

Monographs
on Statistics and
Applied Probability 62

Nonlinear Models for Repeated Measurement Data

Marie Davidian and
David M. Giltinan

MONOGRAPHS ON STATISTICS AND APPLIED PROBABILITY

General Editors

V. Isham, N. Keiding, T. Louis, N. Reid, R. Tibshirani, and H. Tong

- 1 Stochastic Population Models in Ecology and Epidemiology *M.S. Barlett* (1960)
 - 2 Queues *D.R. Cox and W.L. Smith* (1961)
- 3 Monte Carlo Methods *J.M. Hammersley and D.C. Handscomb* (1964)
- 4 The Statistical Analysis of Series of Events *D.R. Cox and P.A.W. Lewis* (1966)
 - 5 Population Genetics *W.J. Ewens* (1969)
- 6 Probability, Statistics and Time *M.S. Barlett* (1975)
 - 7 Statistical Inference *S.D. Silvey* (1975)
- 8 The Analysis of Contingency Tables *B.S. Everitt* (1977)
- 9 Multivariate Analysis in Behavioural Research *A.E. Maxwell* (1977)
 - 10 Stochastic Abundance Models *S. Engen* (1978)
- 11 Some Basic Theory for Statistical Inference *E.J.G. Pitman* (1979)
 - 12 Point Processes *D.R. Cox and V. Isham* (1980)
 - 13 Identification of Outliers *D.M. Hawkins* (1980)
 - 14 Optimal Design *S.D. Silvey* (1980)
- 15 Finite Mixture Distributions *B.S. Everitt and D.J. Hand* (1981)
 - 16 Classification *A.D. Gordon* (1981)
- 17 Distribution-Free Statistical Methods, 2nd edition *J.S. Maritz* (1995)
- 18 Residuals and Influence in Regression *R.D. Cook and S. Weisberg* (1982)
 - 19 Applications of Queueing Theory, 2nd edition *G.F. Newell* (1982)
- 20 Risk Theory, 3rd edition *R.E. Beard, T. Pentikäinen and E. Pesonen* (1984)
 - 21 Analysis of Survival Data *D.R. Cox and D. Oakes* (1984)
- 22 An Introduction to Latent Variable Models *B.S. Everitt* (1984)
 - 23 Bandit Problems *D.A. Berry and B. Fristedt* (1985)
- 24 Stochastic Modelling and Control *M.H.A. Davis and R. Vinter* (1985)
 - 25 The Statistical Analysis of Composition Data *J. Aitchison* (1986)
- 26 Density Estimation for Statistics and Data Analysis *B.W. Silverman* (1986)
 - 27 Regression Analysis with Applications *G.B. Wetherill* (1986)
 - 28 Sequential Methods in Statistics, 3rd edition
G.B. Wetherill and K.D. Glazebrook (1986)
 - 29 Tensor Methods in Statistics *P. McCullagh* (1987)
 - 30 Transformation and Weighting in Regression
R.J. Carroll and D. Ruppert (1988)
 - 31 Asymptotic Techniques for Use in Statistics
O.E. Bandorff-Nielsen and D.R. Cox (1989)
- 32 Analysis of Binary Data, 2nd edition *D.R. Cox and E.J. Snell* (1989)
 - 33 Analysis of Infectious Disease Data *N.G. Becker* (1989)
- 34 Design and Analysis of Cross-Over Trials *B. Jones and M.G. Kenward* (1989)

- 35 Empirical Bayes Methods, 2nd edition *J.S. Maritz and T. Lwin* (1989)
- 36 Symmetric Multivariate and Related Distributions
K.T. Fang, S. Kotz and K.W. Ng (1990)
- 37 Generalized Linear Models, 2nd edition *P. McCullagh and J.A. Nelder* (1989)
- 38 Cyclic and Computer Generated Designs, 2nd edition
J.A. John and E.R. Williams (1995)
- 39 Analog Estimation Methods in Econometrics *C.F. Manski* (1988)
- 40 Subset Selection in Regression *A.J. Miller* (1990)
- 41 Analysis of Repeated Measures *M.J. Crowder and D.J. Hand* (1990)
- 42 Statistical Reasoning with Imprecise Probabilities *P. Walley* (1991)
- 43 Generalized Additive Models *T.J. Hastie and R.J. Tibshirani* (1990)
- 44 Inspection Errors for Attributes in Quality Control
N.L. Johnson, S. Kotz and X. Wu (1991)
- 45 The Analysis of Contingency Tables, 2nd edition *B.S. Everitt* (1992)
- 46 The Analysis of Quantal Response Data *B.J.T. Morgan* (1992)
- 47 Longitudinal Data with Serial Correlation—A state-space approach
R.H. Jones (1993)
- 48 Differential Geometry and Statistics *M.K. Murray and J.W. Rice* (1993)
- 49 Markov Models and Optimization *M.H.A. Davis* (1993)
- 50 Networks and Chaos—Statistical and probabilistic aspects
O.E. Barndorff-Nielsen, J.L. Jensen and W.S. Kendall (1993)
- 51 Number-Theoretic Methods in Statistics *K.-T. Fang and Y. Wang* (1994)
- 52 Inference and Asymptotics *O.E. Barndorff-Nielsen and D.R. Cox* (1994)
- 53 Practical Risk Theory for Actuaries
C.D. Daykin, T. Pentikäinen and M. Pesonen (1994)
- 54 Biplots *J.C. Gower and D.J. Hand* (1996)
- 55 Predictive Inference—An introduction *S. Geisser* (1993)
- 56 Model-Free Curve Estimation *M.E. Tarter and M.D. Lock* (1993)
- 57 An Introduction to the Bootstrap *B. Efron and R.J. Tibshirani* (1993)
- 58 Nonparametric Regression and Generalized Linear Models
P.J. Green and B.W. Silverman (1994)
- 59 Multidimensional Scaling *T.F. Cox and M.A.A. Cox* (1994)
- 60 Kernel Smoothing *M.P. Wand and M.C. Jones* (1995)
- 61 Statistics for Long Memory Processes *J. Beran* (1995)
- 62 Nonlinear Models for Repeated Measurement Data
M. Davidian and D.M. Giltinan (1995)
- 63 Measurement Error in Nonlinear Models
R.J. Carroll, D. Rupert and L.A. Stefanski (1995)
- 64 Analyzing and Modeling Rank Data *J.J. Marden* (1995)
- 65 Time Series Models—In econometrics, finance and other fields
D.R. Cox, D.V. Hinkley and O.E. Barndorff-Nielsen (1996)
- 66 Local Polynomial Modeling and its Applications *J. Fan and I. Gijbels* (1996)
- 67 Multivariate Dependencies—Models, analysis and interpretation
D.R. Cox and N. Wermuth (1996)

