

THIRD EDITION

STATISTICS ALIVE!



WENDY J. STEINBERG • MATTHEW PRICE



$$SS_{\text{bet}} = n_g \sum_1^k (M_g - M_{\text{tot}})^2$$

ANOVA sums of squares between, deviation formula

$$SS_{\text{bet}} = \sum_1^k \frac{(\sum X_g)^2}{n_g} - \frac{(\sum X_{\text{tot}})^2}{N}$$

ANOVA sums of squares between, computational (raw score) formula

$$MS = \frac{SS}{df}$$

mean square

$$F = \frac{MS_{\text{bet obs}} - MS_{\text{bet exp}}}{MS_{\text{with}}}$$

F test

$$\text{HSD} = q \sqrt{\frac{MS_{\text{with}}}{n_g}}$$

Tukey post hoc test

$$\chi^2 = \sum \left[\frac{f_o - f_e}{f_e} \right]^2$$

chi-square

$$E = \frac{RT \times CT}{N}$$

expected frequencies for chi-square test of independence

$$d = \frac{M_1 - M_2}{s_{DV}}$$

Cohen's d effect size, for t test

$$\text{Effect size } r = \sqrt{\frac{t^2}{t^2 + df}}$$

effect size r , for t test

$$\text{Effect size } \eta = \sqrt{\frac{SS_{\text{bet}}}{SS_{\text{tot}}}}$$

effect size eta, for ANOVA

$$\text{Effect size } r = \sqrt{\frac{\chi^2}{(N)(C - 1)}}$$

effect size r , for chi-square goodness of fit

$$\phi = \sqrt{\frac{\chi^2}{N}}$$

effect size phi, for 2×2 chi-square test of independence

$$V = \sqrt{\frac{\phi^2}{(\text{the smaller of } R \text{ or } C) - 1}}$$

Cramer's V effect size, for chi-square test of independence of any size

$$r_{XY} = \frac{\sum z_X z_Y}{N}$$

definitional formula for Pearson r

$$r_{XY} = \frac{\sum (X - M_X)(Y - M_Y)}{N s_X s_Y}$$

deviation formula for Pearson r

$$r_{XY} = \frac{N \sum XY - (\sum X)(\sum Y)}{\sqrt{[N(\sum X^2) - (\sum X)^2][N(\sum Y^2) - (\sum Y)^2]}}$$

computational (raw score) formula for Pearson r

$$Y = bX + a$$

equation for any straight line

$$Y' = r_{XY} \left(\frac{s_Y}{s_X} \right) X - r_{XY} \left(\frac{s_Y}{s_X} \right) M_X + M_Y$$

equation for predicting Y from X

$$s_{YX} = s_Y \sqrt{1 - r_{XY}^2}$$

standard error of prediction

$$Y' = b_1 X_1 + b_2 X_2 + b_3 X_3 + \dots + a$$

equation for multiple regression

*Seldom do most of us meet anyone in our lifetimes as kind, helpful, and guileless as my husband.
I had the good fortune not only to meet such a person, but also to marry him.*

To you, Bob Sanford, this textbook is lovingly dedicated.

—Wendy J. Steinberg

*To my bashert, Peggy Price, thank you for your unending support and love.
The work in this book, and across my career, could not have happened without you in my life.
Thank you for making the lives of David, Leah, and myself that much richer.*

This book is dedicated to you.

—Matthew Price

Sara Miller McCune founded SAGE Publishing in 1965 to support the dissemination of usable knowledge and educate a global community. SAGE publishes more than 1000 journals and over 600 new books each year, spanning a wide range of subject areas. Our growing selection of library products includes archives, data, case studies and video. SAGE remains majority owned by our founder and after her lifetime will become owned by a charitable trust that secures the company's continued independence.

Los Angeles | London | New Delhi | Singapore | Washington DC | Melbourne

Statistics Alive!

Third Edition

Wendy J. Steinberg

Matthew Price

University of Vermont



Los Angeles | London | New Delhi
Singapore | Washington DC | Melbourne



FOR INFORMATION:

SAGE Publications, Inc.
2455 Teller Road
Thousand Oaks, California 91320
E-mail: order@sagepub.com

SAGE Publications Ltd.
1 Oliver's Yard
55 City Road
London EC1Y 1SP
United Kingdom

SAGE Publications India Pvt. Ltd.
B 1/1 Mohan Cooperative Industrial Area
Mathura Road, New Delhi 110 044
India

SAGE Publications Asia-Pacific Pte. Ltd.
18 Cross Street #10-10/11/12
China Square Central
Singapore 048423

Acquisitions Editor: Leah Fargotstein
Editorial Assistant: Natalie Elliott
Production Editor: Bennie Clark Allen
Copy Editor: QuADS Prepress Pvt. Ltd.
Typesetter: C&M Digital (P) Ltd.
Proofreader: Jen Grubba
Indexer: Integra
Cover Designer: Candice Harman
Marketing Manager: Victoria Velasquez

Copyright © 2021 by SAGE Publications, Inc.

All rights reserved. Except as permitted by U.S. copyright law, no part of this work may be reproduced or distributed in any form or by any means, or stored in a database or retrieval system, without permission in writing from the publisher.

All trademarks depicted within this book, including trademarks appearing as part of a screenshot, figure, or other image are included solely for the purpose of illustration and are the property of their respective holders. The use of the trademarks in no way indicates any relationship with, or endorsement by, the holders of said trademarks. SPSS is a registered trademark of International Business Machines Corporation.

Printed in the United States of America

Library of Congress Cataloging-in-Publication Data

Names: Steinberg, Wendy J., author. | Price, Matthew, author.

Title: Statistics alive! / Wendy J. Steinberg, Matthew Price.

Description: Third Edition. | Thousand Oaks : SAGE Publishing, 2020. | Revised edition of Statistics alive!, c2011. | Includes bibliographical references and index.

Identifiers: LCCN 2020010457 | ISBN 9781544328263 (paperback) | ISBN 9781544328249 (epub) | ISBN 9781544328256 (epub) | ISBN 9781544328270 (pdf)

Subjects: LCSH: Statistics. | Social sciences—Statistical methods.

Classification: LCC HA29 .S7974 2020 | DDC 519.5—dc23
LC record available at <https://lccn.loc.gov/2020010457>

This book is printed on acid-free paper.

20 21 22 23 24 10 9 8 7 6 5 4 3 2 1

Detailed Contents

List of Figures	xx
List of Tables	xxiv
Preface	xxvi
Supplemental Material for Use With <i>Statistics Alive!</i>	xxxi
Acknowledgments	xxxii
About the Authors	xxxiii

PART I PRELIMINARY INFORMATION: “FIRST THINGS FIRST”

1

Module 1 Math Review, Vocabulary, and Symbols

2

Getting Started	2
Common Terms and Symbols in Statistics	3
Fundamental Rules and Procedures for Statistics	5
More Rules and Procedures	10

Module 2 Measurement Scales

12

What Is Measurement?	12
Scales of Measurement	12
Nominal Scale	12
Ordinal Scale	13
Interval Scale	14
Ratio Scale	15
Continuous Versus Discrete Variables	17
Real Limits	19

PART II TABLES AND GRAPHS: “ON DISPLAY”

21

Module 3 Frequency and Percentile Tables

22

Why Use Tables?	22
Frequency Tables	23
Relative Frequency or Percentage Tables	27
Grouped Frequency Tables	30
Percentile and Percentile Rank Tables	34
SPSS Connection	38

Module 4 Graphs and Plots

42

Why Use Graphs?	42
Graphing Continuous Data	42
Principles of Graph Construction	42
Histogram	45
Frequency Curve or Line Graph	47
Boxplot or Box-and-Whisker Plot	48

Symmetry, Skew, and Kurtosis	50
Graphing Discrete Data	54
Bar Graph	54
Pie Graph	55
SPSS Connection	60
PART III CENTRAL TENDENCY: "BULL'S-EYE"	65
Module 5 Mode, Median, and Mean	66
What Is Central Tendency?	66
Mode	66
Median	68
Mean	72
Skew and Central Tendency	76
SPSS Connection	80
PART IV DISPERSION: "FROM HERE TO ETERNITY"	81
Module 6 Range, Variance, and Standard Deviation	82
What Is Dispersion?	82
Range	82
Variance	83
Standard Deviation	86
Mean Absolute Deviation	87
Controversy: N Versus $n - 1$	90
SPSS Connection	91
PART V THE NORMAL CURVE AND STANDARD SCORES: "WHAT'S THE SCORE?"	93
Module 7 Percent Area and the Normal Curve	94
What Is a Normal Curve?	94
History of the Normal Curve	95
Uses of the Normal Curve	96
Looking Ahead	100
Module 8 z Scores	101
What Is a Standard Score?	101
Benefits of Standard Scores	101
Calculating z Scores	104
Comparing Scores Across Different Tests	107
SPSS Connection	110
Module 9 Score Transformations and Their Effects	112
Why Transform Scores?	112
Effects on Central Tendency	112
Effects on Dispersion	115
A Graphic Look at Transformations	118

Summary of Transformation Effects	118
Some Common Transformed Scores	120
Looking Ahead	122
PART VI PROBABILITY: “ODDS ARE”	125
Module 10 Probability Definitions and Theorems	126
Why Study Probability?	126
Probability as a Proportion	126
Equally Likely Model	127
Mutually Exclusive Outcomes	128
Addition Theorem	130
Independent Outcomes	131
Multiplication Theorem	133
A Brief Review	134
Probability and Inference	136
Module 11 The Binomial Distribution	137
What Are Dichotomous Events?	137
Finding Probabilities by Listing and Counting	138
Finding Probabilities by the Binomial Formula	144
Finding Probabilities by the Binomial Table	147
Probability and Experimentation	148
Looking Ahead	150
Nonnormal Data	151
PART VII INFERENTIAL THEORY: “OF TRUTH AND RELATIVITY”	155
Module 12 Sampling, Variables, and Hypotheses	156
From Description to Inference	156
Sampling	156
Variables	159
Hypotheses	162
Module 13 Errors and Significance	165
Random Sampling Revisited	165
Sampling Error	165
Significant Difference	166
The Decision Table	166
Type I Error	167
Type II Error	169
Module 14 The z Score as a Hypothesis Test	174
Inferential Logic and the z Score	174
Constructing a Hypothesis Test for a z Score	175
Looking Ahead	176

PART VIII THE ONE-SAMPLE TEST: “ARE THEY FROM OUR PART OF TOWN?”	179
Module 15 Standard Error of the Mean	180
Central Limit Theorem	180
Sampling Distribution of the Mean	184
Calculating the Standard Error of the Mean	185
Sample Size and the Standard Error of the Mean	185
Looking Ahead	191
Module 16 Normal Deviate Z Test	193
Prototype Logic and the Z Test	193
Calculating a Normal Deviate Z Test	194
Examples of Normal Deviate Z Tests	196
Decision Making With a Normal Deviate Z Test	197
Looking Ahead	199
Module 17 One-Sample t Test	200
Z Test Versus t Test	200
Comparison of Z-Test and t -Test Formulas	200
Degrees of Freedom	202
Biased and Unbiased Estimates	202
When Do We Reject the Null Hypothesis?	203
One-Tailed Versus Two-Tailed Tests	204
The t Distribution Versus the Normal Distribution	206
The t Table Versus the Normal Curve Table	207
Calculating a One-Sample t Test	208
Interpreting a One-Sample t Test	210
Looking Ahead	215
SPSS Connection	215
Module 18 Interpreting and Reporting One-Sample t: Error, Confidence, and Parameter Estimates	217
What It Means to Reject the Null	217
Refining Error	217
Decision Making With a One-Sample t Test	220
Dichotomous Decisions Versus Reports of Actual p	221
Parameter Estimation: Point and Interval	222
SPSS Connection	226
PART IX THE TWO-SAMPLE TEST: “OURS IS BETTER THAN YOURS”	229
Module 19 Standard Error of the Difference Between the Means	230
One-Sample Versus Two-Sample Studies	230
Sampling Distribution of the Difference Between the Means	232

