in a meaningful and comprehensible manner makes our research useful to others. There is a fair amount of agreement in the research literature for reporting statistical results from parametric tests. Unfortunately, there is less agreement for nonparametric tests. We have attempted to use the more common reporting techniques found in the research literature.

1.4 RANKING DATA

Many of the nonparametric procedures involve ranking data values. Ranking values is really quite simple. Suppose that you are a math teacher and wanted to find out if students score higher after eating a healthy breakfast. You give a test and compare the scores of four students who ate a healthy breakfast with four students who did not. Table 1.3 shows the results.

TABLE 1.3

Students who ate breakfast	Students who skipped breakfast
87	93
96	83
92	79
84	73

To rank all of the values from Table 1.3 together, place them all in order in a new table from smallest to largest (see Table 1.4). The first value receives a rank of 1, the second value receives a rank of 2, and so on.

TABLE 1.4

Value	Rank
73	1
79	2
83	3
84	4
87	5
92	6
93	7
96	8

Notice that the values for the students who ate breakfast are in bold type. On the surface, it would appear that they scored higher. However, if you are seeking statistical significance, you need some type of procedure. The following chapters will offer those procedures.

1.5 RANKING DATA WITH TIED VALUES

The aforementioned ranking method should seem straightforward. In many cases, however, two or more of the data values may be repeated. We call repeated values *ties*, or tied values. Say, for instance, that you repeat the preceding ranking with a different group of students. This time, you collected new values shown in Table 1.5.

TABLE 1.5

Students who ate breakfast	Students who skipped breakfast
90	75
85	80
95	55
70	90

Rank the values as in the previous example. Notice that the value of 90 is repeated. This means that the value of 90 is a tie. If these two student scores were different, they would be ranked 6 and 7. In the case of a tie, give all of the tied values the average of their rank values. In this example, the average of 6 and 7 is 6.5 (see Table 1.6).

TABLE 1.6

Value	Rank ignoring tied values	Rank accounting for tied values
55	1	1
70	2	2
75	3	3
80	4	4
85	5	5
90	6	6.5
90	7	6.5
95	8	8

Most nonparametric statistical tests require a different formula when a sample of data contains ties. It is important to note that the formulas for ties are more algebraically complex. What is more, formulas for ties typically produce a test statistic that is only slightly different from the test statistic formulas for data without ties. It is probably for this reason that most statistics texts omit the formulas for tied values. As you will see, however, we include the formulas for ties along with examples where applicable.

When the statistical tests in this book are explained using the computer program SPSS® (Statistical Package for Social Scientists), there is no mention of any special treatment for ties. That is because SPSS automatically detects the presence of ties in any data sets and applies the appropriate procedure for calculating the test statistic.

1.6 COUNTS OF OBSERVATIONS

Some nonparametric tests require *counts* (or frequencies) of observations. Determining the count is fairly straightforward and simply involves counting the total number of times a particular observations is made. For example, suppose you ask several children to pick their favorite ice cream flavor given three choices: vanilla, chocolate, and strawberry. Their preferences are shown in Table 1.7.

To find the counts for each ice cream flavor, list the choices and tally the total number of children who picked each flavor. In other words, count the number of children who picked chocolate. Then, repeat for the other choices, vanilla and strawberry. Table 1.8 reveals the counts from Table 1.7.

TABLE 1.7

Participant	Flavor
1	Chocolate
2	Chocolate
3	Vanilla
4	Vanilla
5	Strawberry
6	Chocolate
7	Chocolate
8	Vanilla

TABLE 1.8

Flavor	Count
Chocolate	4
Vanilla	3
Strawberry	1

To check your accuracy, you can add all the counts and compare them with the number of participants. The two numbers should be the same.

1.7 SUMMARY

In this chapter, we described differences between parametric and nonparametric tests. We also addressed assumptions by which nonparametric tests would be favorable over parametric tests. Then, we presented an overview of the nonparametric procedures included in this book. We also described the step-by-step approach we use to explain each test. Finally, we included explanations and examples of ranking and counting data, which are two tools for managing data when performing particular nonparametric tests.