- 68 Statistical Inference—Based on the likelihood *A. Azzalini* (1996)
- 69 Bayes and Empirical Bayes Methods for Data Analysis
B.P. Carlin and T.A. Louis (1996)
- 70 Hidden Markov and Other Models for Discrete-Valued Time Series
I.L. Macdonald and W. Zucchini (1997)
- 71 Statistical Evidence—A likelihood paradigm *R. Royall* (1997)
- 72 Analysis of Incomplete Multivariate Data *J.L. Schafer* (1997)
- 73 Multivariate Models and Dependence Concepts *H. Joe* (1997)
- 74 Theory of Sample Surveys *M.E. Thompson* (1997)
- 75 Retrial Queues *G. Falin and J.G.C. Templeton* (1997)
- 76 Theory of Dispersion Models *B. Jørgensen* (1997)
- 77 Mixed Poisson Processes *J. Grandell* (1997)
- 78 Variance Components Estimation—Mixed models, methodologies and applications
P.S.R.S. Rao (1997)
- 79 Bayesian Methods for Finite Population Sampling
G. Meeden and M. Ghosh (1997)
- 80 Stochastic Geometry—Likelihood and computation
O.E. Barndorff-Nielsen, W.S. Kendall and M.N.M. van Lieshout (1998)
- 81 Computer-Assisted Analysis of Mixtures and Applications—
Meta-analysis, disease mapping and others *D. Böhning* (1999)
- 82 Classification, 2nd edition *A.D. Gordon* (1999)
- 83 Semimartingales and their Statistical Inference *B.L.S. Prakasa Rao* (1999)
- 84 Statistical Aspects of BSE and vCJD—Models for Epidemics
C.A. Donnelly and N.M. Ferguson (1999)
- 85 Set-Indexed Martingales *G. Ivanoff and E. Merzbach* (2000)
- 86 The Theory of the Design of Experiments *D.R. Cox and N. Reid* (2000)
- 87 Complex Stochastic Systems
O.E. Barndorff-Nielsen, D.R. Cox and C. Klüppelberg (2001)
- 88 Multidimensional Scaling, 2nd edition *T.F. Cox and M.A.A. Cox* (2001)
- 89 Algebraic Statistics—Computational Commutative Algebra in Statistics
G. Pistone, E. Riccomagno and H.P. Wynn (2001)
- 90 Analysis of Time Series Structure—SSA and Related Techniques
N. Golyandina, V. Nekrutkin and A.A. Zhigljavsky (2001)
- 91 Subjective Probability Models for Lifetimes
Fabio Spizzichino (2001)
- 92 Empirical Likelihood *Art B. Owen* (2001)
- 93 Statistics in the 21st Century *Adrian E. Raftery, Martin A. Tanner,
and Martin T. Wells* (2001)
- 94 Accelerated Life Models: Modeling and Statistical Analysis
Vilijandas Bagdonavičius and Mikhail Nikulin (2001)
- 95 Subset Selection in Regression, Second Edition
Alan Miller (2002)
- 96 Topics in Modelling of Clustered Data
Marc Aerts, Helena Geys, Geert Molenberghs, and Louise M. Ryan (2002)
- 97 Components of Variance *D.R. Cox and P.J. Solomon* (2002)



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Nonlinear Models for Repeated Measurement Data

Marie Davidian

Associate Professor

Department of Biostatistics

Harvard School of Public Health

Boston

USA

David M. Giltinan

Principal Biostatistician

Genentech Inc.

San Francisco

USA

CHAPMAN & HALL/CRC

A CRC Press Company

Boca Raton London New York Washington, D.C.

Library of Congress Cataloging-in-Publication Data

Catalog record is available from the Library of Congress

This book contains information obtained from authentic and highly regarded sources. Reprinted material is quoted with permission, and sources are indicated. A wide variety of references are listed. Reasonable efforts have been made to publish reliable data and information, but the author and the publisher cannot assume responsibility for the validity of all materials or for the consequences of their use.

Apart from any fair dealing for the purpose of research or private study, or criticism or review, as permitted under the UK Copyright Designs and Patents Act, 1988, this publication may not be reproduced, stored or transmitted, in any form or by any means, electronic or mechanical, including photocopying, micro-filming, and recording, or by any information storage or retrieval system, without the prior permission in writing of the publishers, or in the case of reprographic reproduction only in accordance with the terms of the licenses issued by the Copyright Licensing Agency in the UK, or in accordance with the terms of the license issued by the appropriate Reproduction Rights Organization outside the UK.

The consent of CRC Press LLC does not extend to copying for general distribution, for promotion, for creating new works, or for resale. Specific permission must be obtained in writing from CRC Press LLC for such copying.

Direct all inquiries to CRC Press LLC, 2000 N.W. Corporate Blvd., Boca Raton, Florida 33431.

Trademark Notice: Product or corporate names may be trademarks or registered trademarks, and are used only for identification and explanation, without intent to infringe.

Visit the CRC Press Web site at www.crcpress.com

© 1995 by M. Davidian and D.M. Giltinan

First edition 1995

First CRC press reprint 1998

Originally published by Chapman & Hall

No claim to original U.S. Government works

International Standard Book Number 0-412-98341-9

Printed in the United States of America

5 6 7 8 9 0

Printed on acid-free paper

Contents

Preface	xiii
1 Introduction	1
1.1 Motivating examples	2
1.1.1 Pharmacokinetics of cefamandole	2
1.1.2 Population pharmacokinetics of quinidine	3
1.1.3 Growth analysis for soybean plants	7
1.1.4 Bioassay for relaxin by RIA	8
1.2 Model specification	10
1.3 Outline of this book	13
2 Nonlinear regression models for individual data	17
2.1 Introduction	17
2.2 Model specification	19
2.2.1 Basic nonlinear regression model	20
2.2.2 Classical assumptions	20
2.2.3 Generalizations of the classical framework	21
2.3 Inference	26
2.3.1 Ordinary least squares	27
2.3.2 Generalized least squares	28
2.3.3 Variance function estimation	31
2.3.4 General covariance structures	34
2.3.5 Confidence intervals and hypothesis testing	36
2.4 Computational aspects	43
2.4.1 Least squares estimation	43
2.4.2 Variance function estimation	46
2.4.3 General covariance structures	47
2.5 Examples	48
2.6 Related approaches	55
2.6.1 Generalized linear models	55

2.6.2	Generalized estimating equations	56
2.7	Discussion	57
2.8	Bibliographic notes	61
3	Hierarchical linear models	63
3.1	Introduction	63
3.2	Model specification	64
3.2.1	Examples	64
3.2.2	General linear mixed effects model	67
3.2.3	Bayesian model specification	70
3.3	Inference	72
3.3.1	One-way classification with random effects	72
3.3.2	General linear mixed effects model	76
3.3.3	Bayesian inference	80
3.4	Computational aspects	85
3.4.1	EM algorithm	85
3.4.2	Newton-Raphson and the method of scoring	89
3.4.3	Software implementation	90
3.4.4	Examples revisited	90
3.5	Limitations of the normal hierarchical linear model	92
3.6	Bibliographic notes	95
4	Hierarchical nonlinear models	97
4.1	Introduction	97
4.2	General hierarchical nonlinear model	98
4.2.1	Basic model	98
4.2.2	Intra-individual variation	100
4.2.3	Inter-individual variation	101
4.2.4	Summary	108
4.2.5	Time-dependent covariates	109
4.2.6	Accommodation of multiple responses	110
4.3	Model specification	113
4.3.1	Fully parametric model specification	113
4.3.2	Nonparametric model specification	114
4.3.3	Semiparametric model specification	115
4.3.4	Bayesian model specification	117
4.4	Discussion	118
5	Inference based on individual estimates	125
5.1	Introduction	125
5.2	Construction of individual estimates	126
5.2.1	Generalized least squares	126

5.2.2	Pooled estimation of covariance parameters	127
5.2.3	Confidence regions for covariance parameters	128
5.2.4	Examples	129
5.3	Estimation of population parameters	136
5.3.1	Standard two-stage method	137
5.3.2	Global two-stage method	138
5.3.3	Bayesian method	142
5.3.4	General inter-individual models	142
5.3.5	Choice of individual estimates	144
5.4	Computational aspects	144
5.5	Example	145
5.6	Discussion	147
5.7	Bibliographic notes	150
6	Inference based on linearization	151
6.1	Introduction	151
6.2	First-order linearization	152
6.2.1	Approximate model specification	152
6.2.2	Maximum likelihood	154
6.2.3	Generalized least squares	157
6.2.4	Connection with generalized estimating equations	163
6.3	Conditional first-order linearization	164
6.3.1	Approximate model specification	165
6.3.2	Generalized least squares	166
6.3.3	Alternative derivations and methods	173
6.4	Software implementation	174
6.5	Graphical approaches to model selection	175
6.6	Examples	176
6.7	Discussion	186
6.8	Bibliographic notes	190
7	Nonparametric and semiparametric inference	191
7.1	Introduction	191
7.2	Nonparametric maximum likelihood (NPML)	192
7.2.1	Basic ideas	192
7.2.2	Algorithms for finding the maximum likelihood estimate	196
7.2.3	Incorporation of covariates	198
7.3	Smooth nonparametric maximum likelihood (SNP)	200
7.3.1	Estimation	200
7.3.2	Confidence intervals and hypothesis testing	206