Calculating the Standard Error of the Difference Between the Means	234
Importance of the Size of the Standard Error of the Difference Between the Means	235
Looking Ahead	237
Module 20 t Test With Independent Samples and Equal Sample Sizes	238
A Two-Sample Study	238
Inferential Logic and the Two-Sample t Test	239
Calculating a Two-Sample t Test	241
Interpreting a Two-Sample t Test	244
Looking Ahead	250
SPSS Connection	250
Module 21 t Test With Unequal Sample Sizes	252
What Makes Sample Sizes Unequal?	252
Comparison of Special-Case and Generalized Formulas	252
Calculating a t Test With Unequal Sample Sizes	255
Interpreting a t Test With Unequal Sample Sizes	257
SPSS Connection	262
Module 22 t Test With Related Samples	264
What Makes Samples Related?	264
Comparison of Special-Case and Related-Samples Formulas	266
Advantage and Disadvantage of Related Samples	269
Direct-Difference Formula	269
Calculating a t Test With Related Samples	270
Interpreting a t Test With Related Samples	272
SPSS Connection	276
Module 23 Interpreting and Reporting Two-Sample t: Error, Confidence, and Parameter Estimates	278
What Is Confidence?	278
Refining Error and Confidence	279
Decision Making With a Two-Sample t Test	281
Dichotomous Decisions Versus Reports of Actual p	282
Parameter Estimation: Point and Interval	283
SPSS Connection	289
PART X THE MULTISAMPLE TEST: "OURS IS BETTER THAN YOURS OR THEIRS"	291
Module 24 ANOVA Logic: Sums of Squares, Partitioning, and Mean Squares	292
When Do We Use ANOVA?	292
ANOVA Assumptions	293
Partitioning of Deviation Scores	294

From Deviation Scores to Variances	295
From Variances to Mean Squares	296
From Mean Squares to F	297
Looking Ahead	298
Module 25 One-Way ANOVA: Independent Samples and Equal Sample Sizes	299
What Is a One-Way ANOVA?	299
Inferential Logic and ANOVA	300
Deviation Score Method	303
Sums of Squares Formulas	303
Calculating Sums of Squares	304
Raw Score Method	306
Sums of Squares Formulas	306
Calculating Sums of Squares	307
Remaining Steps for Both Methods: Mean Squares and F	308
Interpreting a One-Way ANOVA	309
The ANOVA Summary Table	311
SPSS Connection	316
PART XI POST HOC TESTS: "SO WHO'S RESPONSIBLE?"	319
Module 26 Tukey HSD Test	320
Why Do We Need a Post Hoc Test?	320
Calculating the Tukey HSD	321
Interpreting the Tukey HSD	323
SPSS Connection	326
Module 27 Scheffé Test	328
Why Do We Need a Post Hoc Test?	328
Calculating the Scheffé	328
Interpreting the Scheffé	332
SPSS Connection	337
PART XII MORE THAN ONE INDEPENDENT VARIABLE: "DOUBLE DUTCH JUMP ROPE"	339
Module 28 Main Effects and Interaction Effects	340
What Is a Factorial ANOVA?	340
Factorial ANOVA Designs	340
Number and Type of Hypotheses	342
Main Effects	342
Interaction Effects	344
Looking Ahead	346
Module 29 Factorial ANOVA	355
Review of Factorial ANOVA Designs	355
Data Setup and Preliminary Expectations	356

Sums of Squares Formulas	357
Calculating Factorial ANOVA Sums of Squares: Raw Score Method	360
Factorial Mean Squares and F s	362
Interpreting a Factorial F Test	363
The Factorial ANOVA Summary Table	364
SPSS Connection	371
PART XIII NONPARAMETRIC STATISTICS: “WITHOUT FORM OR VOID”	375
Module 30 One-Variable Chi-Square: Goodness of Fit	376
What Is a Nonparametric Test?	376
Chi-Square as a Goodness-of-Fit Test	377
Formula for Chi-Square	377
Inferential Logic and Chi-Square	378
Calculating a Chi-Square Goodness of Fit	379
Interpreting a Chi-Square Goodness of Fit	381
Looking Ahead	382
SPSS Connection	385
Module 31 Two-Variable Chi-Square: Test of Independence	388
Chi-Square as a Test of Independence	388
Prerequisites for a Chi-Square Test of Independence	389
Formula for a Chi-Square	390
Finding Expected Frequencies	390
Calculating a Chi-Square Test of Independence	392
Interpreting a Chi-Square Test of Independence	392
SPSS Connection	397
PART XIV EFFECT SIZE AND POWER: “HOW MUCH IS ENOUGH?”	399
Module 32 Measures of Effect Size	400
What Is Effect Size?	400
For Two-Sample t Tests	401
For ANOVA F Tests	403
For Chi-Square Tests	404
For a Goodness-of-Fit Test	404
For a Test of Independence	405
Module 33 Power and the Factors Affecting It	410
What Is Power?	410
Factors Affecting Power	412
Size of Type I Error	413
Size of the Actual Difference Between the Means	414
Amount of Error Variance	414
Sample Size	415
Putting It Together: Alpha, Power, Effect Size, and Sample Size	416
Looking Ahead	419

PART XV CORRELATION: “WHITHER THOU GOEST, I WILL GO” 421

Module 34 Relationship Strength and Direction 422

Experimental Versus Correlational Studies	422
Plotting Correlation Data	424
Relationship Strength	427
Relationship Direction	428
Linear and Nonlinear Relationships	431
Outliers and Their Effects	434
Looking Ahead	435
SPSS Connection	437

Module 35 Pearson r 439

What Is a Correlation Coefficient?	439
Calculation of a Pearson r	441
Formulas for Pearson r	442
z-Score Scatterplots and r	443
Calculating Pearson r : Deviation Score Method	445
Interpreting a Pearson r Coefficient	446
Looking Ahead	447
SPSS Connection	452

Module 36 Correlation Pitfalls 453

Effect of Sample Size on Statistical Significance	453
Statistical Significance Versus Practical Importance	454
Effect of Restriction in Range	455
Effect of Sample Heterogeneity or Homogeneity	456
Effect of Unreliability in the Measurement Instrument	457
Correlation Versus Causation	457

PART XVI LINEAR PREDICTION: “YOU'RE SO PREDICTABLE” 461

Module 37 Linear Prediction 462

Correlation Permits Prediction	462
Logic of a Prediction Line	463
Equation for the Best-Fitting Line	468
Using a Prediction Equation to Predict Scores on Y	471
Another Calculation Example	472
SPSS Connection	477

Module 38 Standard Error of Prediction 480

What Is a Confidence Interval?	480
Correlation and Prediction Error	480
Distribution of Prediction Error	482
Calculating the Standard Error of Prediction	482

Using the Standard Error of Prediction to Calculate Confidence Intervals	483
Factors Influencing the Standard Error of Prediction	486
Another Calculation Example	489
Module 39 Introduction to Multiple Regression	491
What Is Regression?	491
Prediction Error, Revisited	491
Why Multiple Regression?	492
The Multiple Regression Equation	493
Multiple Regression and Predicted Variance	494
Hypothesis Testing in Multiple Regression	496
An Example	498
The General Linear Model	503
SPSS Connection	505
 PART XVII REVIEW: “SAY IT AGAIN, SAM”	 509
Module 40 Selecting the Appropriate Analysis	510
Review of Descriptive Methods	510
Tables and Graphs	510
Descriptive Statistics	512
Review of Inferential Methods	512
Parametric Test Statistics	513
Nonparametric Test Statistics	515
Effect Size and Power	516
 Appendix A: Normal Curve Table	 518
Appendix B: Binomial Table	522
Appendix C: <i>t</i> Table	524
Appendix D: <i>F</i> Table (ANOVA)	525
Appendix E: Studentized Range Statistic (for Tukey HSD)	528
Appendix F: Chi-Square Table	529
Appendix G: Correlation Table	530
Appendix H: Odd Solutions to Textbook Exercises	531
References	582
Index	583

List of Figures

Figure 4.1	Template for a Vertical Graph	43
Figure 4.2	Three Graphs of BeachTown University SAT Scores	44
Figure 4.3	BeachTown University SAT Scores: Graphed According to Convention	45
Figure 4.4	Template for Creating a Histogram for Height of Basketball Team Members (in Inches)	46
Figure 4.5	Histogram for Statistics Test Scores in 5-Point Intervals	47
Figure 4.6	Histogram for Statistics Test Scores in 10-Point Intervals	47
Figure 4.7	Statistics Score Histogram With Frequency Curve Overlay and With Overlay Alone	48
Figure 4.8	Example of a Boxplot or a Box-and-Whisker Plot for Some Data	48
Figure 4.9	Template for Creating a Frequency Curve for Height of Basketball Team Members (in Inches)	49
Figure 4.10	Normal Curve	51
Figure 4.11	Skewed Distributions	51
Figure 4.12	Kurtotic Distributions	52
Figure 4.13	Bimodal Distribution	52
Figure 4.14	Uniform Distribution	52
Figure 4.15	Bar Graph for Reported Religious Affiliations in a Northeastern University Class	55
Figure 4.16	Pie Graph for Reported Religious Affiliations in a Northeastern University Class	56
Figure 5.1	Scores Balancing at Mean of 0	74
Figure 5.2	Scores Balancing at Mean of +6	74
Figure 5.3	Position of Mean, Median, and Mode in Normally Distributed Data	77
Figure 5.4	Position of Mode, Median, and Mean in Skewed Score Distributions	77
Figure 5.5	Position of the Mean in a Bimodal Distribution	78
Figure 7.1	Inflection Points in the Normal Curve	94
Figure 7.2	Standard Deviations in the Normal Curve	96
Figure 7.3	Normal Curve Percentages, to Two Decimal Places	97
Figure 7.4	Normal Curve Showing Means and Standard Deviations for a Variety of Common Scores	98
Figure 8.1	Normal Curve, With z Scores Indicated	104
Figure 8.2	Shaded Area Below z	105
Figure 8.3	Normal Curve, With Percentage Area Indicators	106
Figure 9.1	Effects of Score Transformations	118
Figure 9.2	Normal Curve With Standard Deviations and z Scores	120
Figure 9.3	Normal Curve With Various Transformed Score Scales	122
Figure 10.1	Embedded Outcomes	128
Figure 10.2	Overlapping Outcomes	129
Figure 11.1	Histogram for Two Tosses	139
Figure 11.2	Histogram for Three Tosses	141
Figure 11.3	Histograms for 4, 5, and 10 Tosses	142

