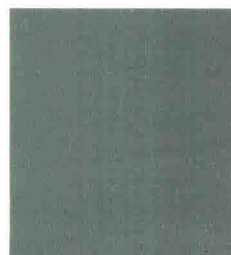
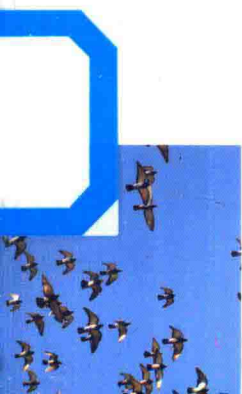
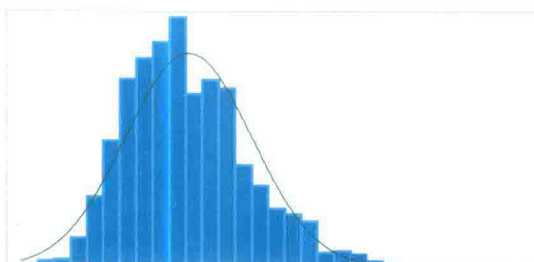
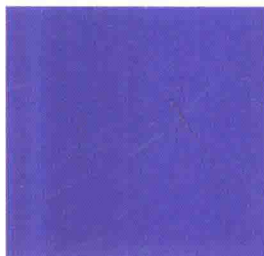
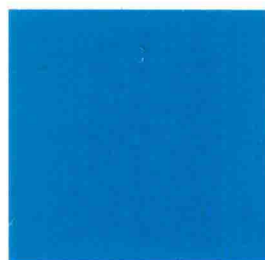
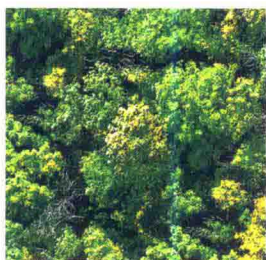
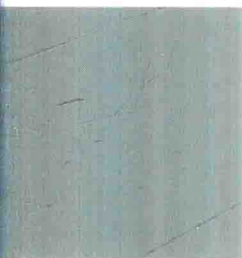


Joseph M. Hilbe

MODELING COUNT DATA



MODELING COUNT DATA

JOSEPH M. HILBE

Arizona State University
and
Jet Propulsion Laboratory,
California Institute of Technology



CAMBRIDGE
UNIVERSITY PRESS

CAMBRIDGE
UNIVERSITY PRESS

32 Avenue of the Americas, New York, NY 10013-2473, USA

Cambridge University Press is part of the University of Cambridge.

It furthers the University's mission by disseminating knowledge in the pursuit of education, learning, and research at the highest international levels of excellence.

www.cambridge.org

Information on this title: www.cambridge.org/9781107611252

© Joseph M. Hilbe 2014

This publication is in copyright. Subject to statutory exception and to the provisions of relevant collective licensing agreements, no reproduction of any part may take place without the written permission of Cambridge University Press.

First published 2014

Printed in the United States of America

A catalog record for this publication is available from the British Library.

ISBN 978-1-107-02833-3 Hardback

ISBN 978-1-107-61125-2 Paperback

Additional resources for this publication at www.cambridge.org/9781107611252

Cambridge University Press has no responsibility for the persistence or accuracy of URLs for external or third-party Internet web sites referred to in this publication and does not guarantee that any content on such web sites is, or will remain, accurate or appropriate.

MODELING COUNT DATA

This definitive entry-level text, authored by a leading statistician in the field, offers clear and concise guidelines on how to select, construct, interpret, and evaluate count data. Written for researchers with little or no background in advanced statistics, the book presents treatments of all major models, using numerous tables, insets, and detailed modeling suggestions. It begins by demonstrating the fundamentals of modeling count data, including a thorough presentation of the Poisson model. It then works up to an analysis of the problem of overdispersion and of the negative binomial model, and finally to the many variations that can be made to the base count models. Examples in Stata, R, and SAS code enable readers to adapt models for their own purposes, making the text an ideal resource for researchers working in health, ecology, econometrics, transportation, and other fields.