7.4	Other approaches	208
7.5	Software implementation	209
7.6	Example	209
7.7	Discussion	214
8	Bayesian inference	217
8.1	Introduction	217
8.2	Model framework	218
8.2.1	General three-stage hierarchical model	218
8.2.2	Simple examples	219
8.3	Inference	220
8.3.1	Inference for the general hierarchical model	220
8.3.2	Conditional distributions	221
8.3.3	From conditionals to marginals: the Gibbs sampler	222
8.3.4	Sampling from conditional distributions	224
8.4	Description of the Gibbs sampler	225
8.4.1	Bivariate case	225
8.4.2	More than two variables	228
8.4.3	Sampling from conditional distributions	229
8.4.4	Convergence and implementation	230
8.5	More complex models	231
8.6	Software implementation	235
8.7	Discussion	235
9	Pharmacokinetic and pharmacodynamic analysis	237
9.1	Introduction	237
9.2	Background	238
9.2.1	Population pharmacokinetics	238
9.2.2	Pharmacodynamic modeling	239
9.3	Population pharmacokinetics of quinidine	242
9.4	Pharmacokinetics of IGF-I in trauma patients	253
9.5	Dose escalation study of argatroban	262
9.6	Discussion	272
10	Analysis of assay data	275
10.1	Introduction	275
10.2	Background on assay methods	276
10.2.1	Radioimmunoassay	276
10.2.2	Enzyme-linked immunosorbent assay	277
10.3	Model framework	278
10.3.1	Model for the mean response	278

10.3.2 Modeling intra-assay response variation	280
10.4 Assessing intra-assay precision	281
10.4.1 Estimation of variance parameters	281
10.4.2 Intra-assay precision profiles	282
10.5 Assay sensitivity and detection limits	288
10.6 Confidence intervals for unknown concentrations	290
10.6.1 Methods of forming confidence intervals	290
10.6.2 Examples	291
10.6.3 Simulation studies	294
10.7 Bayesian calibration	294
10.8 Discussion	297
10.9 Bibliographic notes	298
11 Further applications	299
11.1 Introduction	299
11.2 Comparison of soybean growth patterns	299
11.3 Prediction of bole volume of sweetgum trees	310
11.4 Prediction of strong motion of earthquakes	319
12 Open problems and discussion	327
References	333
Author index	349
Subject index	355

TO OUR PARENTS

Preface

Analysis of repeated measurement data is a recurrent challenge to statisticians engaged in biological and biomedical applications. For example, data from clinical trials are often longitudinal in nature, with repeated measures of response taken over time. In pharmacokinetic studies, serial measurements of drug concentrations are taken from each participant. By definition, growth studies involve repeated measurements over time. In other applications, measurements on experimental units may be repeated across some other dimension, e.g. spatially rather than temporally.

Methods for *linear* modeling of repeated measurement data are well developed and well documented in the statistical literature; recent accounts include those by Crowder and Hand (1990), Lindsey (1993), and Diggle, Liang and Zeger (1994). In many biological applications, such as pharmacokinetic analysis and studies of growth and decay, however, *nonlinear* modeling is required for meaningful analysis. For this type of modeling, statistical approaches are less well understood, and discussion of appropriate methodology is scattered across a wide literature. Recent years have seen more attention to nonlinear repeated measurement data in the statistical literature; however, the economy of style imposed by many journals means that the material is sometimes presented in a manner that does not make it readily accessible to practicing statisticians. The result is that, although nonlinear modeling of repeated measurement data represents an area of some practical importance, it is one that still appears to engender a good deal of confusion among data analysts.

Our purpose in writing this monograph is to provide a clear delineation of currently available modeling approaches and inferential methods for nonlinear repeated measures. The goal is to make the material accessible to a wide audience. The book is targeted mainly to practicing biostatisticians in industry and academia, and to

graduate students in statistics or biostatistics. We have attempted to keep the exposition at an intermediate level, however, so that the majority of the material should also be accessible to pharmacokineticists and to researchers in the clinical and biological sciences.

The model framework that forms the basis for the inferential methods discussed in the book is that of the hierarchical nonlinear model for continuous response data. This may be viewed as an extension of standard nonlinear modeling techniques to accommodate multiple levels of variability (within and among individuals). Alternatively, it may be regarded as a generalization of the hierarchical linear model framework to include models that are nonlinear in parameters of interest. Hierarchical nonlinear modeling may thus be expected to inherit the computational difficulties intrinsic to both nonlinear regression and to hierarchical linear models. We have certainly found this to be true in practice; computational issues can be formidable at times, so that inference within this framework is not an enterprise to be undertaken lightly. We have included several case studies in later chapters; these represent 'real-life' data sets. We have tried to report analyses in sufficient detail to give the reader a realistic sense, not only of the potential scope and utility of the methods discussed, but also of the potential difficulties. Because computational aspects play a key role in the implementation of all the techniques described in this book, we have tried to include some discussion of available software in each of the relevant chapters. Most of the data sets considered in this book will be available on Statlib, together with code to implement model fits discussed in the text using various software packages.

Several friends and colleagues helped us while writing this book. We are grateful to Sharon Baughman, Doug Bates, Eric Chi, Art DeVault, Jim Frane, Tim Gregoire, Karen Higgins, Debbi Kotlovker, Cynthia Ladd, Nishit Modi, James Reimann, Alan Schumitzky, Anastasios Tsiatis, Jon Wakefield, and Fong Wang-Clow for comments on earlier drafts of the manuscript. Thanks go to researchers at Genentech, Inc., and elsewhere for permission to use their data in the book. Special thanks are due to Alan Hopkins and to Genentech for granting the second author a leave of absence to complete the manuscript and for encouragement throughout the writing process. Moral support was provided by Peter Compton, Carol Deasy, Ellen Gilkerson, Debbi Kotlovker, James Reimann, and Georgia Thompson. Finally, words cannot adequately express our debt of gratitude to Butch Tsiatis, without whose keen statistical insight, ongoing moral and culinary support, and unflinching

good humor this manuscript would never have been completed.

This book was typeset using $\text{\LaTeX}$; figures were created using Splus.

Boston and San Francisco
March 1995

Marie Davidian
David M. Giltinan



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

CHAPTER 1

Introduction

Data consisting of repeated measurements taken on each of a number of individuals arise commonly in biological and biomedical applications. For example, in longitudinal clinical studies, measurements are taken on each of a number of subjects over time. Similarly, participants in pharmacokinetic experiments undergo serial blood sampling following administration of a test agent. Pharmacodynamic studies may involve repeated measurement of physiological effect in the same subject in response to differing doses of a drug. By definition, studies of growth and decay involve repeated measurements taken on sample units, which could be human or animal subjects, plants, or cultures.

Modeling data of this kind usually involves characterization of the relationship between the measured response, y , and the repeated measurement factor, or covariate, x . In many applications, the proposed systematic relationship between y and x is nonlinear in unknown parameters of interest. In some cases, the relevant nonlinear model may be derived on physical or mechanistic grounds. In other contexts, a nonlinear relationship may be used simply to provide an empirical description of the data.