Figure 11.4	Histogram for Five Tosses	146
Figure 11.5	Histogram With a Line Graph Superimposed	149
Figure 11.6	Skewed Histogram for Six Tosses of a Weighted Coin	152
Figure 13.1	Elements of a Decision Table	167
Figure 13.2	Null Hypothesis for Amount of Water Drunk	168
Figure 13.3	Correct Decision Under a True Null Hypothesis	168
Figure 13.4	Disparate Samples Under a True Null Hypothesis	169
Figure 13.5	Decision Table Showing Type I Error	169
Figure 13.6	False Null Hypothesis for Amount of Water Drunk	170
Figure 13.7	Correct Decision Under a False Null Hypothesis	170
Figure 13.8	Homogeneous Samples Under a False Null Hypothesis	171
Figure 13.9	Decision Table Showing Type II Error	171
Figure 15.1	Frequency Graph of Original Population	183
Figure 15.2	Template for Graphing Frequency of Sample Means	183
Figure 15.3	Sampling Distribution of the Mean	184
Figure 15.4	Numbered Slips of Paper	188
Figure 15.5	Size of Standard Error of the Mean for Two Different Sample Sizes	190
Figure 15.6	Sampling Distribution of the Mean, Showing the Value of σ_M	191
Figure 15.7	Sampling Distribution of the Mean, Showing σ_M and the Observed Sample Mean	191
Figure 16.1	Sampling Distribution of the Mean, Showing the Value and Location of the Normal Deviate Z	195
Figure 17.1	Regions of Rejection and Retention	203
Figure 17.2	Region of Rejection in a One-Tailed Test	204
Figure 17.3	Regions of Rejection in a Two-Tailed Test	205
Figure 17.4	Region of Rejection for a One-Tailed Test of Ms. Teachwell's Class	205
Figure 17.5	Change in t Distribution as a Function of Sample Size	206
Figure 17.6	Position of Ms. Teachwell's Class Relative to the Statewide Population	210
Figure 17.7	Ms. Teachwell's Class's t Relative to the .05 Region of Rejection	211
Figure 17.8	Ms. Teachwell's Class's t Relative to the .01 Region of Rejection	212
Figure 18.1	Ms. Teachwell's Class, Shown as an Unusual Sample From the Population	220
Figure 18.2	Ms. Teachwell's Class, Shown as a Typical Sample From the Population	221
Figure 19.1	Diagram of the Research Hypothesis	231
Figure 19.2	Diagram of the Null Hypothesis	231
Figure 19.3	Sampling Distribution of the Difference Between the Means	233
Figure 20.1	Sampling Distribution of the Difference Between the Means, Showing the Value of $S_{M_1-M_2}$	243
Figure 20.2	Sampling Distribution of the Difference Between the Means, Showing the Location of the Calculated t	245
Figure 21.1	Location of Observed Difference Between the Sample Means	257
Figure 22.1	Location of Observed Difference Between the Sample Means	272
Figure 23.1	True Null Hypothesis, With Unusual Samples	281
Figure 23.2	False Null Hypothesis, With Typical Samples	282
Figure 24.1	Three Different Treatment Populations	294
Figure 24.2	The F Distribution	297

Figure 25.1	Three Different Treatment Populations	304
Figure 28.1	Unconnected Group Means for a Two-Way ANOVA With Two Main Effects and No Interaction Effect	342
Figure 28.2	Connected Group Means for a Two-Way ANOVA With Two Main Effects and No Interaction Effect	343
Figure 28.3	Unconnected Group Means for a Two-Way ANOVA With an Interaction Effect and No Main Effects	345
Figure 28.4	Connected Group Means for a Two-Way ANOVA With an Interaction Effect and No Main Effects	345
Figure 28.5	Graph of Connected Group Means for a Two-Way ANOVA With a Less Obvious Interaction Effect	347
Figure 29.1	Three Different Treatment Populations	357
Figure 29.2	Two Separate Independent Variables	358
Figure 31.1	Blank Table of Observed and Expected Program Placement for Boys and Girls	391
Figure 31.2	Completed Table of Observed and Expected Program Placement for Boys and Girls	392
Figure 33.1	Rejecting the Null Hypothesis: Type I Error Versus a Treatment Effect	410
Figure 33.2	Retaining the Null Hypothesis: Type II Error Versus No Treatment Effect	411
Figure 33.3	Decision Table With Power Cell	411
Figure 33.4	Power and the Amount of Type I Error	413
Figure 33.5	Power and Directionality of the Alternative Hypothesis	413
Figure 33.6	Power and the Actual Difference Between the Means	414
Figure 34.1	Drawing a Scatterplot	424
Figure 34.2	Scores on Quiz 1 and Quiz 2	425
Figure 34.3	Scatterplots of Perfect and Zero Relationships	428
Figure 34.4	Two Scatterplots of Moderate Relationships	429
Figure 34.5	Scores on Quiz 1 and Quiz 2	430
Figure 34.6	Two Linear Relationships	431
Figure 34.7	An Inverted-U Relationship	432
Figure 34.8	An S-Shaped Relationship	432
Figure 34.9	Curvilinear Data Fitted With Linear and Curvilinear Lines	433
Figure 34.10	Data Without and With an Outlier	434
Figure 35.1	Four Quadrants in a Cartesian Coordinate System	443
Figure 35.2	Raw Score Scatterplot Versus z Scores in a Cartesian Coordinate System for Quiz 1 and Quiz 2	444
Figure 36.1	Effect of Restricted Range on Correlation	455
Figure 36.2	Effect of Homogeneity on Correlation	456
Figure 36.3	Effect of Inappropriate Heterogeneity on Correlation	456
Figure 37.1	Linear Relationship Between Degrees Celsius and Degrees Fahrenheit	463
Figure 37.2	Grid for Temperature Plot	464
Figure 37.3	Moderate Relationship Between SAT Score and Freshman GPA	465
Figure 37.4	Multiple Scores on Y With a Less Than Perfect Relationship Between X and Y	465
Figure 37.5	Prediction Line for SAT-GPA Data	466
Figure 37.6	Prediction Line as a Series of Mean Y Scores for Various X Scores	467

Figure 37.7	Many Possible Lines for Predicting Y From X	467
Figure 37.8	Linear Conversion of Degrees Celsius to Degrees Fahrenheit	469
Figure 37.9	The Linear Prediction Equation	470
Figure 37.10	Double Prediction Lines When One Variable Is Skewed and the Other Variables Are Normally Distributed	472
Figure 38.1	More Prediction Error for Smaller Correlations	481
Figure 38.2	Normal Distribution of Prediction Errors Around the Prediction Line	482
Figure 38.3	Percentage of Normally Distributed Scores Around $\hat{Y}$	483
Figure 38.4	Range of Predicted Scores on Y Under Different Correlation Values Between X and Y	487
Figure 38.5	Size of the Standard Error of Prediction When the Standard Deviation of Y Is Large and When the Standard Deviation of Y Is Small	488
Figure 39.1	Venn Diagram Demonstrating Predictor Intercorrelation in Multiple Regression	495
Figure 39.2	Venn Diagram Demonstrating a Method for Partitioning Common Variance Among Predictors in Multiple Regression	495
Figure 40.1	Flowchart for Selecting a Descriptive Display or Statistic	511
Figure 40.2	Flowchart for Selecting an Inferential Statistic	514

List of Tables

Table 2.1	Statistics Test Scores and Ranks	13
Table 3.1	Statistics Test Scores	22
Table 3.2	Statistics Test Scores in Descending Order	23
Table 3.3a	Frequency for Statistics Test Scores Containing All Possible Values	23
Table 3.3b	Frequency for Statistics Test Scores Containing Only Values That Were Actually Obtained	24
Table 3.4	Cumulative Frequency for Statistics Test Scores	27
Table 3.5	Relative Frequency for Statistics Test Scores	28
Table 3.6	Cumulative Relative Frequency for Statistics Test Scores	29
Table 3.7	Grouped Frequency for Statistics Scores: 5-Point Intervals	31
Table 3.8	Grouped Frequency for Statistics Scores: 20-Point Intervals	32
Table 3.9	Cumulative Relative Frequency for Statistics Test Scores	35
Table 3.10	Real Limits of Score X	36
Table 3.11	Selected Portion of Cumulative Relative Frequency Table	38
Table 4.1	Ungrouped and Grouped Statistics Test Scores	43
Table 5.1	Real Limits of Pop-Quiz Scores for 90 Students	70
Table 8.1	Bob's Scores on Tests A and B	102
Table 8.2	Bob's Scores on Tests A to C	102
Table 8.3	A Portion of the Normal Curve Table	105
Table 8.4	Raw Scores, Means, and Standard Deviations for Tests A to C	108
Table 11.1	A Portion of the Binomial Table	147
Table 16.1	A Portion of the Normal Curve Table	195
Table 17.1	A Portion of the t Table	207
Table 17.2	A Portion of the Normal Curve Table	207
Table 17.3	A Portion of the t Table	211
Table 18.1	A Portion of the t Table	218
Table 20.1	A Portion of the t Table	245
Table 23.1	A Portion of the t Table	279
Table 25.1	Depression Level After Treatment	299
Table 25.2	Depression Level After Treatment: Deviation Score Setup	304
Table 25.3	Depression Level After Treatment: Raw Score Setup	307
Table 25.4	A Portion of the F Table	310
Table 25.5	ANOVA Summary Table for Effects of Three Different Treatments on Depression Level	311
Table 26.1	Depression Level After Treatment	321
Table 26.2	ANOVA Summary Table for Effects of Three Different Treatments on Depression Level	322
Table 26.3	A Portion of the Studentized Range Table	322
Table 26.4	Tukey Matrix for the Effect of Three Treatments on Depression Level	324
Table 27.1	ANOVA Summary Table for Effects of Three Different Treatments on Depression Level	329
Table 27.2	Depression Level After Treatment: Raw Score Calculation	329

Table 27.3	A Portion of the F Table for Use With the Scheffé Test	333
Table 28.1	Design Diagram for a Two-Way ANOVA	341
Table 28.2	Design Diagram for a Three-Way ANOVA	341
Table 28.3	Group Means for a Two-Way ANOVA With Two Main Effects and No Interaction Effect	343
Table 28.4	Group Means for a Two-Way ANOVA With an Interaction Effect and No Main Effects	346
Table 28.5	Group Means for a Two-Way ANOVA With a Less Obvious Interaction Effect	347
Table 29.1	Design Diagram for a Two-Way ANOVA	355
Table 29.2	Design Table for a Two-Way ANOVA	355
Table 29.3	Data Table for a Two-Way ANOVA of Sleep Deprivation and Memory Training on Word Recall	356
Table 29.4	Row and Column Factors in a Two-Way ANOVA	358
Table 29.5	Data Table for Use With Raw Score Formulas	360
Table 29.6	A Portion of the F Table	363
Table 29.7	Factorial ANOVA Summary Table	364
Table 30.1	Observed and Expected Registrations in Each Political Party in the District	380
Table 30.2	A Portion of the Chi-Square Table	381
Table 31.1	A Portion of the Chi-Square Table	393
Table 33.1	Sample Size Needed per Group to Obtain .80 or .90 Power, Given Various Cohen's d s or Effect Size η s, for a Two-Sample Independent t Test With Equal Sample Sizes and .05 Nondirectional α	417
Table 33.2	Sample Size Needed per Group to Obtain .80 or .90 Power, Given Various Effect Size η s, for an Independent Sample's One-Way ANOVA With Equal Sample Sizes and .05 α	417
Table 33.3	Total Sample Size Needed to Obtain .80 or .90 Power, Given Various Effect Size ϕ s or V s, for a Two-Variable Chi-Square Test of Independence at .05 α and 1 to 5 df	418
Table 35.1	A Portion of the Pearson r Correlation Table	447
Table 39.1	SPSS Multiple Regression Output Showing Predicted Variance	499
Table 39.2	SPSS Multiple Regression Output Showing Coefficients for the Prediction Equation	499
Table 39.3	SPSS Multiple Regression Output Showing ANOVA Summary Table	500

Preface

I did poorly many years ago when I took my first undergraduate statistics course. I never understood what we were doing or why we were doing it. How ironic it is that I then went on to earn a PhD in measurement and that my favorite course to teach is statistics. I believe that the memory of that first dismal undergraduate course, together with graduate coursework in educational and cognitive psychology, has helped me understand what works and what doesn't work for effective statistical instruction.