Joseph M. Hilbe is a solar system ambassador with NASA's Jet Propulsion Laboratory, California Institute of Technology; an adjunct professor of statistics at Arizona State University; an emeritus professor at the University of Hawaii; and an instructor for Statistics.com, a web-based continuing-education program in statistics. He is currently president of the International Astrostatistics Association, and he is an elected Fellow of the American Statistical Association, for which he is the current chair of the section on Statistics in Sports. Author of several leading texts on statistical modeling, Hilbe also serves as the coordinating editor for the Cambridge University Press series Predictive Analytics in Action.

Other Statistics Books by Joseph M. Hilbe

Generalized Linear Models and Extensions (2001, 2007, 2013 – with J. Hardin)

Generalized Estimating Equations (2002, 2013 – with J. Hardin)

Negative Binomial Regression (2007, 2011)

Logistic Regression Models (2009)

Solutions Manual for Logistic Regression Models (2009)

R for Stata Users (2010 – with R. Muenchen)

Methods of Statistical Model Estimation (2013 – with A. Robinson)

A Beginner's Guide to GLM and GLMM with R: A Frequentist and Bayesian Perspective for Ecologists (2013 – with A. Zuur and E. Ieno)

Quasi-Least Squares Regression (2014 – with J. Shults)

Practical Predictive Analytics and Decisioning Systems for Medicine (2014 – with L. Miner, P. Bolding, M. Goldstein, T. Hill, R. Nisbit, N. Walton, and G. Miner)

Preface

Modeling *Count Data* is written for the practicing researcher who has a reason to analyze and draw sound conclusions from modeling count data. More specifically, it is written for an analyst who needs to construct a count response model but is not sure how to proceed.

A count response model is a statistical model for which the dependent, or response, variable is a count. A count is understood as a nonnegative discrete integer ranging from zero to some specified greater number. This book aims to be a clear and understandable guide to the following points:

- How to recognize the characteristics of count data
- Understanding the assumptions on which a count model is based
- Determining whether data violate these assumptions (e.g., overdispersion), why this is so, and what can be done about it
- Selecting the most appropriate model for the data to be analyzed
- Constructing a well-fitted model
- Interpreting model parameters and associated statistics
- Predicting counts, rate ratios, and probabilities based on a model
- Evaluating the goodness-of-fit for each model discussed

There is indeed a lot to consider when selecting the best-fitted model for your data. I will do my best in these pages to clarify the foremost concepts and problems unique to modeling counts. If you follow along carefully, you should have a good overview of the subject and a basic working knowledge needed for constructing an appropriate model for your study data. I focus on understanding the nature of the most commonly used count models and

on the problem of dealing with both over- and underdispersion, as well as on Poisson and negative binomial regression and their many variations. However, I also introduce several other count models that have not had much use in research because of the unavailability of commercial software for their estimation. In particular, I also discuss models such as the Poisson inverse Gaussian, generalized Poisson, varieties of three-parameter negative binomial, exact Poisson, and several other count models that will provide analysts with an expanded ability to better model the data at hand. Stata and/or R software and guidelines are provided for all of the models discussed in the text.

I am supposing that most people who will use this book start with little to no background in modeling count response data, although readers are expected to have a working knowledge of a major statistical software package, as well as a basic understanding of statistical regression. I provide an overview of maximum likelihood and iterative reweighted least squares (IRLS) regression in Sections 1.4.2 and 1.4.3, which assume an elementary understanding of calculus, but I consider these two sections as optional to our discussion. They are provided for those who are interested in how the majority of models we discuss are estimated. I recommend that you read these sections, even if you do not have the requisite mathematical background. I have attempted to present the material so that it will still be understood. Various terms are explained in these sections that will be used throughout the text.