The presence of repeated observations on an individual requires particular care in characterizing the random variation in the data. It is important to recognize two sources of variability explicitly: random variation among measurements *within* a given individual and random variation *among* individuals. Inferential procedures accommodate these different variance components within the framework of an appropriate hierarchical statistical model. When the postulated relationship between y and x is linear in the unknown parameters, the relevant framework is that of the classical linear mixed effects model. An alternative is provided by Bayesian inferential methods for a suitable hierarchical linear model. There is an extensive literature on hierarchical linear models; Searle, Casella,

and McCulloch (1992) provide a comprehensive overview.

Methods for repeated measurement data where the relationship between y and x is *nonlinear* in the unknown parameters are less well developed. Treatment of existing techniques is scattered through a wide literature. The purpose of this monograph is to provide a unified presentation of methods and issues for nonlinear repeated measurement data. We begin by considering several examples to motivate our subsequent development.

1.1 Motivating examples

1.1.1 Pharmacokinetics of cefamandole

Figure 1.1 shows data from a pilot study to investigate the pharmacokinetics of cefamandole, a cephalosporin antibiotic (Aziz *et al.*, 1978). In this experiment, a dose of 15 mg/kg body weight of cefamandole was administered by ten-minute intravenous infusion to six healthy male volunteers. Blood samples were collected from each subject at each of 14 time points post-dose. Drug concentrations in plasma were determined for each sample by high-performance liquid chromatography (HPLC). The resulting plasma concentration-time profiles for each subject are plotted in the figure.

In characterizing the pharmacokinetics of a drug, it is common to represent the body as a system of compartments and to assume that the rates of transfer between compartments follow first-order or linear kinetics. Solution of the resulting differential equations shows that the relationship between drug concentration and time may be described by a sum of exponential terms. For instance, the biexponential equation

$$C(t) = \beta_1 \exp(-\beta_2 t) + \beta_3 \exp(-\beta_4 t), \quad \beta_1, \dots, \beta_4 > 0, \quad (1.1)$$

where $C(t)$ is drug plasma concentration and t is time post-dose, follows from the assumption of a two-compartment model to describe kinetics following intravenous injection (Gibaldi and Perrier, 1982).

The data in Figure 1.1 exhibit similarly shaped profiles for each subject, with possibly different parameter values for different subjects. Variation within each subject about the model in equation (1.1) arises mainly due to the HPLC assay. It is commonly recognized that intra-subject variation of this kind tends to increase with plasma concentration level (Beal and Sheiner, 1988).

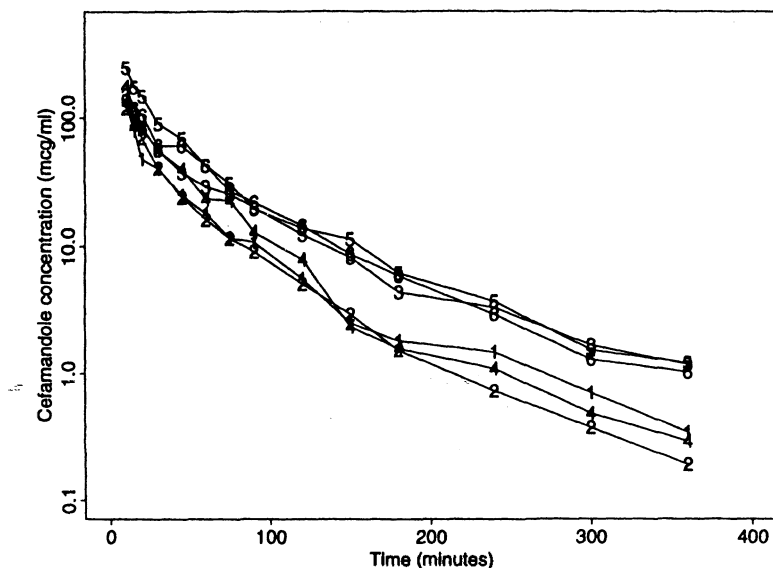


Figure 1.1. *Plasma concentration-time profiles for six subjects, cefamandole data.*

In pilot volunteer studies like this one, primary objectives are to establish an appropriate kinetic model, to obtain preliminary information on values of the model parameters, and to assess the nature of intra-subject variation. Typically, results from the analysis of a pilot study are used as a basis for subsequent investigation of kinetics in a larger, more heterogeneous patient population.

1.1.2 Population pharmacokinetics of quinidine

Data from a clinical study of the pharmacokinetics of the antiarrhythmic agent quinidine reported by Verme *et al.* (1992) consist of quinidine concentration (mg/L) measurements for 136 hospitalized patients (135 men, 1 woman) treated for either atrial fibrillation or ventricular arrhythmias with oral quinidine therapy. A total of 361 quinidine concentration measurements ranging from one to 11 observations per patient were obtained by enzyme immunoassay during the course of routine clinical treatment.

Measurements were taken within a range of 0.08 hours to 70.5

Table 1.1. *Partial data for two subjects, pharmacokinetic study of quinidine. Units for measurements are given in the text.*

<i>time</i>	<i>conc.</i>	<i>dose</i>	<i>SS</i> ¹	<i>age</i>	<i>wt.</i>	<i>creat.</i> ²	<i>glyco.</i> ³
<i>Subject 2</i>							
0.0	—	166	—	58	85	> 50	82
6.0	—	166	—	58	85	> 50	82
12.0	—	166	—	58	85	> 50	82
18.0	—	166	—	58	85	> 50	82
25.0	1.2	—	—	58	85	> 50	82
height = 69, race = Latin, nonsmoker, ethanol abuse, moderate congestive heart failure							
<i>Subject 10</i>							
0.0	—	201	8	73	79	< 50	254
2.2	3.9	—	8	73	79	< 50	254
288.0	—	201	8	73	79	< 50	176
290.0	5.4	—	8	73	79	< 50	176
504.0	—	201	8	73	79	< 50	150
506.0	2.8	—	8	73	79	< 50	150
816.0	—	201	8	73	79	< 50	127
816.2	—	201	8	73	79	< 50	127
817.0	3.1	—	8	73	79	< 50	127
1241.0	—	201	8	73	79	< 50	98
1249.0	—	201	8	73	79	< 50	98
7897.0	—	201	8	74	82	> 50	158
7897.8	1.6	—	8	74	82	> 50	158
height = 69, race = Caucasian, nonsmoker, no ethanol abuse or congestive heart failure							

¹ if numeric, subject has achieved steady state with given dosing interval; if blank, subject has not achieved steady state

² creatinine clearance

³ α_1 -acid glycoprotein concentration

hours after dose. Table 1.1 shows partial data records for two patients selected from the total of 136. Demographic and physiological covariate information was collected for each patient over an observation period ranging from 0.13 hours to 8095.0 hours. The following variables were available for the majority of patients: weight, height, and age, as well as information on race (Latin, Caucasian, Black), smoking status (yes, no), ethanol abuse (yes, no, previously), and status with respect to congestive heart failure

(severe, moderate, mild or none). Weight, age, smoking, and cardiac status were recorded periodically during the study. Creatinine clearance (ml/min), a measure of renal function, and α_1 -acid glycoprotein concentration (mg/dL), the level of a circulating molecule that binds quinidine, were also measured periodically on all patients, although information on creatinine clearance was recorded in a categorical rather than continuous manner. Baseline albumin concentration (g/dL) measurements were available for some, but not all, patients. Oral quinidine may be administered in two different forms; in this study, it was given as quinidine sulfate to 53 patients, as quinidine gluconate to 57 patients, and in both forms to 26 patients. Doses were adjusted for differences in salt content between the two forms by conversion of both forms to milligrams of quinidine base.