Given my eye for effective instruction, I was never able to find a statistics textbook that suited me. Many textbooks are mathematically dense, spending too much time deriving formulas and not enough time on the formulas' logic and practical use. Other textbooks go to the opposite extreme, being mere consumer-oriented overviews. As a result of my dissatisfaction with the available textbooks, over the years I developed dozens of pages of handouts, exercises, and mini-lessons, which I provided to students as a supplement to whatever textbook I was using. To my surprise, my students soon told me that my lectures and supplemental materials were all that they needed to understand the subject and that they had given up reading the textbook altogether! Finally, it dawned on me: Why not incorporate my teaching lectures and supplemental materials into a textbook of my own? Hence, this textbook.

I believe that even students with a modest mathematical background (not beyond basic algebra) can master both the mechanics and the underlying logic of hypothesis testing, which is the heart of statistics. This is the type of student I have been unusually successful in reaching, without compromising course content or rigor. Despite the light tone throughout, this textbook is definitely *not* "Stats Lite." Coverage is comprehensive, and students are well prepared for advanced courses.

Content Distinctives

This textbook has been written for undergraduates taking a first course in statistics. Frankly, content doesn't vary much from one introductory textbook to another. Thus, this textbook includes frequency distributions and their shapes; measurement scales; measures of central tendency and dispersion; the normal curve and z scores; hypothesis formation; inferential statistics such as t , F , and χ^2 ; correlation; and bivariate regression. However, in this textbook, the concepts of sampling error, significant differences, and Type I and Type II errors are stressed throughout. Also, effect size and power, often shortchanged in other textbooks, each get substantive treatment. New to this edition is a conceptual overview of multiple regression as a bridge to higher-level courses.

A first course in statistics should focus on univariate inferential hypothesis testing. Without a solid understanding of inferential logic, the student is not prepared to know when to use which statistic or understand the meaning of any statistic he or she does compute. Thus, the primary emphasis in this textbook is on the underlying logic of hypothesis testing. Every topic in the textbook meets this criterion: Is understanding that necessary and helpful to mastering inferential hypothesis testing? Toward that end, the following steps have been taken:

- *Probability theory is minimized:* Many statistics textbooks spend up to one fourth of their content covering probability topics. Although probability theory and integral calculus underlie the curves by which statistical inferences are interpreted, deriving the probabilities is daunting and unnecessary in a social science statistics course. Social science students need only enough probability theory to understand that the values in the inferential charts they consult are theory based. If you teach probability as part of your statistics course, I have included enough theory and practice to make the connection from the bar chart for hit-or-miss binary events to the continuous normal curve formed by connecting the midpoints of those discrete probabilities. If you are nonplussed by probability theory, you may skip the topic and its modules without loss to the textbook's flow or inferential logic.
- *Mathematical proofs are minimized:* Few social science students are prepared to understand that the mean is the vertex of the parabola of squared errors. Few seek to assure themselves that

the raw score, deviation score, and z-score correlation formulas are mathematically equivalent. So why intimidate otherwise capable students? My experience is that such instruction is counterproductive. Therefore, you will find few proofs in this textbook.

- The proper sample size in the denominator for the standard deviation when using a statistic inferentially is $n - 1$. On that, we all agree. Debate arises when the statistic is being used descriptively: Shall we use N ? Or shall we use $n - 1$, as we will have to use that for later inferential instruction anyhow? A survey of textbooks currently on the market shows about a 50:50 split in presentation in this matter. A similar preference split likely exists among instructors. Handheld calculators give the user a choice: One button says $/N$, and the other button says $/n - 1$. IBM® SPSS® on the other hand, does not give the user a choice. The SPSS algorithm uses $n - 1$ as the sample size in the denominator for the standard deviation (and the subsequent z score), whether the statistic is being used descriptively or inferentially. In this edition of the textbook, I have chosen not to take sides in the debate. Thus, while inferential statistics, of course, are calculated using $n - 1$, I present *descriptive standard deviation solutions using N and also using $n - 1$* . Instructors can take their pick. Instructors using SPSS will need to teach the $n - 1$ algorithm if hand-calculated answers are to agree with the software answers.

- After learning about measures of central tendency and dispersion, students are shown the *effects of score transformations* (adding/subtracting and multiplying/dividing) on the mean and standard deviation. This leads to the understanding that area in the normal curve is constant regardless of actual mean and standard deviation values. This understanding is helpful for most statistics students, and it is essential for those who will later study or interpret standardized educational or psychological tests that are scored in T , CEEB (College Entrance Examination Board score), or similar rescaled units. Depending on your goals for the course, you may include or omit this topic and its modules.

- The z score is the bridge to inferential statistics in that the tabled percentage for any given score is based on infinite sampling. Thus, after learning how to compute and interpret individual z scores, students *reconceptualize the z score as a test of a null hypothesis*: “There is no significant difference between this person’s score and the average of the population of scores from which this score came.” Tabled percentages then take on a new meaning—from percentages actually scoring above or below a given score to the probability that a given score actually came from a given population. In this way, the z formula is reconceptualized using inferential logic.

- This textbook emphasizes the *logical similarity* of all experimental test statistics. That is, from z to one-sample t , to two-sample t , to F , to χ^2 , the *numerator always* is the difference between what you got and what you expected to get, and the *denominator always* scales that difference in terms of random dispersion units. Once that logic is grasped, students should be able to interpret just about any inferential statistic they later run into, even one they have not previously learned. This logic is the missing element in so many statistics textbooks. Thus, each time the student learns a new inferential test statistic and computes an example, I ask the same two questions: (1) “What did you expect to get, and did you get it?” (found in the numerator, and the answer is no, probably not), followed by (2) “Is the observed difference from expectation a lot, or is it only a little?” (scaled against the sampling error term in the denominator, then checked in a table). I use the same systematic approach in teaching the sampling distributions themselves, building from a raw score distribution to a sampling distribution of the mean, to a sampling distribution of the difference between the means, to a within-groups mean square. For inferential statistics, the logic is the message.

- Correlation and prediction, even when used inferentially, follow a different logic (scaled over total variance) from traditional hypothesis testing (scaled over error variance). In order not to break the logic that I repeat throughout the modules, *correlation and linear prediction are given their own section* at the completion of the modules on traditional hypothesis testing. However, the correlation and prediction modules are written in such a way that the instructor who prefers to teach this topic immediately after descriptive statistics can do so without an awkward break in the flow.

- An important feature in this textbook is an *introduction to multiple regression and the general linear model (GLM)*. This module follows naturally from bivariate regression at the conclusion of the textbook. The topic is presented and discussed conceptually only; examples are shown in

SPSS but not calculated. The GLM follows the proportional logic discussed above, which is very different from the logic of hypothesis testing using t and F . Because the GLM is nearly always the focus of a second course in statistics, many students struggle in their second course without this conceptual bridge. Instructors in programs requiring that students take a second course in statistics are advised to assign this module for preparatory purposes. Instructors in terminal programs requiring only a single course in statistics may omit this module with no deficit in instruction.

- Throughout, I show the *silliness of selecting a dichotomous, all-or-none decision criterion*. Computers have access to all possible Type I probabilities, not just the two or three probability columns provided in textbook tables. Thus, computers give the exact probability of having made a Type I error and then leave it to the user to determine whether the hypothesis was met, using whatever criterion the user deems appropriate. Journals, too, do not report dichotomous decisions. Rather, they report the incurred error levels and let the readers determine if the study meets their own statistical and practical significance standards. Thus, after calculating the test statistic, I guide students in interpolating the incurred Type I error, which inevitably falls between the tabled levels. This extra instruction, rare in introductory textbooks, makes clear the meaning of the p values that students will inevitably see in computer output and in journal articles. It also leads naturally to a discussion of confidence levels, confidence intervals, and design implications.
- Many statistics textbooks still do not address *effect size*. This is an unfortunate omission because, after spending so much of the course learning how to calculate statistical significance, students tend to go away thinking that statistical significance means that the treatment “worked.” This, of course, is not necessarily so. Thus, once students understand the logic, calculations, and interpretations of hypothesis testing, I discuss effect size. This includes the difference between statistical significance and practical significance, various measures of effect size, and how much of an effect size is “enough.” Moreover, textbooks that do discuss effect size tend to discuss only one measure: Cohen’s d . This is surprising because Cohen’s d can be used only with t tests, not with analysis of variance or with chi-square. This textbook presents several measures of effect size—at least one for each major test statistic.
- I take the same comprehensive approach in discussing *power*. This includes the factors that affect power, where those factors “sit” in the test statistic formulas (numerator vs. denominator), how much power is enough, and the sampling and design flaws that increase or decrease power. I then interrelate alpha, power, and effect size—the push-and-pull effect they have on one another. Students practice estimating power under different conditions of alpha, effect size, and N .
- I also discuss typical journal *publication guidelines* for statistical significance, power, and effect size, so that students can more effectively distinguish a publishable study from one that is not publishable. And I let them know that these are the considerations that the researcher needs to take into account when designing the study, rather than be disappointed at the data analysis stage.
- A significant element in this textbook is *SPSS instruction and output*. Examples are not only solved manually within the textbook narrative but also shown as software output in the new “SPSS Connection” sections. These sections not only show output as the student would see it but also give detailed point-and-click instructions for obtaining the output. Indeed, the instructions are so adequately detailed that instructors may find that no separate SPSS instruction manual is needed. Placement of the output and instructions at the end of the modules allows instructors not teaching SPSS to easily skip these pages.
- In the first edition of this textbook, answers to all exercises were provided to the instructor, who could choose whether or not to provide them to the students. In this third edition, answers to the odd-numbered exercises are provided at the rear of the textbook for the student’s convenience, with answers to the even-numbered exercises reserved for the instructor. The number of exercises also has been doubled, giving students additional practice.
- PowerPoint slides for instructional purposes have been redone to more fully capture key instructional points and to allow for presentation of important figures and diagrams.

Style Distinctives

This textbook is written in short *modules* rather than long modules. Studies of today's students have shown that they would rather read many short modules than fewer long modules. This allows the instructor to assign modules that better match the intended instruction. It also permits students to master a distinct block of information before progressing to the next block.