Seasoned statisticians can also learn new material from the text, but I have specifically written it for researchers or analysts, as well as students at the upper-division to graduate levels, who want an entry-level book that focuses on the practical aspects of count modeling. The book is also addressed to statistical and predictive analytics consultants who find themselves faced with a project involving the modeling of count data, as well as to anyone with an interest in this class of statistical models. It is written in guidebook form, with lots of bullet points, tables, and complete statistical programming code for all examples discussed in the book.

Many readers of this book may be acquainted with my text *Negative Binomial Regression* (Cambridge University Press), which was first published in 2007. A substantially enhanced second edition was published in 2011. That text addresses nearly every count model for which there existed major statistical software support at the time of the book's publication. *Negative Binomial Regression* was primarily written for those who wish to understand the mathematics behind the models as well as the specifics and applications of each

model. I recommend it for those who wish to go beyond the discussions found in *Modeling Count Data*.

I primarily use two statistical software packages to demonstrate examples of the count models discussed in the book. First, the Stata 13 statistical package (<http://www.stata.com>) is used throughout the text to display example model output. I show both Stata code and output for most of the modeling examples. I also provide R code (www.r-project.org) in the text that replicates, as far as possible, the Stata output. R output is also given when helpful. There are also times when no current Stata code exists for the modeling of a particular procedure. In such cases, R is used. SAS code for a number of the models discussed in the book is provided in the Appendix. SAS does not support many of the statistical functions and tests discussed later in the book, but its count-modeling capability is growing each year. I will advise readers on the book's web site as software for count models is developed for these packages. I should mention that I have used Stat/Transfer 12 (2013, Circle Systems) when converting data between statistical software packages. The user is able to convert between 37 different file formats, including those used in this book. It is a very helpful tool for those who must use more than one statistical or spreadsheet file.

Many of the Stata statistical models discussed in the text are offered as a standard part of the commercial package. Users have also contributed count model "commands" for the use of the greater Stata community. Developers of the user-authored commands used in the book are acknowledged at the first use of the software. James Hardin and I have both authored and coauthored a number of the more advanced count models found in the book. Many derive from our 2012 text *Generalized Linear Models and Extensions, 3rd edition* (Stata Press; Chapman & Hall/CRC). Several others in the book are based on commands we developed in 2013 for journal article publications. I should also mention that we also coauthored the current version of Stata's **glm** command (2001), although Stata has subsequently enhanced various options over the past 12 years as new versions of Stata were released. Several of the R functions and scripts used in the book were coauthored by Andrew Robinson and me for use in our book (Hilbe and Robinson 2013). Data sets and functions for this book, as well as for Hilbe (2011), are available in the **COUNT** package, which may be downloaded from any **CRAN** mirror site. I also recommend installing **msme** (Hilbe and Robinson), also available on **CRAN**. I have also posted all of my user-authored Stata commands and functions, as well as all data sets used in the book, on the book's web site at the following

address: http://works.bepress.com/joseph_hilbe/. The book's page with Cambridge University Press is at www.cambridge.org/9781107611252.

The data files used for examples in the book are real data. The **rwm1984** and **medpar** data sets are used extensively throughout the book. Other data sets used include **titanic**, **heart**, **azcabgptca**, **smoking**, **fishing**, **fasttrakg**, **rwm5yr**, **nuts**, and **azprocedure**. The data are defined where first used. The **medpar**, **rwm5yr**, and **titanic** data are used more than other data in the book. The **medpar** data are from the 1991 Arizona Medicare files for two diagnostic groups related to cardiovascular procedures. I prepared **medpar** in 1993 for use in workshops I gave at the time. The **rwm5yr** data consist of 19,609 observations from the German Health Reform data covering the five-year period of 1984–1988. Not all patients were in the study for all five years. The count response is the number of visits made by a patient to the doctor during that calendar year. The **rwm1984** data were created from **rwm5yr**, with only data from 1984 included – one patient, one observation. The well-known **titanic** data set is from the 1912 *Titanic* ship disaster survival data. It is in grouped format with *survived* as the response. The predictors are *age* (adult vs. child), *gender* (male vs. female), and *class* (1st-, 2nd-, and 3rd-class passengers). Crew members have been excluded.