One possible characterization of quinidine disposition is to use a one-compartment open model with first-order absorption (Verme *et al.*, 1992). Let $C(t)$ be the concentration of quinidine and let $C_a(t)$ be the apparent concentration of quinidine in the absorption depot at time t . Written in recursive form, the model is:

For the non-steady state at a dosage time $t = t_\ell$

$$\begin{aligned} C_a(t_\ell) &= C_a(t_{\ell-1}) \exp\{-k_a(t_\ell - t_{\ell-1})\} + FD_\ell/V \\ C(t_\ell) &= C(t_{\ell-1}) \exp\{-k_e(t_\ell - t_{\ell-1})\} + C_a(t_{\ell-1}) \frac{k_a}{k_a - k_e} \\ &\quad \times \left[\exp\{-k_e(t_\ell - t_{\ell-1})\} - \exp\{-k_a(t_\ell - t_{\ell-1})\} \right]. \end{aligned} \quad (1.2)$$

For the steady state at a dosage time, $t = t_\ell$

$$\begin{aligned} C_a(t_\ell) &= (FD_\ell/V)(1 - \exp\{-k_a\tau_{ss}\})^{-1} \\ C(t_\ell) &= (FD_\ell/V) \frac{k_a}{k_a - k_e} \\ &\quad \times \left[(1 - \exp\{-k_e\tau_{ss}\})^{-1} - (1 - \exp\{-k_a\tau_{ss}\})^{-1} \right]. \end{aligned} \quad (1.3)$$

Between dosage times, $t_\ell < t < t_{\ell+1}$

$$\begin{aligned} C(t) &= C(t_\ell) \exp\{-k_e(t - t_\ell)\} + C_a(t_\ell) \frac{k_a}{k_a - k_e} \\ &\quad \times \left[\exp\{-k_e(t - t_\ell)\} - \exp\{-k_a(t - t_\ell)\} \right]. \end{aligned} \quad (1.4)$$

In (1.2)–(1.4), t_ℓ , $\ell = 0, 1, \dots$, are the times at which doses D_ℓ

are administered, $C_a(t_0) = FD_0/V$, $C(t_0) = 0$, F is the fraction of dose available, k_a is the absorption rate constant, $k_e = Cl/V$ is the elimination rate constant, Cl is the clearance, V is the apparent volume of distribution, $Cl, V, k_a > 0$, and τ_{ss} is the steady-state dosing interval.

The data in this example share certain characteristics with the cefamandole data: a common nonlinear model form for all subjects, values of the pharmacokinetic parameters that may differ from subject to subject, and probable heterogeneity of assay variation within subject. There are obvious differences as well. In contrast to the experimental setting, where essentially complete concentration-time profiles are collected for each subject, the quinidine data, collected in a clinical setting, are sparse, with relatively few observations per subject. This difference between data collected in a controlled experimental setting and routine clinical data is fairly typical. Clinical data may be available for a much greater number of patients, but data on any one individual patient tend to be sparse.

A second major difference between the quinidine and cefamandole data sets is the availability of demographic and physiological information that may help to explain inter-subject differences in the disposition of quinidine. This difference between experimental and clinical data is also usually the case. Early Phase I data are gathered frequently from a small number of healthy volunteers in a carefully controlled setting. Routine clinical data typically come from a much more extensive and heterogeneous patient population. This allows broader inferences to be drawn, although the task is complicated by the relative paucity of information at the individual subject level.

The major question of interest in the quinidine and similar clinical studies is identification of the demographical and physiological factors affecting drug disposition in a broad patient population. A thorough understanding of this issue may afford important clinical benefit: the dosage regimen for a given patient may be individualized based on relevant physiological and demographic information for the patient if a model is available relating drug disposition to measured patient covariates. Thus, more accurate titration of dosage may be feasible, avoiding possible suboptimal therapeutic benefit resulting from underdosing and minimizing potential toxicity associated with overdosing. Accurate dosing is of particular importance in drugs with a low therapeutic index, where the window of desirable serum concentrations is relatively narrow. For all

drugs, however, it is clearly beneficial to maximize understanding of factors affecting drug disposition.

Given a model which accurately predicts a subject's pharmacokinetic parameters as a function of physiological and demographic characteristics, there will still remain a random, or unexplained, component of variation among individuals. Quantifying this inter-subject variation in pharmacokinetic parameters is a secondary objective in population analysis of data such as those from the quinine study. Achieving the analysis goals is complicated by several issues: (i) the sparsity of information on individual subjects; (ii) the generally nonlinear dependence of response on the relevant pharmacokinetic parameters; and (iii) inter-subject variability. Methods that allow valid inference in the face of these difficulties form the core of this book.

1.1.3 Growth analysis for soybean plants

Figure 1.2 shows data from an experiment to compare growth patterns of two genotypes of soybean: Plant Introduction #416937 (P), an experimental strain, and Forrest (F), a commercial variety.

In this study, data were collected in each of three consecutive years (1988–1990). At the beginning of the growing season in each year, 16 plots were planted with seeds, eight plots with each genotype. To assess growth, each plot was sampled eight to ten times at approximate weekly intervals. At each sampling time, six plants were randomly selected from each plot, leaves from these plants were weighed, and average leaf weight per plant was calculated for the plot. Different plots in different sites were used in different years.

Inspection of Figure 1.2 indicates that the usual logistic function

$$y = \frac{\beta_1}{1 + \beta_2 \exp(\beta_3 x)}, \quad \beta_1, \beta_2 > 0, \quad \beta_3 < 0, \quad (1.5)$$

provides a reasonable representation of average leaf weight y and time x . It is evident from the figure that considerable variation in the parameters β_1 , β_2 and β_3 that characterize the growth pattern exists among plots for a given genotype. For this kind of growth data, it is reasonable to expect serial correlation among measurements within the same plot. In addition, intra-plot variability may be expected to increase with the average level of response. Thus, these data have several features in common with the previous examples: (i) a nonlinear dependence of response on parameters of

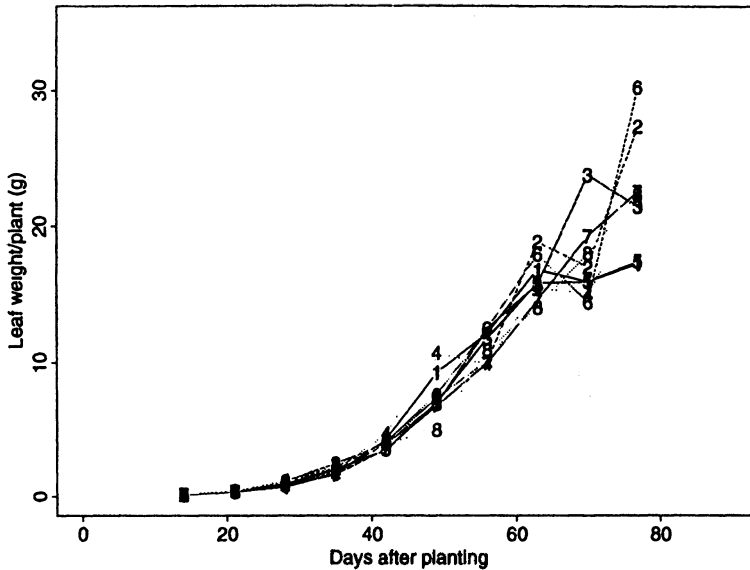


Figure 1.2. Average leaf weight-time profiles for 8 plots planted with Plant Introduction #416937, 1988.

interest; (ii) similarly shaped profiles for each plot; and (iii) heterogeneous within-plot variability, which may also include serial correlation. As with the cefamandole example, sufficient data are available for each plot to allow fitting of model parameters based on data from that plot only.

Variation in growth characteristics may depend on several factors. The primary objective of the experiment was comparison of growth characteristics (initial leaf weight, limiting leaf weight, and growth rate) for the two genotypes. Weather patterns differed considerably over the three years: 1988 was unusually dry, 1989 was wet, and conditions in 1990 were normal. Comparison of growth across weather patterns was also of interest.