This is not to say that you can select modules at random. Statistics “builds.” Hence, most modules must be taken in a prescribed order. Still, a few modules can be omitted or their order changed. For example, you may choose to omit the modules on score transformations, probability, one-sample *t* tests, and variations on two-sample *t* tests, and you may even stop short of analysis of variance—all with minimal loss to the overall logic. Finally, I chose to place the correlation modules after completing the *t*, *F*, and χ^2 instruction in order not to interrupt the repetitive inferential logic. However, because some instructors prefer to teach correlation right after descriptive statistics, I intentionally worded the correlation modules so that they can be taught in either location without awkward transitions.

The textbook contains the following additional features:

- Each module begins with a set of *learning objectives* and a list of *terms and symbols* unique to that module. These assist the student by providing both a scaffold for what to expect of that day's reading and a reference for where to find key information in the future. These features also assist you in knowing where key concepts first appear.
- Frequent *Check Yourself!* boxes throughout the text *reinforce learning* soon after key concepts have been taught. Thus, students quickly test their understanding as they progress.
- *Practice exercises are dispersed* throughout the modules as subtopics are covered rather than appended to the end of each module. Also, rather than overload the textbook with exercises, a modest number of exercises appear in the textbook, with additional exercises appearing in the *Instructor Resource Guide* and many more in the *Student Study Guide*. Answers to all the textbook exercises are located in the *Instructor Resource Guide*, thereby allowing the instructor to choose between using textbook exercises for homework and testing, or providing students with the exercise answer file for self-study purposes.
- The tone is informal and *conversational*. I ask questions: “If it were up to you, which of the three formulas would you want to use?” And I confirm students' competence: “By now you are so proficient at this task that you are probably already looking for a table in which to look it up.” The intent is for me to appear to be right beside the student, in the manner of a classroom teacher. Indeed, students have said that I write like I talk.
- I make extensive use of *humor*. Cartoons are dispersed throughout, and quips are placed in the margin sidebars as their associated topics are covered. This is admittedly a bit of form over substance, but it serves as a stress buster (a person cannot be amused and anxious at the same time) and also appeals to the quick-transitions learning style of today's student. In trial runs, even graduate students enjoyed the humor and reported that they skimmed each module for its jokes before reading the module's content.

Final Note From Wendy

It is my sincere hope that you and your students find this book to be, as the title states, *Statistics Alive!*

Sincerely,
Wendy J. Steinberg

Revisions Made to the Third Edition

The third edition was written under different, and unfortunate, circumstances from the first two. As Wendy began work on this edition, she fell ill. Wendy was committed to this text and to revising

it to ensure that it remained one of the most approachable presentations of statistics available. Unfortunately, she was only able to begin her work on this revision before she passed away. I, Matthew Price, have been involved with *Statistics Alive!* since the first edition. I was personally selected by Wendy to review the drafts of each edition and write the companion study guide because she felt that we wrote in a similar voice. Thus, I was asked by SAGE and given the blessing of Wendy's family to edit this third edition. In taking over the revision, I have done my best to improve on what was already an incredible book. I hope that students and instructors will benefit from the updates to this edition.

Here is an overview of the updates that were made to this edition:

- The learning objectives throughout the text have been revised to assist students in identifying relevant topics from each module.
- More attention is given to the rationale and theory behind hypothesis testing. Similarly, the focus on computation by hand has been reduced.
- The discussion of plotting has been revised such that box-and-whisker plots are now discussed and methods of plotting that are no longer in use have been removed.
- The notation throughout the text has been revised to be more consistent with other texts and articles in the area of statistics. This change should allow the information provided in this book to be more easily translated to other resources.
- The discussion of confidence intervals and interval estimates has been expanded throughout the text.

As Wendy closed her prior introduction, I also am open to feedback. If, in using the materials, you find any errors, please let me know. I also welcome any other feedback (compliments or criticism) regarding your experience using this textbook. It would be my pleasure to receive your comments and to incorporate your suggestions in the next edition. You may use the feedback card at the end of this textbook to contact me.

—Matthew Price

Supplemental Material for Use With the Third Edition of *Statistics Alive!*

Student Study Guide

This affordable *Student Study Guide to Accompany Statistics Alive!* (third edition) will help students get the added review and practice they need to improve their skills and master the material for the course. The newly revised study guide is broken down by module and includes the following: summaries, learning objectives, and practice exercises (which consist of computation, true/false, short-answer, and multiple-choice questions).

Student Study Site

An open-access, free student study site provides additional support to students to help them prepare for class and exams. Datasets for SPSS Connection sections are available for additional practice. Variable lists are also available. Access the student study site edge.sagepub.com/steinberg3e.

Instructor's Resource Site

Helpful teaching aids are provided for professors, particularly those new to teaching statistics and to using *Statistics Alive!* (third edition). Included on the password-protected portion of the companion website are the following:

- Teaching tips, with suggested class activities
- Test bank (including computation, true/false, and multiple-choice questions)
- Answers to all test bank questions, textbook exercises, and even-numbered questions
- Instructor's manual, including solutions and activities.
- Accessible PowerPoint® slides for teaching

Visit the SAGE Edge site edge.sagepub.com/steinberg3e.

Acknowledgments

Only an intravenous administration of Valium will get me through this course!

—Comment written by Shelly W. on an
index card that I asked students to
complete on the first day of class

*This book was inspired by all the Shelly Ws who have suffered through statistics courses over the
years. May this textbook forever end your suffering.*

—Wendy J. Steinberg

I would like to thank specific people who contributed to the completion of this third edition.

To my past students: I am deeply sorry for all of those hand calculations, but I stand by how important they are to this topic. Thank you for going through that and working with me on this textbook.

To Bennie Clark Allen, production editor at SAGE Publications for the textbook; to Elizabeth Cruz, editorial assistant, and Tyler Huxtable and Megan O’Heffernan, cartoon copyright acquisitions sleuths; and to the copyediting team at QuADS Prepress—Rajasree Ghosh, Shamila Swamy, and Rajeswari Krithivasan: Thank you for all that you have done! It was a pleasure working with you all.

To Abbie Rickard, acquisitions editor for psychology at SAGE Publications: Thank you for guiding this work through an unknown process with an author who was very new to it all.

And special thanks to Zoe Brier, author of the *Student Study Guide*. You are always a pleasure to work with and have done an exceptional job in taking over the study guide. Thank you for all your help and long nights in getting this done—I could not have left it in more capable hands!

And a very special thanks to Wendy J. Steinberg for bringing this book into the world and introducing so many students to the study of stats! You have made such a difference in this world!

—Matthew Price
Spring 2020

About the Authors

Wendy J. Steinberg entered academia midcareer, having spent the first part of her career in high-stakes test development. She held a PhD in educational psychology with dual concentrations, one in measurement and the other in development and cognition. Teaching was her passion. She viewed education as a sacred task that teachers and students alike should treat with reverence. She wanted this textbook in the hands of every statistics student so that tears will be banished forever from the classroom. A portion of the sale of each textbook goes to charity.

Matthew Price holds a PhD in clinical psychology and has spent his career pursuing two goals. The first is helping victims of trauma, and the second is teaching statistics. From his time in undergraduate statistics, he saw the challenge that this topic posed to many talented students. He has since spent many late nights making heads or tails out of how to teach the probability of heads and tails in an approachable and enjoyable manner. He is honored to assist in writing this textbook.

Module 1. Math Review, Vocabulary, and Symbols

Module 2. Measurement Scales

Math Review, Vocabulary, and Symbols

Terms: case, subject, sample, population, statistic, parameter, variable, constant, summation, PEMDAS

Symbols: $X, X_1, X^2, \bar{X}, M, p, q, N, n, k, \approx, <, >, ||, \Sigma, \sqrt{}$

Identify and define common statistical terms and symbols

Understand the common algebraic rules to manipulate equations

Review the basic arithmetic functions, rules, and procedures necessary to learn statistics

Getting Started

Statistics is a language. As with any language, before we can communicate effectively, we must master that language's vocabulary and symbol system. This first module, then, introduces essential statistical vocabulary and symbols. Additional vocabulary and symbols will be introduced as you progress through the textbook.

Following the vocabulary and symbols, the remainder of this module is a math refresher. A surprising proportion of students need a math review or, at the least, benefit from it. Note that the review is basic. It is intended merely to make sure that everyone using this textbook—including the math phobic and those who have been out of school for some time—remembers the fundamentals.

Nevertheless, do not be misled by this math review into thinking that this is a “remedial math” type of textbook. It is not. This textbook is a comprehensive, college-level treatment of statistical concepts and calculations. While the first few modules may seem elementary, the material quickly builds in difficulty. Still, the mathematics never becomes overwhelming. The difficulty of statistics seldom lies with the mathematics. Rather, the difficulty lies with the inferential logic. Even in its most difficult modules, this textbook requires only elementary algebra. In most units, basic arithmetic is sufficient. You are fully capable of mastering the material.

Your classroom teacher is, of course, the primary instructor for your course. However, we have written this textbook in a conversational tone, as if we were standing alongside you, guiding you in your understanding of statistics. In that sense, we, too, are your teachers. So, as your coteacher, let us offer a few words of advice.

- First, carefully read all the assigned text material. Because statistics is logic driven, missing a piece of the argument can be deadly to your later understanding. Think of this textbook as an integral part of your class instruction.
- Second, take notes. Have a highlighter and a pen nearby when reading the textbook.
- Third, when you come up against a *Check Yourself!* box, don't just skip over it. Rather, check yourself! The boxes are strategically placed throughout the text to ensure mastery of important material before progressing to the next topic.

- Finally, take time to laugh. Despite rumors to the contrary, the study of statistics doesn't have to be boring. We have peppered the text with cartoons and the margins with quips and quotes. Let these be tension breakers, reminding you that even statistics has a lighter side. For example, did you know that one of the statistics used to test experimental hypotheses was originally created to test the quality of beer? Or that the only certain conclusion with any inferential statistic is that we're not certain? Or that from knowledge of the amount of ice cream eaten by members of a community, we can predict the number of drownings in that community quite accurately?

How can these things be? You will have to read on to find out. We hope that you find your study to be . . . *Statistics Alive!*

Common Terms and Symbols in Statistics

To work with statistics, we need to use the notation commonly accepted by those in the field. Most statistical formulas or problems contain one or more of these terms or symbols.

Case. This is an individual unit under study. It could be a person, an animal, a car, a soft drink, a lightbulb, and so on. For example, a study of car-stopping distances might include 300 cases (cars).

Subject or Participant. When the cases are human beings, they are often referred to as subjects or participants. For example, a study of running speed in college sophomores might include 80 subjects (80 sophomore students).

Sample. This is the group of participants in a study. Usually, they are a subset of a larger group. For example, we might measure the height of a sample of 100 elderly women.

Population. This is the larger group of participants about which we want to draw a conclusion. For example, although we may have sampled only 100 elderly women, we might want to draw a conclusion about the height for the whole population of elderly women of which the 100 women were a part.

Statistic. This is a summary number (e.g., an average) for a sample. Our statistical average might be 63 in.

Parameter. This is a summary number (e.g., an average) for a population. Our parametric average might be 63.5 in.

Variable. When the value of the trait being measured varies from case to case, that trait is referred to as a variable. For example, when measuring the running speed of college sophomores, running speed is a variable because each student is expected to run at a different speed.