I advise the reader that there are parts of Chapter 3 that use or adapt text from the first edition of *Negative Binomial Regression* (Hilbe 2007a), which is now out of print, as it was superseded by Hilbe (2011). Chapter 2 incorporates two tables that were also used in the first edition. I received very good feedback regarding these sections and found no reason to change them for this book. Now that the original book is out of print, these sections would be otherwise lost.

I wish to acknowledge five eminent colleagues and friends in the truest sense who in various ways have substantially contributed to this book, either indirectly while working together on other projects or directly: James Hardin, director of the Biostatistics Collaborative Unit and professor, Department of Statistics and Epidemiology, University of South Carolina School of Medicine; Andrew Robinson, director, Australian Centre of Excellence for Risk Analysis (ACERA), Department of Mathematics and Statistics, University of Melbourne, Australia; Alain Zuur, senior statistician and director of Highland Statistics Ltd., UK; Peter Bruce, CEO, Institute for Statistics Education (Statistics.com); and John Nelder, late Emeritus Professor of Statistics, Imperial College, UK. John passed away in 2010, just shy of his eighty-sixth birthday; our many discussions over a 20-year period are sorely missed. He definitely

spurred my interest in the negative binomial model. I am fortunate to have known and to have worked with these fine statisticians. Each has enriched my life in different ways.

Others who have contributed to this book's creation include Valerie Troiano and Kuber Dekar of the Institute for Statistics Education; Professor William H. Greene, Department of Economics, New York University, and author of the Limdep econometrics software; Dr. Gordon Johnston, Senior Statistician, SAS Institute, author of the SAS Genmod Procedure; Professor Milan Hejtmanek, Seoul National University, and Dr. Digant Gupta, M.D., director, Outcomes Research, Cancer Treatment Centers of America, both of whom provided long hours reviewing early drafts of the book manuscript. Helen Wheeler, production editor for Cambridge University Press, is also gratefully acknowledged. A special acknowledgment goes to Patricia Branton of Stata Corp., who has provided me with statistical support and friendship for almost a quarter of a century. She has been a part of nearly every text I have written on statistical modeling, including this book.

There have been many others who have contributed to this book as well, but space limits their express acknowledgment. I intend to list all contributors on the book's web site. I invite readers to contact me regarding comments or suggestions about the book. You may email me at hilbe@asu.edu or at the address on my BePress web site listed earlier.

Finally, I must also acknowledge Diana Gillooly, senior editor for mathematical sciences with Cambridge University Press, who first encouraged me to write this monograph. She has provided me with excellent feedback in my attempt to develop a thoroughly applied book on count models. Her help with this book has been invaluable and goes far beyond standard editorial obligations. I also wish to thank my family for yet again supporting my writing of another book. My appreciation goes to my wife, Cheryl L. Hilbe, my children and grandchildren, and our white Maltese dog, Sirr, who sits close by my side for hours while I am typing. I dedicate this book to Cheryl for her support and feedback during the time of this book's preparation.

Joseph M. Hilbe
Florence, Arizona
August 12, 2013

Contents

Preface	xi
Chapter 1	
Varieties of Count Data	1
Some Points of Discussion	1
1.1 What Are Counts?	1
1.2 Understanding a Statistical Count Model	3
1.2.1 Basic Structure of a Linear Statistical Model	3
1.2.2 Models and Probability	7
1.2.3 Count Models	9
1.2.4 Structure of a Count Model	16
1.3 Varieties of Count Models	18
1.4 Estimation – the Modeling Process	22
1.4.1 Software for Modeling	22
1.4.2 Maximum Likelihood Estimation	23
1.4.3 Generalized Linear Models and IRLS Estimation	31
1.5 Summary	33
Chapter 2	
Poisson Regression	35
Some Points of Discussion	35
2.1 Poisson Model Assumptions	36
2.2 Apparent Overdispersion	39