1.1.4 Bioassay for relaxin by RIA

Determination of the concentration of a particular protein in an unknown sample frequently relies on immunoassay or bioassay techniques. Bioassay methods are generally based on a relevant measure

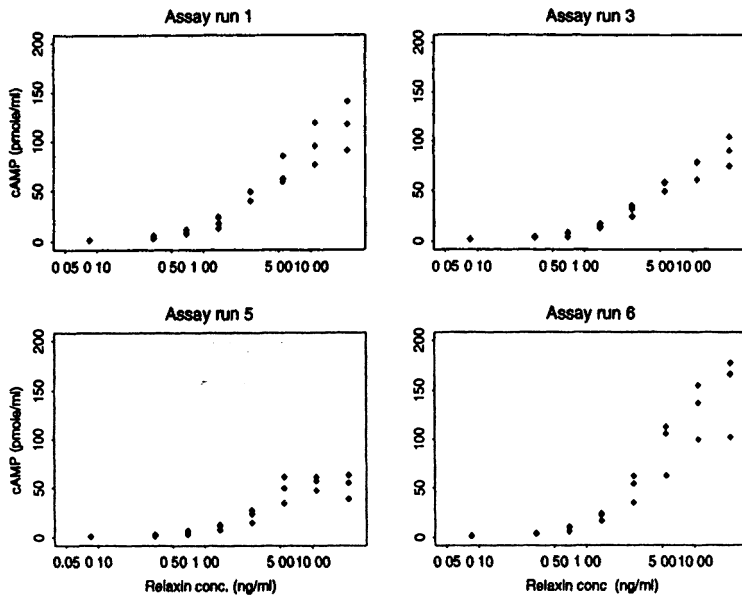


Figure 1.3. Assay response (cAMP)-concentration data for four runs of the relaxin bioassay.

of bioactivity of the protein in question and involve measurement of activity at several known (standard) concentrations of the protein (analyte). The resulting concentration-response curve is used to determine the protein concentration in unknown samples by inverse regression (calibration).

Figure 1.3 shows concentration-response data obtained for standard concentrations in four runs of a bioassay for the therapeutic protein relaxin (Fei *et al.*, 1990). For this assay, bioactivity of relaxin is measured by increased generation and release of intracellular adenosine-3', 5'-cyclic monophosphate (cAMP) by normal human uterine endometrial cells in the presence of relaxin. (cAMP is an enzyme that plays a key role in regulating glycogen metabolism in the cell.) For each a total of nine runs, triplicate cAMP measurements were determined by radioimmunoassay for each of seven known relaxin concentrations. A single measurement at zero standard was also available for each run; by convention, the response at zero concentration has been plotted at two dilutions below the

lowest standard in Figure 1.3. A standard choice of dose-response model in describing this kind of assay data is the four-parameter logistic function

$$y = \beta_1 + \frac{\beta_2 - \beta_1}{1 + \exp\{\beta_4(\log x - \beta_3)\}}, \quad \beta_1, \beta_2 > 0. \quad (1.6)$$

The response y in (1.6) is cAMP level (pmoles/ml), and x represents the known relaxin concentration (ng/ml). The parameters in (1.6) have the following interpretation: β_1 and β_2 represent response at 'infinite' and zero concentration, respectively; β_3 is the $\log EC_{50}$ value, that is, the logarithm of the concentration that gives a response midway between β_1 and β_2 ; and β_4 is a slope parameter governing the steepness of the concentration-response curve. It is evident from Figure 1.3 that values of the parameters may vary considerably from run to run. It is also clear from the plots that within-assay variability depends strongly on response level.

These data share several features with the previous examples: (i) nonlinear dependence of y on regression parameters; (ii) similarly shaped concentration-response profiles for each run, with possibly different parameter values from run to run; and (iii) within-run variability that increases with the level of the response.

What inferential questions are of interest for assay data of this kind? The primary focus is typically on calibration, or inverse regression; that is, estimation of the concentration of analyte in an unknown sample based on the observed response for that sample. Ancillary questions pertain to assay precision and performance. For example, with what precision are unknown samples calibrated? How may one form accurate confidence intervals about estimated concentrations? What is the 'acceptable range' of the assay, where calibrated estimates are sufficiently precise? What is the lowest limit of reliable assay measurement? Can one exploit the similarity among assay runs to improve calibration? These issues will be discussed in detail in Chapter 10. For now, we remark that the inferential challenge is the same as that in the previous examples, to address the questions of interest within a framework that correctly accommodates both the inter- and intra-assay variation.

1.2 Model specification

In the examples of the previous section, several common features of the data may be identified:

- (i) Repeated response measurements taken on a number of different individuals (subjects, plots, or assay runs).
- (ii) Nonlinear dependence of the response y on a set of unknown parameters, β , for each individual.
- (iii) Response profiles that are similarly shaped across individuals, but that may have different values of the parameter vector β for different individuals.
- (iv) A pattern of within-individual variability that is not necessarily homogeneous. Possible deviations from constant variation, or homoscedasticity, include a dependence of the variability on mean response, serial correlation among measurements within an individual, or both.
- (v) Inter-individual variability between regression parameters that may be considered to be random, to be systematically related to individual-specific characteristics, or a combination of both.

To incorporate these features in an inferential setting, a useful strategy is to build a *hierarchical*, or *staged*, model. Full details are presented in Chapter 4; here, we sketch an outline of the approach, using a two-stage model.

The first stage specifies the mean and covariance structure for a given individual. For the i th of m individuals, assume that the $(n_i \times 1)$ response vector y_i , satisfies

$$E(y_i|\beta_i) = f_i(\beta_i) = \begin{bmatrix} f(x_{i1}, \beta_i) \\ \vdots \\ f(x_{in_i}, \beta_i) \end{bmatrix},$$

$$\text{Cov}(y_i|\beta_i) = R_i. \quad (1.7)$$

In (1.7), the function f characterizes the systematic dependence of the response on the repeated measurement conditions for the i th individual, summarized in the covariate vectors $x_{i1}, \dots, x_{in_i}$. The regression function f depends in a nonlinear fashion on a regression parameter β_i specific to the i th individual and has the same basic functional form for all individuals; different individual response patterns are accommodated through the possibility of different β_i values for different individuals as well as through the individual-specific covariate vectors x_{ij} . The matrix R_i is a covariance matrix summarizing the pattern of random variability associated with the data for the i th individual. Along with an assumption about the distribution of y_i , the first stage model (1.7) thus describes both

systematic and random features of the data at the *intra-individual* level.

Inter-individual variation is characterized in a second stage, consisting of a *model* for variation in the regression parameters β_i . This variation can be modeled using a distributional assumption for the β_i at various levels of complexity. For instance, one might specify a parametric 'regression' model for the β_i , e.g.,

$$\beta_i = A_i\beta + b_i. \quad (1.8)$$

In (1.8), β_i is assumed to depend linearly and systematically on a vector of parameters β and on individual-specific information, such as physiological and demographic characteristics in the quinidine study or genotype and weather condition in the soybean growth study, summarized in a design matrix A_i . The 'error' b_i corresponds to the random component of inter-individual variation, which might be taken to have mean zero and covariance matrix D . One could add the further restriction that the distribution of the β_i belongs to a particular parametric family, for example, the multivariate normal distribution. In the example, this would amount to the assumption that

$$\beta_i \sim \mathcal{N}(A_i\beta, D).$$

We shall refer to the case where the variation in the inter-individual random parameters is specified by a parametric model like (1.8) and the random component is assumed to belong to a particular distributional family as the *fully parametric* model specification. At the other end of the spectrum, one might make no assumptions at all about the form or distribution of the β_i . We refer to this as the *nonparametric* model specification. An intermediate possibility is to specify a parametric model for the β_i as in (1.8), but to avoid the assumption of a particular distributional family for the random component. We refer to this kind of model specification as being *semiparametric*. Finally, the kind of two-stage model that we consider may be arrived at naturally from a Bayesian perspective. In the Bayesian view, individual-specific regression parameters β_i are considered to arise from a distribution whose mean and covariance are drawn from an appropriate prior distribution.