Constant. When the value of the trait being measured is the same for all cases, that trait is referred to as a constant. In a mechanized (machine-driven) study of car-stopping distance, for example, researchers might figure how long it takes the average person to raise a foot to the brake pedal and then add that value as a constant to the measured stopping distances.

Uppercase Letters. These usually represent variables, scores that vary from case to case. An example would be X , which might stand for each subject's running score. Every subject would have a different X score.

Lowercase Letters. These usually stand for constants, values that are the same for each case. An example would be c , which, in the study of car-stopping distance, stands for the average time it takes to raise a foot to the brake pedal.

Bar Over a Letter. This represents the average of a variable. An example would be “ $\bar{X}$ ” (pronounced “X-bar”), which is the average score on the variable X .

M. This letter is reserved for the mean, known in lay language as the average. Wait. Didn’t we say above that $\bar{X}$ is the symbol for the average? Yes, either symbol, $\bar{X}$ or M , is used to represent an average. However, M is the more recent symbol and is the only symbol now accepted by APA journals. Nevertheless, you should learn to recognize the $\bar{X}$ symbol, as it will appear in older textbooks, in textbooks aimed at students outside the social sciences, and in journal articles published prior to the change.

p. This letter is reserved for the probability of an event occurring. An example would be the probability of rolling a 3 on a standard die. Because a die has six sides, numbered one through six, the probability of rolling any given number is $1/6$. In decimal form, the probability is .167.

q. This letter is reserved for the probability of an event *not* occurring, or in other words, $1 - p$. An example would be the probability of not rolling a 3 on a standard die. The probability is $1 - p$, which is $1 - 1/6$, which is $5/6$. In decimal form, the probability is .833.

N, n. This letter is reserved for the number of cases. An example would be the number of students in your statistics class. If there are 50 people in your class, then $N = 50$. Typically, n is the number of cases in a sample, and N is the number of cases in a population.

k. This letter is often used to refer to the number of groups or categories. For example, if you were doing a study where you were comparing a new treatment for depression (1) to a placebo (2), then $k = 2$, because there were two groups.

Subscripts. Subscripts refer to particular subjects or cases. Continuing with the variable “ X ,” “ X_1 ” would be the score of the first subject, “ X_2 ” the score of the second subject, and so on.

Wavy Parallel Lines ($\approx$). This symbol means *about* or *approximately*. For example, we might say that the coat cost $\approx \$200$ or that you expect to experience a snack attack in ≈ 15 min.

Less Than and Greater Than ($<$ and $>$). These symbols mean *less than* and *greater than*, respectively. For example, we might say that the coat cost $< \$200$ or that you expect > 15 min to pass before you experience a snack attack.

Summation (Σ). This means to add the scores of all cases. For example, ΣX means to add up all scores on the variable X .

Multiplication Indicators. There are several ways to indicate that two numbers are to be multiplied. As a child, you learned to write “ 3×4 .” However, when substituting letters for numerals in algebra, the use of “ $\times$ ” to indicate multiplication becomes confusing. Thus, a second method to indicate multiplication is to put parentheses around each number individually. For example, “ $(3)(4)$ ” means to multiply 3 and 4. Another way is to put a midlevel dot or asterisk between the numbers. For example, “ $3 * 4$ ” means to multiply 3 and 4. A fourth way, which can be used only when the numbers are symbolized by letters, is to write the numbers next to each other without a space. For example, “ ab ” means to multiply a and b .

Reciprocal. A reciprocal is “one divided by the number.” Thus, the reciprocal of 3 is $1/3$, and the reciprocal of 6 is $1/6$.

Superscripted Number or Exponent. This indicates the number of times a number should be multiplied by itself. For example, 3^2 (pronounced “three squared”) means to multiply 3 twice, which is 3×3 , which is 9.

Radical Sign ($\sqrt{}$). This says to find the square root of the number under the radical sign. The square root is the number that when multiplied by itself yields the number under the radical sign. For example, “ $\sqrt{4}$ ” is 2, because 2 times itself is 4. And “ $\sqrt{25}$ ” is 5, because 5 times itself is 25.

Check Yourself!

Which terms in this section were new to you? Reread the definition for any term that you did not already know.

Fundamental Rules and Procedures for Statistics

A few rules and procedures are fundamental to the study of statistics. Again, this review is intended for those whose math skills are rusty.

- *Ordering operations (PEMDAS):* When several arithmetic operations are called for in a formula, perform them in the following order: (1) Complete any operations inside the **P**arentheses first, then (2) all **E**xponents (and remember that square roots are considered exponents), then (3) all **M**ultiplication, then (4) all **D**ivision, then all (5) **A**ddition, and finally (6) **S**ubtraction. You can use the first letter of each operation to create the acronym **PEMDAS**, which stands for parentheses, exponents, multiplication, division, addition, and subtraction. This order is how to proceed with all operations in mathematics and statistics. For example, $4 + 2 \times 3 = 4 + 6 = 10$. But $(4 + 2) \times 3 = 6 \times 3 = 18$. For another example, $2 + 3^2/2 = 2 + (3^2/2) = 2 + (9/2) = 2 + 4.5 = 6.5$. But $(2 + 3^2)/2 = (2 + 9)/2 = 11/2 = 5.5$.

- *Multiplying by a reciprocal:* Dividing by a number is the same as multiplying by its reciprocal. For example, 6 divided by 3 is equal to 6 times $1/3$. Both are equal to 2. Some statistical formulas substitute one function for the other. If you are expecting division and don't see it, check for multiplication by a reciprocal. For example, some formulas multiply by $1/N$ rather than divide by N .

- If the signs of two numbers being multiplied differ, the result will be negative. However, if the signs of the two numbers are the same (whether positive or negative), the result will be positive. You may have learned this rule as “a negative times a negative is a positive.” For example, $(+2)(+4) = +8$, and $(-2)(-4) = +8$. But $(+2)(-4) = -8$, because the signs of the numbers being multiplied differ.

- *Converting between fractions, decimals, and percentages:* To convert a fraction to a decimal, divide the numerator by the denominator. For example, $1/4$ is 1 divided by 4, which, when you carry out the division, is 0.25. To convert a decimal to a fraction, remove the decimal point and place the number over 1 followed by as many zeros as there were original digits. For example, 0.25 is $25/100$, which further reduces to $1/4$. To convert a decimal to a percentage, multiply by 100 and add a percentage sign. (To multiply by 100, move the decimal place two places to the right.) For example, 0.25 is 25%.

- *Deciding on the number of decimal places:* In math, “precision” refers to the exactness of a calculation. In the social sciences, two decimal places are usually considered adequate for a final answer. Answers carried to additional decimal places imply a degree of precision that we rarely have in the social sciences. However, most of the formulas in this textbook require that you perform a series of operations before arriving at a final answer. Rarely will either the final answer or the intervening steps be in whole numbers. If final answers are to be reported to two decimal places, you should carry all the preliminary steps to three decimal places and then round the final answer to two decimal places. Never round to whole numbers during the preliminary steps, as a series of such roundings may significantly alter the final answer.

- *Rounding numbers:* If the last digit is greater than 5, round up by adding one to the preceding digit. If the last digit is less than 5, leave the preceding digit unchanged. For example, 46.268 rounds up to 46.27, and 46.263 rounds down to 46.26. If the last digit is exactly 5, whether or not you round depends on the preceding digit. If the preceding digit is odd, round up by adding one to it; if the preceding digit is even, leave it unchanged. For example, 32.635 rounds to 32.64, but 32.645 also rounds to 32.64.

- *Reordering terms to solve for an unknown:* If the quantity for which we want to solve is not alone on one side of an equation, we must first isolate it. For example, if we have the equation $24 = 16 + a$, we must isolate a on one side or the other of the equal sign. We do this by subtracting 16 from both sides. This gives us $24 - 16 = a$. This, of course, is 8.
- Sometimes, a formula is given to solve for one term, but we instead want to solve for some other term within the formula. Again, we must reorder the terms to isolate the desired term. Take the formula for a z score, which you will come across in Module 8:

$$z = \frac{X - M}{s}$$

- The formula solves for z . But what if we know z and want to instead solve for X ? We have to reorder the terms to isolate X . First, we get rid of the s in the denominator by multiplying both sides of the equation by s . This gives us $zs = X - M$. Then, we add M to both sides of the equation. This gives us $M + zs = X$. If we like, we can flip the terms on both sides of the equation so that X is to the left of the equal sign: $X = M + zs$. Either version is correct.

Check Yourself!

Which rules and procedures in this section were new to you? Reread the explanation for any rule or procedure that you did not already know.

Check Yourself!

Here are two formulas for a variance (a statistic that you will study later in the course):

$$s^2 = \sum (X - M)^2 (1/n) \quad \text{and} \quad s^2 = \frac{\sum (X - M)^2}{n}$$

Explain why both formulas lead to the same answer.



Reprinted with permission of the Carolina Biological Supply Co.

Practice

1. Complete the following chart:

Fraction	Percentage	Decimal
$\approx 33\frac{1}{3}$		
	12.5%	
		0.667

2. Complete the following chart:

Fraction	Percentage	Decimal
$1/6$		
	10%	
		0.143

3. Complete the following chart:

Fraction	Percentage	Decimal
$1/20$		
	95%	
		0.25

4. Complete the following chart:

Fraction	Percentage	Decimal
$1/50$		
	0.99%	
		0.01

5. Round the following numbers to two decimal places:

- a. 26.412 _____
 b. 62.745 _____
 c. 36.846 _____

6. Round the following numbers to two decimal places:

- a. 77.935 _____
 b. 1086.267 _____
 c. 39.633 _____

(Continued)

(Continued)

7. Round the following numbers to two decimal places:

a. 95.555 _____

b. 0.023 _____

c. 48.950 _____

8. Round the following numbers to two decimal places:

a. 110.001 _____

b. 12.635 _____

c. 276.772 _____

9. Solve the following equations applying the rules of order of operations:

a. $4 \times 5 + 3 \times 2 =$ _____

b. $4(5 + 3) \times 2 =$ _____

c. $((4 \times 5) + 3)(2) =$ _____

d. $\sqrt{4(5+3)}(2) =$ _____

e. $4^2(5 + 3)(2) =$ _____

10. Solve the following equations applying the rules of order of operations:

a. $4\sqrt{5^2}(3 \times 2) =$ _____

b. $(4)(5 + 3 \times 2) =$ _____

c. $4\sqrt{5} + (3 \times 2) =$ _____

d. $4^2(5 + 3 \times 2) =$ _____

e. $(4)(5 + 3)2^2 =$ _____

11. Solve the following equations applying the rules of order of operations:

a. $4\sqrt{36} \times (2 \times 3) + 2 =$

b. $6^2/12 - 4/2 =$

c. $\sqrt{8-4} \times 6/2 - 1 =$

d. $[(30)(0.5) - 6] \times 2 + 4 =$

e. $5 + 3 \times 7 - 4 =$

12. Solve the following equations applying the rules of order of operations:

a. $\sqrt{5 \times 6 + 6 - 3^2} - 2 =$

b. $(4 + 7^2 - 17)/6 =$

c. $6 - 2 \times 2 + 9/3 =$

d. $100 - 10 \times 5 - 5 =$

e. $\sqrt{(12-8) \times 16} - 2 =$

13. Reexpress the following equations, substituting reciprocals for division. Then, solve in decimal form.

Equation	Reexpressed in Reciprocals	Solution
$16/5 - 0.246 =$	$\underline{\hspace{2cm}} =$	
$68 + 68/3 =$	$\underline{\hspace{2cm}} =$	
$2/3 - 1/5 =$	$\underline{\hspace{2cm}} =$	