The chapters that follow will present step-by-step directions for performing these statistical procedures both by manual, computational methods and by computer analysis using SPSS. In the next chapter, we address procedures for comparing data samples with a normal distribution.

1.8 PRACTICE QUESTIONS

1. Male high school students completed the 1-mile run at the end of their 9th grade and the beginning of their 10th grade. The following values represent the differences between the recorded times. Notice that only one student’s time improved (−2:08). Rank the values in Table 1.9 beginning with the student’s time difference that displayed improvement.

TABLE 1.9

Participant	Value	Rank
1	0:36	
2	0:28	
3	1:41	
4	0:37	
5	1:01	
6	2:30	
7	0:44	
8	0:47	
9	0:13	
10	0:24	
11	0:51	
12	0:09	
13	−2:08	
14	0:12	
15	0:56	

2. The values in Table 1.10 represent weekly quiz scores on math. Rank the quiz scores.

TABLE 1.10

Participant	Score	Rank
1	100	
2	60	
3	70	
4	90	
5	80	
6	100	
7	80	
8	20	
9	100	
10	50	

3. Using the data from the previous example, what are the counts (or frequencies) of passing scores and failing scores if a 70 is a passing score?

1.9 SOLUTIONS TO PRACTICE QUESTIONS

- 1. The value ranks are listed in Table 1.11. Notice that there are no ties.
- 2. The value ranks are listed in Table 1.12. Notice the tied values. The value of 80 occurred twice and required averaging the rank values of 5 and 6.

TABLE 1.11

Participant	Value	Rank
1	0:36	7
2	0:28	6
3	1:41	14
4	0:37	8
5	1:01	13
6	2:30	15
7	0:44	9
8	0:47	10
9	0:13	4
10	0:24	5
11	0:51	11
12	0:09	2
13	-2:08	1
14	0:12	3
15	0:56	12

TABLE 1.12

Participant	Score	Rank
1	100	9
2	60	3
3	70	4
4	90	7
5	80	5.5
6	100	9
7	80	5.5
8	20	1
9	100	9
10	50	2

$$(5 + 6) \div 2 = 5.5$$

The value of 100 occurred three times and required averaging the rank values of 8, 9, and 10.

$$(8 + 9 + 10) \div 3 = 9$$

3. Table 1.13 shows the passing scores and failing scores using 70 as a passing score. The counts (or frequencies) of passing scores is $n_{\text{passing}} = 7$. The counts of failing scores is $n_{\text{failing}} = 3$.

TABLE 1.13

Participant	Score	Pass/Fail
1	100	Pass
2	60	Fail
3	70	Pass
4	90	Pass
5	80	Pass
6	100	Pass
7	80	Pass
8	20	Fail
9	100	Pass
10	50	Fail

TESTING DATA FOR NORMALITY

2.1 OBJECTIVES

In this chapter, you will learn the following items:

- How to find a data sample's kurtosis and skewness and determine if the sample meets acceptable levels of normality.
- How to use SPSS® to find a data sample's kurtosis and skewness and determine if the sample meets acceptable levels of normality.
- How to perform a Kolmogorov–Smirnov one-sample test to determine if a data sample meets acceptable levels of normality.
- How to use SPSS to perform a Kolmogorov–Smirnov one-sample test to determine if a data sample meets acceptable levels of normality.

2.2 INTRODUCTION

Parametric statistical tests, such as the *t*-test and one-way analysis of variance, are based on particular assumptions or parameters. The data samples meeting those parameters are randomly drawn from a normal population, based on independent observations, measured with an interval or ratio scale, possess an adequate sample size (see Chapter 1), and approximately resemble a normal distribution. Moreover, comparisons of samples or variables should have approximately equal variances. If data samples violate one or more of these assumptions, you should consider using a nonparametric test.

Examining the data gathering method, scale type, and size of a sample are fairly straightforward. However, examining a data sample's resemblance to a normal distribution, or its normality, requires a more involved analysis. Visually inspecting a graphical representation of a sample, such as a stem and leaf plot or a box and whisker plot, might be the most simplistic examination of normality. Statisticians advocate this technique in beginning statistics; however, this measure of normality does not suffice for strict levels of defensible analyses.