2.3	Constructing a “True” Poisson Model	41
2.4	Poisson Regression: Modeling Real Data	48
2.5	Interpreting Coefficients and Rate Ratios	55
2.5.1	How to Interpret a Poisson Coefficient and Associated Statistics	55
2.5.2	Rate Ratios and Probability	59
2.6	Exposure: Modeling over Time, Area, and Space	62
2.7	Prediction	66
2.8	Poisson Marginal Effects	68
2.8.1	Marginal Effect at the Mean	69
2.8.2	Average Marginal Effects	70
2.8.3	Discrete Change or Partial Effects	71
2.9	Summary	73

Chapter 3

Testing Overdispersion

74

	Some Points of Discussion	74
3.1	Basics of Count Model Fit Statistics	74
3.2	Overdispersion: What, Why, and How	81
3.3	Testing Overdispersion	81
3.3.1	Score Test	84
3.3.2	Lagrange Multiplier Test	87
3.3.3	Chi2 Test: Predicted versus Observed Counts	88
3.4	Methods of Handling Overdispersion	92
3.4.1	Scaling Standard Errors: Quasi-count Models	92
3.4.2	Quasi-likelihood Models	96
3.4.3	Sandwich or Robust Variance Estimators	99
3.4.4	Bootstrapped Standard Errors	105
3.5	Summary	106

Chapter 4

Assessment of Fit

108

	Some Points of Discussion	108
4.1	Analysis of Residual Statistics	108
4.2	Likelihood Ratio Test	112

4.2.1	Standard Likelihood Ratio Test	112
4.2.2	Boundary Likelihood Ratio Test	114
4.3	Model Selection Criteria	116
4.3.1	Akaike Information Criterion	116
4.3.2	Bayesian Information Criterion	119
4.4	Setting up and Using a Validation Sample	122
4.5	Summary and an Overview of the Modeling Process	123
4.5.1	Summary of What We Have Thus Far Discussed	124

Chapter 5

Negative Binomial Regression 126

	Some Points of Discussion	126
5.1	Varieties of Negative Binomial Models	126
5.2	Negative Binomial Model Assumptions	128
5.2.1	A Word Regarding Parameterization of the Negative Binomial	133
5.3	Two Modeling Examples	136
5.3.1	Example: rwm1984	136
5.3.2	Example: medpar	148
5.4	Additional Tests	152
5.4.1	General Negative Binomial Fit Tests	152
5.4.2	Adding a Parameter – NB-P Negative Binomial	153
5.4.3	Modeling the Dispersion – Heterogeneous Negative Binomial	156
5.5	Summary	160

Chapter 6

Poisson Inverse Gaussian Regression 162

	Some Points of Discussion	162
6.1	Poisson Inverse Gaussian Model Assumptions	162
6.2	Constructing and Interpreting the PIG Model	165
6.2.1	Software Considerations	165
6.2.2	Examples	165
6.3	Summary – Comparing Poisson, NB, and PIG Models	170

Chapter 7	
Problems with Zeros	172
Some Points of Discussion	172
7.1 Counts without Zeros – Zero-Truncated Models	173
7.1.1 Zero-Truncated Poisson (ZTP)	174
7.1.2 Zero-Truncated Negative Binomial (ZTNB)	177
7.1.3 Zero-Truncated Poisson Inverse Gaussian (ZTPIG)	180
7.1.4 Zero-Truncated NB-P (ZTNBP)	182
7.1.5 Zero-Truncated Poisson Log-Normal (ZTPLN)	183
7.1.6 Zero-Truncated Model Summary	184
7.2 Two-Part Hurdle Models	184
7.2.1 Poisson and Negative Binomial Logit Hurdle Models	185
7.2.2 PIG-Logit and Poisson Log-Normal Hurdle Models	192
7.2.3 PIG-Poisson Hurdle Model	194
7.3 Zero-Inflated Mixture Models	196
7.3.1 Overview and Guidelines	196
7.3.2 Fit Tests for Zero-Inflated Models	197
7.3.3 Fitting Zero-Inflated Models	197
7.3.4 Good and Bad Zeros	198
7.3.5 Zero-Inflated Poisson (ZIP)	199
7.3.6 Zero-Inflated Negative Binomial (ZINB)	202
7.3.7 Zero-Inflated Poisson Inverse Gaussian (ZIPIG)	206
7.4 Summary – Finding the Optimal Model	207
Chapter 8	
Modeling Underdispersed Count Data – Generalized Poisson	210
Some Points of Discussion	210
Chapter 9	
Complex Data: More Advanced Models	217
Types of Data and Problems Dealt with in This Chapter	217
9.1 Small and Unbalanced Data – Exact Poisson Regression	218
9.2 Modeling Truncated and Censored Counts	224
9.2.1 Truncated Count Models	225
9.2.2 Censored Count Models	229
9.2.3 Poisson-Logit Hurdle at 3 Model	231