14. Reexpress the following equations, substituting reciprocals for division. Then, solve in decimal form.

Equation	Reexpressed in Reciprocals	Solution
$98 - 75/3 =$	$\underline{\hspace{2cm}} =$	
$672 + 2/6 =$	$\underline{\hspace{2cm}} =$	
$8/3 + 5/4 =$	$\underline{\hspace{2cm}} =$	

15. Reexpress the following equations, substituting reciprocals for division. Then, solve in decimal form.

Equation	Reexpressed in Reciprocals	Solution
$6/3 - 0.95 =$	$\underline{\hspace{2cm}} =$	
$12.50 - 8/2 =$	$\underline{\hspace{2cm}} =$	
$25/5 + 6 =$	$\underline{\hspace{2cm}} =$	

16. Reexpress the following equations, substituting reciprocals for division. Then, solve in decimal form.

Equation	Reexpressed in Reciprocals	Solution
$155 + 5/2 =$	$\underline{\hspace{2cm}} =$	
$3.28 - 3/2 =$	$\underline{\hspace{2cm}} =$	
$720/9 - 45 =$	$\underline{\hspace{2cm}} =$	

17. Rearrange the equation to solve for the indicated unknown.

- a. $64/4 + b = 30$ Solve for b : $\underline{\hspace{2cm}}$
 b. $0.30c = 20$ Solve for c : $\underline{\hspace{2cm}}$
 c. $Y = bX + a$ Solve for a : $\underline{\hspace{2cm}}$

18. Rearrange the equation to solve for the indicated unknown.

- a. $b + 20 = 0.5$ Solve for b : $\underline{\hspace{2cm}}$
 b. $F = 1.8c + 32$ Solve for c : $\underline{\hspace{2cm}}$
 c. $50 = 60/a$ Solve for a : $\underline{\hspace{2cm}}$

(Continued)

(Continued)

19. Rearrange the equation to solve for the indicated unknown.

a. $b + 24/8 = 3$ Solve for b : _____

b. $T = 12c - 4$ Solve for c : _____

c. $12 = 3a$ Solve for a : _____

20. Rearrange the equation to solve for the indicated unknown.

a. $7 = 0.5b + 2.5$ Solve for b : _____

b. $64 = 4c$ Solve for c : _____

c. $3a = 0.25V$ Solve for a : _____

More Rules and Procedures

In this textbook, derivations and proofs are kept to a minimum. However, a fuller understanding of the formulas and of the relationship between various statistics requires mathematical manipulations beyond the basic rules just presented. If your instructor intends to derive one formula from another (say, a raw score formula from a deviation score formula) or prove theorems (say, the binomial theorem), some additional mathematical rules are necessary. Knowing these rules is especially appropriate for students intending to take additional courses in statistics. Let your instructor be your guide on the need for this section.

Rules for summation across and within parentheses

1. The summation of a constant is N times the constant.

$$\Sigma(a) = Na$$

For example, if $a = 6.24$ and there are 3 subjects, then

$$\Sigma(a) = Na = 3 \times 6.24 = 18.72$$

2. The summation of a constant times a variable is the constant times the sum of the variable.

$$\Sigma(aX) = a\Sigma X$$

For example, if $a = 6.24$ and if $X_1 = 2$, $X_2 = 4$, and $X_3 = 6$, then

$$\Sigma(aX) = a\Sigma X = 6.24 \times (2 + 4 + 6) = 6.24 \times 12 = 74.88$$

3. The summation of two or more terms within parentheses is the same as their independent summation.

$$\Sigma(X + Y) = \Sigma X + \Sigma Y$$

$$\Sigma(X + a) = \Sigma X + Na$$

$$\Sigma(X + aY) = \Sigma X + a\Sigma Y$$

For example, if $X_1 = 2$, $X_2 = 4$, and $X_3 = 6$; if $Y_1 = 1$, $Y_2 = 3$, and $Y_3 = 5$; and if $a = 6.24$, then

$$\Sigma(X + Y) = \Sigma X + \Sigma Y = (2 + 4 + 6) + (1 + 3 + 5) = 12 + 9 = 21$$

$$\Sigma(X + a) = \Sigma X + Na = (2 + 4 + 6) + (3 \times 6.24) = 12 + 18.72 = 30.72$$

$$\Sigma(X + aY) = \Sigma X + a\Sigma Y = (2 + 4 + 6) + (6.24)(1 + 3 + 5) = 12 + (6.24)(9) = 12 + 56.16 = 68.16$$

4. An exponent applied to a product within parentheses can be distributed to each term within the parentheses. For example, $(ab)^2 = (ab)(ab) = aa \times bb = a^2 \times b^2$

$$\text{Thus, } (ab)^2 = a^2 \times b^2$$

$$\text{For example, } (4 \times 3)^2 = 4^2 \times 3^2 = 16 \times 9 = 144.$$

5. When there is a binomial in the parentheses, there are rules for expanding the binomial, such as $(a + b)^2 = aa + ab + ab + bb = a^2 + 2ab + b^2 = a^2 + b^2 + 2ab$

Thus, $(a + b)^2 = a^2 + b^2 + 2ab$

For example, $(4 + 3)^2 = 4^2 + 3^2 + 2(4 \times 3) = 16 + 9 + 2(12) = 25 + 24 = 49$

$(a - b)^2 = aa - ab - ab + bb = a^2 - 2ab + b^2 = a^2 + b^2 - 2ab$

Thus, $(a - b)^2 = a^2 + b^2 - 2ab$

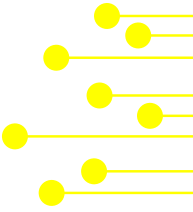
For example, $(4 - 3)^2 = 4^2 + 3^2 - 2(4 \times 3) = 16 + 9 - 2(12) = 25 - 24 = 1$

Practice

21. Expand the following expressions using the rules for summation within parentheses and for binomial expansion (*Note:* Numerals are always constants; thus, N is always a constant):
- a. $\Sigma(X + Y)^2$ _____
 - b. $\Sigma(bXY + 1)^2$ _____
 - c. $\Sigma(XY + 1)$ _____
22. Expand the following expressions, using the rules for summation within parentheses and for binomial expansion (*Note:* Numerals are always constants; thus, N is always a constant):
- a. $\Sigma(1 + N)$ _____
 - b. $\Sigma(2X + Y)$ _____
 - c. $\Sigma(aX)^2$ _____

That should be all the math you need in this textbook. Of course, higher math is necessary for the fullest understanding of statistical formulas and for further study in statistics. However, a first course in statistics rarely goes beyond what is presented here.

Visit the study site at edge.sagepub.com/steinberg3e for SPSS datasets and variable lists.



Terms: nominal, ordinal, interval, ratio, continuous, discrete, real limits

Symbols: LL, UL

Define a scale of measurement

Classify data according to their scale of measurement

Distinguish between discrete and continuous variables

Establish real limits for continuously scored data

What Is Measurement?

Measurement is the process of assigning numbers to the observations on some variable. However, not all numbers have the same numeric properties. For example, one football player may wear a jersey with No. 7 on it, and a second player may wear a jersey with No. 18 on it. However, you would not say that the second player has 11 points more of any trait than the first player. The jersey numbers simply don't mean that. Or you may earn the top score on a statistics test, your friend Lauren the second highest score, and your friend Sidney the third highest score. However, you cannot assume that the score difference between you and Lauren is the same as the score difference between Lauren and Sidney. Again, the rankings simply don't mean that. Or one person may score 75 on an intelligence test, and a second person may score 150 on the same intelligence test. While we would agree that the second person is probably considerably more intelligent than the first one, we cannot say that the second person is twice as intelligent as the first person. Again, the scores simply don't mean that.

Scales of Measurement

What, then, can we say about each of these situations? What we can say depends on the scale on which the data were measured. Stevens (1946) suggested that variables are measured on one of four scales: nominal, ordinal, interval, or ratio. Each of these scales allows us to draw different conclusions about the meaning of subjects' scores. Furthermore, the statistics that can be calculated on the data are partially dependent on the measurement scale in the data. Thus, before we can select and compute statistics, we need to know the scale on which the data were measured.

Nominal Scale

A **nominal** scale classifies cases into categories. For that reason, it is also sometimes called a categorical scale. Here are some examples:

m = *male*, f = *female*
 1 = *married*, 2 = *divorced*, 3 = *separated*, 4 = *never married*
 tel = *owns a telephone*, notel = *does not own a telephone*

The nominal scale is the lowest level of measurement. It does not measure how much of a measured trait a person possesses. It merely categorizes the person as a “this” or a “that.” Even when numbers are used to represent the categories, they are designations only, much like social security numbers or the numbers on a football jersey. Because the designations have no numeric meaning, it makes no sense to perform arithmetic operations on them. For example, you would learn nothing by knowing the average telephone number for the members of your statistics class or the average jersey number for the members of a football team.

Most of the statistics that you will use throughout this textbook require that the numbers have numeric meaning. Thus, most of the procedures that you will learn will not be appropriate for data that are measured on only a nominal level. However, certain statistical procedures have been developed to analyze data in a nominal form. For example, the chi-square statistic, which you will learn about later in the course, is computed on nominal data.

Ordinal Scale

An **ordinal** scale ranks people according to the degree to which they possess some measured trait. Persons are first measured on some attribute (e.g., height). Then, they are assigned ranks according to how much of the attribute they possess. Here are some examples:

1 = *tallest*, 2 = *second tallest*, 3 = *third tallest*, and so on.
 1 = *highest GPA* (grade point average), 2 = *second highest GPA*, 3 = *third highest GPA*, and so on.
 1 = *fastest runner*, 2 = *second fastest runner*, 3 = *third fastest runner*, and so on.

To create ranks, place the scores in order by value (usually descending) and then assign the highest score a rank of 1, the second highest score a rank of 2, and so on. Table 2.1 presents scores and ranks for 10 students on a 100-item statistics exam given as a pretest before the course began.

With an ordinal scale, the difference in test scores between two adjacent ranks may not be the same as the difference in test scores between any other two adjacent ranks. Compare the test scores between adjacent ranks in Table 2.1. The test score difference between the first two ranks is 1 point

Table 2.1 Statistics Test Scores and Ranks	
Score	Rank
73	1
72	2
60	3
59	4
57	5
56	6
52	7
48	8
36	9
28	10

(73 – 72), but between the next two ranks, the difference in test scores is 12 points (72 – 60). Clearly, certain aspects of relative performance are lost when scores are expressed as ranks.

Most of the statistics used throughout this textbook require a scale in which adjacent intervals on the measurement scale imply equal intervals between adjacent scores throughout the scale. Because this is not the case for ordinal data, most of the statistics will not be appropriate for data that are measured on an ordinal scale. However, ranks do tell us more than nominal classifications. A higher-ranked person does have more of the measured trait, even if distance on the underlying scores is inconsistent. Therefore, certain statistical procedures have been developed to analyze data when they are in ordinal form. For example, a Spearman rho coefficient measures the degree of relationship between two sets of ranks.