Each of the four kinds of model specifications for the random parameters – parametric, nonparametric, semiparametric, or Bayesian – leads to different inferential approaches. One other factor is a major determinant of the inferential technique to be applied: the relative amount of information that is available per individual. We

have seen in the context of pharmacokinetics that two quite distinct scenarios are possible: (i) sparse information on each of a large number of individuals and (ii) rich information on each of a small number of individuals. Intermediate scenarios are also a possibility, of course. Not surprisingly, the sampling design plays a role in determining what analysis methods may be employed; certain methods that can be used when sampling on an individual basis is relatively dense are not applicable to the sparse data case.

We shall use the term 'individual' throughout this book to refer to the experimental unit over which repeated measurements are available. In the repeated measurement literature, use of the term 'subject' is common in this context. Other terms include 'block' and 'stratum;' the term 'cluster' has also been proposed (Lindstrom and Bates, 1990). We avoid the latter term due to its more common usage elsewhere in the statistical literature, and use 'individual' in preference to 'subject' because of its greater generality.

1.3 Outline of this book

The goal of this book is to provide a unified presentation of modeling strategies and inferential procedures for the types of continuous repeated measurement data exemplified by the data sets described in section 1.1. This kind of data has recently received considerable attention, but discussion of inferential methods is scattered across a wide variety of sources, both in the statistical and subject-matter literature. We present methods suitable for *continuous* data of this type; techniques for binary or discrete data are not discussed. By providing a clear delineation of models and methods, we hope to make the relevant techniques more accessible both to statisticians and investigators faced with the challenge of analyzing nonlinear repeated measurement data.

For the most part, we have tried to keep exposition at an intermediate level, with emphasis throughout on applications. Familiarity with regression analysis at the level of a text such as Draper and Smith (1981) and with statistical inference at a first-year graduate level should provide adequate background for most of the material in this book. Chapters 7 and 8 are at a somewhat higher mathematical level than the remainder of the book, but can be omitted without significant loss of continuity.

Hierarchical *nonlinear* modeling provides the central framework for everything else in this book. A review of nonlinear regression methods is given in Chapter 2, which, in the context of repeated

measurement data, is relevant to data from a single individual only. Many of the techniques applicable to hierarchical nonlinear models are direct extensions of methods for individual data. Chapter 3 provides a comprehensive review of hierarchical *linear* models; a good understanding of this material is helpful in making the transition to the nonlinear case. Chapter 4 is central to the rest of the book; here, we lay out the various approaches to hierarchical nonlinear modeling in some detail. As mentioned above, these may be categorized as (i) fully parametric, (ii) semi- and nonparametric, and (iii) Bayesian. Each of these perspectives leads to different inferential strategies, discussed in Chapters 5–8. In Chapters 5 and 6, inferential procedures for the fully parametric (normal theory) hierarchical nonlinear model are presented. Chapter 5 discusses *two-stage* methods, which are applicable only in the case where sufficient data are available for each individual to allow estimation of individual-specific regression parameters based on data for that individual only. Methods presented in Chapter 6 are based on some type of *linearization* of the model and may also be applied to the ‘sparse data’ case. Chapter 7 is devoted to semiparametric and nonparametric inference. Bayesian methods are described in Chapter 8. Each of Chapters 9 and 10 treats a particular area of application in detail. Chapter 9 discusses population pharmacokinetic and pharmacodynamic modeling. The analysis of assay data, with particular reference to immunoassays and bioassays, is covered in Chapter 10. Case studies in the areas of crop science, forestry, and seismology, illustrating the general applicability of the methods, are presented in Chapter 11. The book concludes with a discussion of open issues and general comments (Chapter 12).

Schematically, the organization of the material may be represented as follows:

Chapter 1	Introduction
Chapters 2 and 3	Background material
Chapter 4	Model specification
Chapters 5–8	Inferential methods
Chapters 9–11	Applications and case studies
Chapter 12	Conclusion

It is inevitable that different readers will approach this book with different backgrounds and prior exposure to this material. Depending on background and the primary focus of the reader's interest, different reading strategies are possible. For instance, a researcher in pharmacokinetics, whose primary interest is in population pharmacokinetic and pharmacodynamic modeling, might

choose to include the following sections on a first reading: Chapter 1, Chapter 2, Chapter 3 (excluding sections 3.2.3, 3.3.3, 3.4, and 3.6), sections 4.1, 4.2, and 4.3.1, Chapters 5 and 6, and Chapter 9. This would provide comprehensive coverage of the fully parametric approach to population modeling. Other approaches, such as the nonparametric or Bayesian paradigms in Chapters 7 and 8, could be covered on subsequent readings.

Similarly, the above sequence might provide a useful first approach to statisticians unfamiliar with this material. For these readers, it would also be useful to include Chapters 10 and 11 to achieve a better understanding of the scope of the methods. Readers with interests in applications in biometry not directly related to pharmacokinetics or pharmacodynamics might cover the following chapters: Chapters 1–3 (excluding sections 3.2.3 and 3.3.3), Chapters 4–6, and Chapter 11, possibly including Chapter 10, depending on interest.

Computational aspects play an important role in the practical implementation of all of the techniques described in this book. For this reason, we have generally included a section in each of the relevant chapters that discusses software implementation. Some appreciation of potential computational difficulties is necessary if the methods discussed in this book are to be implemented sensibly.



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

CHAPTER 2

Nonlinear regression models for individual data

2.1 Introduction

Individual repeated measurement data often exhibit a relationship between response and measurement factors that is best characterized by a model nonlinear in its parameters. In some settings, an appropriate nonlinear model may be derived on the basis of theoretical considerations. In other situations, a nonlinear relationship may be employed to provide an empirical description of the data. In this chapter, as a prelude to discussion of the hierarchical nonlinear model for data from several individuals, we review techniques for nonlinear modeling and inference for data from a single individual only.

To fix ideas, consider the data in Table 2.1, taken from a study reported by Kwan *et al.* (1976) of the pharmacokinetics of indomethacin following bolus intravenous injection of the same dose in six human volunteers. For each subject, plasma concentrations of indomethacin were measured at 11 time points ranging from 15 minutes to 8 hours post-injection. In this chapter, we focus on the data for the fifth subject only for purposes of illustration; the concentration-time profile for this individual is shown in Figure 2.1. The usual approach to derivation of a suitable model for pharmacokinetic data of this kind is predicated upon the assumption of a compartment model for the human body, as described in section 1.1.1. This approach suggests that, except for random intra-individual variation, the biexponential function

$$y = \beta_1 \exp(-\beta_2 x) + \beta_3 \exp(-\beta_4 x), \quad \beta_1, \dots, \beta_4 > 0,$$

is a reasonable representation of plasma concentration y as a function of time x .