9.3	Counts with Multiple Components – Finite Mixture Models	232
9.4	Adding Smoothing Terms to a Model – GAM	235
9.5	When All Else Fails: Quantile Count Models	238
9.6	A Word about Longitudinal and Clustered Count Models	239
9.6.1	Generalized Estimating Equations (GEEs)	239
9.6.2	Mixed-Effects and Multilevel Models	241
9.7	Three-Parameter Count Models	245
9.8	Bayesian Count Models – Future Directions of Modeling?	248
9.9	Summary	252
Appendix: SAS Code		255
Bibliography		269
Index		277

CHAPTER 1

Varieties of Count Data

SOME POINTS OF DISCUSSION

- What are counts? What are count data?
- What is a linear statistical model?
- What is the relationship between a probability distribution function (PDF) and a statistical model?
- What are the parameters of a statistical model? Where do they come from, and can we ever truly know them?
- How does a count model differ from other regression models?
- What are the basic count models, and how do they relate with one another?
- What is overdispersion, and why is it considered to be the fundamental problem when modeling count data?

1.1 WHAT ARE COUNTS?

When discussing the modeling of count data, it's important to clarify exactly what is meant by a count, as well as "count data" and "count variable." The word "count" is typically used as a verb meaning to enumerate units, items, or events. We might count the *number* of road kills observed on a stretch of highway, *how many* patients died at a particular hospital within 48 hours of having a myocardial infarction, or *how many* separate sunspots were observed in March 2013. "Count data," on the other hand, is a plural noun referring

to observations made about events or items that are enumerated. In statistics, count data refer to observations that have only nonnegative integer values ranging from zero to some greater undetermined value. Theoretically, counts can range from zero to infinity, but they are always limited to some lesser distinct value – generally the maximum value of the count data being modeled. When the data being modeled consist of a large number of distinct values, even if they are positive integers, many statisticians prefer to model the counts as if they were continuous data. We address this issue later in the book.

A “count variable” is a specific list or array of count data. Again, such observations can only take on nonnegative integer values. However, in a statistical model, a response variable is understood as being a random variable, meaning that the particular set of enumerated values or counts could be other than they are at any given time. Moreover, the values are assumed to be independent of one another (i.e., they show no clear evidence of correlation). This is an important criterion for count model data, and it stems from the fact that the observations of a probability distribution are independent. On the other hand, predictor values are fixed; that is, they are given as facts, which are used to better understand the response.

We will be primarily concerned with four types of count variables in this book. They are:

1. A count or enumeration of events
2. A count of items or events occurring within a period of time or over a number of periods
3. A count of items or events occurring in a given geographical or spatial area or over various defined areas
4. A count of the number of people having a particular disease, adjusted by the size of the population at risk of contracting the disease

Understanding how count data are modeled, and what modeling entails, is discussed in the following section. For readers with little background in linear models, I strongly suggest that you read through Chapter 1 even though various points may not be fully understood. Then re-read the chapter carefully. The essential concepts and relationships involved in modeling should then be clear. In Chapter 1, I have presented the fundamentals of modeling, focusing on normal and count model estimation from several viewpoints, which should at the end provide the reader with a sense of how the modeling process is to be understood when applied to count models. If certain points are still