Interval Scale

With an **interval** scale, the distances between adjacent scores are equal and consistent throughout the scale. Equal intervals on the scale imply equal amounts of the variable being measured. For this reason, the interval scale is sometimes referred to as the equal-interval scale. Here are some examples:

Scores on the final exam in this course

Scores on an intelligence test

Degrees Fahrenheit (°F) or Celsius (°C)

Scores on certain personality or career interest tests

Because interval scales are consistent throughout the scale, it makes sense to compare scores by adding or subtracting them. For example, an intelligence score of 120 is 10 points higher than an intelligence score of 110, just as an intelligence score of 60 is 10 points higher than an intelligence score of 50. In either case, the higher-scoring person has 10 more points of scaled intelligence than the lower-scoring person.

However, interval scales have no absolute zero point—the point at which a person would have none of the measured attribute. That is, even if a person scored 0 on an intelligence test, we would not say that the person has no intelligence. This is because the test's starting point is fixed and arbitrary, whereas the starting point of the actual trait itself—in this case, intelligence—is unknown. This is typically the case in psychological or educational measurement. Similarly, 0 °F or 0 °C does not indicate a lack of temperature—it refers to a low temperature.

Because of the lack of an absolute zero point in an interval measurement scale, it makes no sense to compare interval scores by multiplying or dividing them. Although a person who scores 120 on an intelligence test has scored twice as high as the person who scores 60, the person who scores 120 does not have twice as much intelligence as the person who scores 60.

Most of the statistics that you will learn throughout this textbook require at least an interval measurement scale. Fortunately, most data in social science research are also interval scaled. Therefore, the statistics that you will learn are ideally suited for most educational and psychological data.

One controversy in educational and psychological research is whether personality and career interest tests are interval scaled or only ordinal scaled. This is because individual test items typically ask test takers to use a noninterval scale. For example, test takers might be asked to rate the frequency with which they feel that “life has no meaning” with the following scale: *never*, *occasionally*, *often*, or *always*. Or the scale might ask test takers to rate the degree to which they agree with the statement “I enjoy working out-of-doors” with the following scale: *strongly agree*, *slightly agree*, *neither agree nor disagree*, *slightly disagree*, or *strongly disagree*. We cannot say that the difference in frequency between *never* and *occasionally* is the same as the difference in frequency between *often* and *always*. Similarly, we cannot say that the difference in degree between *strongly agree* and *slightly agree* is the same as the difference in degree between *slightly disagree* and *strongly disagree*. Therefore, the scale appears to be ordinal.

On the other hand, at the completion of the personality or career interest test, the test taker receives a score, much like the score you might receive on a test in this statistics course. Moreover, the scores are typically normed against a large number of test takers. This gives any one person's score both position and distance when compared with the known distribution of all scores. These features are typical of an interval scale, not an ordinal scale. Because the scores contain features of both ordinal and interval scales, the debate over the appropriate level of measurement, and hence the appropriate statistics to use to describe the scores, is ongoing.

Ratio Scale

A **ratio** scale is like an interval scale, in that the distance between adjacent scores is equal throughout the distribution. However, unlike an interval scale, in a ratio scale there is an absolute zero point. That is, there is a point at which a person does not have any of the measured traits. Because the trait's starting point is known, the scale reflects that zero point.

Ratio measurement typically applies to measures in the physical sciences. Here are some examples:

- Height
- Weight
- Distance
- Time
- Temperature measured in degrees Kelvin

Because there is an absolute zero point, it makes sense to compare scores by multiplying or dividing them. For example, a person who weighs 110 lb not only weighs 55 lb more than a person who weighs 55 lb but also is twice as heavy. A person who makes a standing broad jump of 3 ft not only jumps 3 ft less than one who jumps 6 ft but also jumps only half as far. On the temperature scale of Kelvin, 0 degrees means an absolute lack of energy—no heat at all!

Although most data in psychology and education are interval scaled, some are ratio scaled. For example, measures of reaction time or physical performance are ratio scaled. Any statistic that is appropriate for interval-scaled data is also appropriate for ratio-scaled data.

Finally, it is possible to measure data on more than one scale. That is, higher-level data can be measured on a lower-level scale. For example, your interval score on a statistics test can be nominally categorized as either pass or fail. However, lower-level data cannot be measured on a higher-level scale. For example, your sex (male or female) cannot be ranked. It can only be nominally categorized.

Check Yourself!

Give examples of nominal, ordinal, interval, and ratio data. Then, convert your ratio data example into interval, ordinal, and nominal scales.

Practice

1. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Feet of snow	
Brands of carbonated soft drinks	
Class rank at graduation	
GPA	
Speed of a baseball pitch	

(Continued)

(Continued)

2. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Eye color	
Time taken to solve a puzzle	
Genre of favorite television program	
Level of depression	
Position in a starting lineup	
Number of angels that can fit on the head of a pin (Oops . . . angels may not be subject to earthly measurement constraints, so you may skip this one.)	

3. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Hair length	
Species of tree	
Military rank	
Political party membership	
Yearly income	

4. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Favorite radio station	
How much sleep you had last night	
How many calories you ate today	
Athletic ability	
Position in a graduation procession	

5. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Literature genres	
Relative academic position in one's graduating class	

Measure	Highest Scale
Hair length	
Job salary	
Self-confidence	

6. Indicate the first letter (N, O, I, R) of the *highest* possible scale for each of the following measures, where N is the lowest and R is the highest:

Measure	Highest Scale
Type of disease	
Business profit	
Pass or fail status	
Conscientiousness	
Car brands	

Continuous Versus Discrete Variables

Continuous variables are variables whose values theoretically could fall anywhere between adjacent scale units. The data that are measured on a ratio scale are always continuously scored. A person's height or weight or even the time a person spends talking on the phone can fall anywhere between the scale units.

Some interval scale data are continuous variables. For example, although people do not score fractional points on IQ tests, depression inventories, or measures of talkativeness, they theoretically could score fractional points if the tests were so graded.

Discrete data, on the other hand, are values that cannot even theoretically fall between adjacent scale units. Some interval scale data are discrete. Examples are the number of blue ribbons won, number of children in a family, or number of photographs taken. With discrete interval data, we can perform all the arithmetic operations on the data that we would with any other interval-scaled set of scores. That is, we can speak of twice as many blue ribbons won, the average number of children in families, or half as many photographs taken. At the same time, we recognize that individual scores cannot fall between the scale units. That is, no single person can earn a partial ribbon, have a partial child, or take a partial photograph.

Check Yourself!

Is the number of courses students register for in a semester a discrete variable or a continuous variable? Is the GPA that students earn in those courses a discrete variable or a continuous variable?

Practice

7. Are the following variables discrete or continuous? Mark “D” or “C” to indicate your answer:

Measure	Variable Type
Inches of rainfall	
GPA	
Speed of a baseball pitch	
Time taken to solve a puzzle	
Level of depression	
Number of angels that can fit on the head of a pin (Oops . . . how'd that get in here again? You may skip this one.)	

8. Are the following variables discrete or continuous? Mark “D” or “C” to indicate your answer:

Measure	Variable Type
Hair length	
Yearly income	
How much sleep you had last night	
How many calories you ate today	
Athletic ability	

9. Are the following variables discrete or continuous? Mark “D” or “C” to indicate your answer:

Measure	Variable Type
Number of TVs in the household	
Level of extraversion	
Color saturation level	
How tired you feel right now	
How many awards you won as a child	

10. Are the following variables discrete or continuous? Mark “D” or “C” to indicate your answer:

Measure	Variable Type
Amount of irritation you felt today	
Number of times you felt irritated today	
How many different ice cream flavors you have tasted	
How much weight you have gained or lost in the past year	
How badly you need a vacation right now	

Real Limits

For continuous variables, scores can theoretically fall anywhere between scale units, even if actual scores fall only at specified locations along the scale. In other words, even though a continuous scale uses the units of 1, 2, and 3, it is theoretically possible to have a score between 1 and 2. As an example, imagine that a syllabus for a Statistics course says that students who have a mean (which is a statistics way of saying average) score above a 90 on three tests, whose possible scores range from 0 to 100, will receive an A. David earns an 85, 86, and 98 on those three tests, giving him a mean test score of approximately 89.67 (we will cover how to compute a mean in Module 5). His mean falls right between 89 and 90! So should David receive an A? You might be inclined to say “No” because 89.67 is less than 90. However, because this scale is continuous, it is important to consider the score’s **real limits**.

On a continuous scale, each score point (e.g., 90) actually refers to a range of possible scores that include all of the values from 90 to the adjacent scores. In our current example, the adjacent scores are 89 and 91. Real limits are set by using half the scale’s unit. Thus, the real limits of any particular score are the score plus or minus (symbolized as $\pm$) one-half scale unit. In our current example, the test scores were scaled by 1 point, so the real limit for each point is ± 0.5 . That means the cutoff of 90 actually corresponds to all the values that could fall between 89.5 and 90.5. The lower real limit (**LL**) is one-half unit below the score (in this case, 89.5), and the upper real limit (**UL**) is one-half unit above the score (in this case, 90.5). No scores actually fall at the real limits. Returning to our example, the score of 90 on a continuous scale suggests that the lowest possible mean on those three tests necessary to get an A is 89.5, which indicates that David’s mean grade of 89.67 is enough to earn an A. Depending on how grades are assigned in your class, this information may be very helpful later on when talking to your instructor about grades.

Check Yourself!

For a test measured on a scale of 10 to 100 in 10-point intervals, what is the value of the real limit? What, then, are the real limits (**LL** and **UL**) for a score of 70?

Real limits will be important in calculating medians and percentile ranks from tabulated data and for constructing histograms for continuous data. You will learn about these in the next few modules.

Practice

11. What are the real limits of the following scores?

Score	Size of Each Scale Interval	Real Limits
IQ = 116	1 point	_____ and _____
Height = 66.5	1/2 in.	_____ and _____
Age = 20	10 years	_____ and _____
Driving speed = 60	5 mph	_____ and _____
Olympic performance = 9.7	1/10 point	_____ and _____

(Continued)

(Continued)

12. What are the real limits of the following scores?

Score	Size of Each Scale Interval	Real Limits
GPA = 3.2	1/10 point	_____ and _____
Test score = 83	1 point	_____ and _____
Miles to work = 20	10 miles	_____ and _____
Feet of snow = 3.25	1/4 ft	_____ and _____
Shoe size = 11	Whole size	_____ and _____

13. What are the real limits of the following scores?

Score	Size of Each Scale Interval	Real Limits
GPA = 3.20	1/100 point	_____ and _____
Test score = 83.4	1/10 point	_____ and _____
Miles to work = 20	5 miles	_____ and _____
Feet of snow = 3	Whole foot	_____ and _____
Shoe size = 11.0	Half size	_____ and _____

14. What are the real limits of the following scores?

Score	Size of Each Scale Interval	Real Limits
IQ = 116.0	1/2 point	_____ and _____
Height = 66.5	1/10 in.	_____ and _____
Age = 20	Whole years	_____ and _____
Driving speed = 60	10 mph	_____ and _____
Olympic performance = 9.70	1/100 point	_____ and _____

Visit the study site at edge.sagepub.com/steinberg3e for SPSS datasets and variable lists.