Table 2.1. *Plasma concentrations ($\mu\text{g/ml}$) following intravenous injection of indomethacin for six human subjects.*

time (hrs)	Subject					
	1	2	3	4	5	6
0.25	1.50	2.03	2.72 ¹	1.85	2.05	2.31
0.50	0.94	1.63	1.49	1.39	1.04	1.44
0.75	0.78	0.71	1.16	1.02	0.81	1.03
1.00	0.48	0.70	0.80	0.89	0.39	0.84
1.25	0.37	0.64	0.80	0.59	0.30	0.64
2.00	0.19	0.36	0.39	0.40	0.23	0.42
3.00	0.12	0.32	0.22	0.16	0.13	0.24
4.00	0.11	0.20	0.12	0.11	0.11	0.17
5.00	0.08	0.25	0.11	0.10	0.08	0.13
6.00	0.07	0.12	0.08	0.07	0.10	0.10
8.00	0.05	0.08	0.08	0.07	0.06	0.09

¹ outlier; not included in later analyses

Individual data that follow a nonlinear model often exhibit response variation that changes systematically with the level of the response. Heterogeneous variation occurs in nearly all fields of application, including those that are the focus of this book. For example, assay data typically exhibit intra-run variance that is an increasing function of the response level. Variation in pharmacokinetic data and data from growth studies is also widely acknowledged to be related systematically to mean response.

Another complication for repeated measurement data may arise from the tendency for observations on a given individual to be related. When measurement is repeated over time, serial correlation may be evident; in other contexts, correlation patterns may be due to factors such as adjacent positioning of samples on an assay microtiter plate, similarity in genetic composition of litter-mates, or spatial orientation of field samples.

Because of the frequency with which these features arise in practice, the overview of nonlinear regression modeling and inference given in this chapter includes detailed discussion of generalizations of the classical ordinary least squares approach to allow for heterogeneous response variance, often called heteroscedasticity, and for correlation. In section 2.2, the nonlinear regression model framework is introduced, and we discuss inference, including methods for

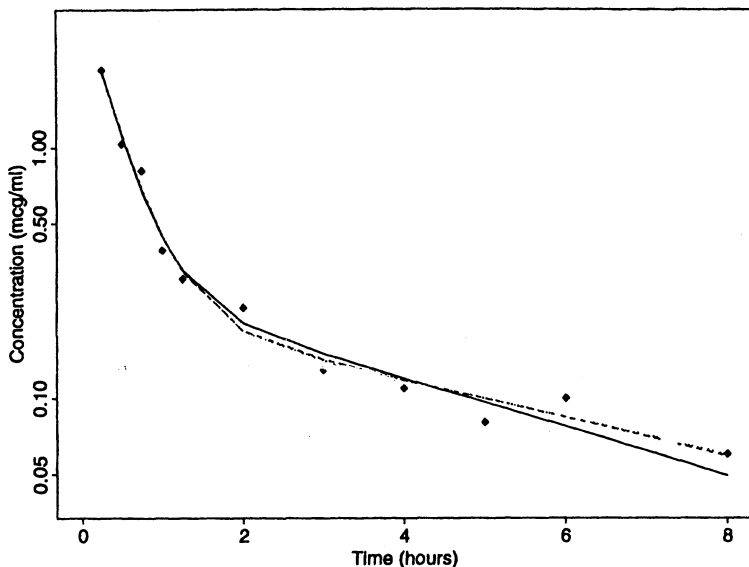


Figure 2.1. Plasma concentration-time profile, indomethacin data, subject 5. Fits of model (2.43) are superimposed; see section 2.5. The solid line is the OLS fit ($\theta = 0$), the dotted line is the GLS fit with $\theta = 1$, and the dashed line is the GLS fit with θ estimated by PL ($\theta = 0.82$).

handling heterogeneity of variance and correlation, in section 2.3. Notes on computational methods are given in section 2.4, and two examples are considered in section 2.5. In section 2.6, we comment on related approaches. Section 2.7 contains a brief discussion, and the chapter concludes with bibliographic notes in section 2.8.

2.2 Model specification

The classical nonlinear regression model described in this section is a direct extension of the linear case. We begin our discussion of this model in section 2.2.1 with a statement of the basic statistical model and notation. In section 2.2.2, we set out the classical assumptions and describe how they may be violated in practice. We describe a number of generalizations in section 2.2.3 to account for departures from the assumptions. For simplicity, the index i for individual is suppressed throughout this chapter.

2.2.1 Basic nonlinear regression model

The basic model for a response variable y has two main components: the nonlinear function characterizing mean response and a specification for intra-individual response variance. Let y_j denote the response taken at the j th covariate value x_j , $j = 1, \dots, n$. The x_j are most often viewed as fixed quantities, and, in the case where they are random, the model and assumptions below are understood to be conditional on their values. In the repeated measurement context, the response vector $y = [y_1, \dots, y_n]'$ summarizes the information available for a single individual taken at values $X = (x'_1, \dots, x'_n)'$ of the repeated measurement factor(s).

The model for the j th observation is usually written as

$$y_j = f(x_j, \beta) + e_j. \quad (2.1)$$

In (2.1), the regression function f depends on β ($p \times 1$), the vector of regression parameters, in a nonlinear fashion. For the indomethacin data,

$$f(x, \beta) = \beta_1 \exp(-\beta_2 x) + \beta_3 \exp(-\beta_4 x), \quad (2.2)$$

with $\beta = [\beta_1, \dots, \beta_4]'$. The random errors e_j reflect uncertainty in the measured response. For repeated measurement data, the vector $e = [e_1, \dots, e_n]'$ thus summarizes the uncertainty for all observations on a given individual.

2.2.2 Classical assumptions

The classical nonlinear regression framework specifies that data arise according to (2.1) together with the following assumptions:

- (i) The errors e_j have mean zero.
- (ii) The errors e_j are uncorrelated.
- (iii) The errors e_j have common variance σ^2 , $\text{Var}(e_j) = \sigma^2$, and are identically distributed for all x_j .
- (iv) The errors e_j are normally distributed.

The first assumption ensures that the model f for mean response is correctly specified. This assumption is rarely called into question, as it is usually the case that the form of the covariate-response relationship is fairly well understood, especially for nonlinear relationships, where the model may result directly from theoretical considerations.

The remaining three assumptions are fairly restrictive and may not hold in some applications. Assumption (iv) is a reflection of

the emphasis of much of statistical methodology on the historical Gaussian paradigm and forms the basis for the standard approach to inference. In applications for which nonlinear models are appropriate, however, this assumption may be particularly unrealistic. For instance, for given x_j , the distribution of biological data may be heavily skewed or subject to a higher intensity of outlying observations than would be expected under normality.

If the assumption of normality (iv) does hold, along with that of uncorrelated errors (ii), then the errors are independently distributed. Thus, the assumption of independence is usually made directly, and under the full set of assumptions, the y_j are independently normally distributed with

$$E(y_j) = f(x_j, \beta), \quad \text{Var}(y_j) = \sigma^2. \quad (2.3)$$

For repeated measurement data, however, even the less stringent specification of uncorrelated errors may be unrealistic; for example, measurements taken over time on a given individual may be serially related.

Assumption (iii), that of constant intra-individual response variance, is violated frequently in practice. For example, growth data often exhibit constant coefficient of variation rather than constant variance; that is, variance proportional to the square of the mean response. In this case, a more appropriate assumption than (2.3) would be

$$E(y_j) = f(x_j, \beta), \quad \text{Var}(y_j) = \sigma^2 \{f(x_j, \beta)\}^2, \quad (2.4)$$

where the scale parameter σ is the coefficient of variation. A general discussion of extensions of the classical model to accommodate heterogeneity of variance is given in the next section.

2.2.3 Generalizations of the classical framework

The classical nonlinear regression framework may be generalized to accommodate departures from the assumptions in a variety of ways. The following exposition is by no means exhaustive. We adopt the perspective taken in Chapters 2 and 3 of Carroll and Ruppert (1988) and Chapter 6 of Seber and Wild (1989); specifically, that of modeling systematic response variance and correlation patterns explicitly. We focus on this strategy because the hierarchical nonlinear model discussed in the remainder of the book adapts readily to account for these features. Other approaches are noted in section 2.7.