

Springer Series in Statistics

Advisors:

P. Bickel, P. Diggle, S. Fienberg, U. Gather,
I. Olkin, S. Zeger

Michael S. Hamada

Alyson G. Wilson

C. Shane Reese

Harry F. Martz

Bayesian Reliability



Springer

Michael S. Hamada
Los Alamos National Laboratory
Los Alamos, NM 87545, USA
hamada@lanl.gov

C. Shane Reese
Department of Statistics
Brigham Young University
Provo, UT 84602, USA
reese@stat.byu.edu

Alyson G. Wilson
Los Alamos National Laboratory
Los Alamos, NM 87545, USA
agw@lanl.gov

Harry F. Martz
Los Alamos National Laboratory
Los Alamos, NM 87545, USA
martz@rt66.com

ISSN 0172-7397

ISBN 978-0-387-77948-5

e-ISBN 978-0-387-77950-8

DOI: 10.1007/978-0-387-77950-8

Library of Congress Control Number: 2008930561

© 2008 Springer Science+Business Media, LLC

All rights reserved. This work may not be translated or copied in whole or in part without the written permission of the publisher (Springer Science+Business Media, LLC, 233 Spring Street, New York, NY 10013, USA), except for brief excerpts in connection with reviews or scholarly analysis. Use in connection with any form of information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed is forbidden.

The use in this publication of trade names, trademarks, service marks, and similar terms, even if they are not identified as such, is not to be taken as an expression of opinion as to whether or not they are subject to proprietary rights.

Printed on acid-free paper

springer.com

To Jung Hee, Christina, and Alexandra - M.S.H.

To Carol Ann - H.F.M.

To Wendy, Madison, Brittany, Bryon, and Mom - C.S.R.

To Greg - A.G.W.

Preface

In this book, we present modern methods and techniques for analyzing reliability data from a Bayesian perspective. The acceptance and application of Bayesian methods in virtually all branches of science and engineering have significantly increased over the past few decades. This increase is largely due to advances in simulation-based computational tools for implementing Bayesian methods. We extensively use such tools here.

We focus our attention on assessing the reliability of components and systems with particular attention to models containing explanatory variables. Such models include failure time regression models, accelerated testing models, and degradation models. We also pay special attention to Bayesian goodness-of-fit testing, model validation, reliability test design, and assurance test planning. Throughout the book we use Markov chain Monte Carlo (MCMC) algorithms for implementing Bayesian analyses. MCMC makes the Bayesian approach to reliability computationally feasible and conceptually straightforward; this is an important advantage in complex settings where classical approaches fail or become too difficult for practical implementation.

We intend this book to be primarily a reference collection of modern Bayesian methods in reliability for use by reliability practitioners. To this end, we have included more than 70 illustrative examples. Most have a real data component, and several of the corresponding datasets have not previously been published. We note, however, that space constraints have made it impractical to fully detail model diagnostics and goodness-of-fit procedures in all examples.

This book can also be used as a textbook for an undergraduate or graduate course in reliability. Therefore, we have included more than 165 exercises to further illustrate and emphasize text material. We base many of the exercises on real data. A solution manual for the exercises that also contains code for the examples is available for instructors at <http://www.springer.com>.

As a prerequisite, readers should have a basic knowledge of probability and statistics, as presented in a first course in applied statistics. In particular, prior familiarity with probability distributions, statistical estimation, and regression

analysis is useful. We present fundamental notions of reliability in Chap. 1, so prior knowledge of reliability concepts is not required. Basic calculus and matrix algebra concepts are also required.

Noteworthy highlights of the book include the following:

- Development and use of Bayesian goodness-of-fit and model selection methods,
- Introduction and use of Bayesian hierarchical models for reliability estimation,
- Consideration of a Bayesian fault tree analysis method supporting data acquisition at all levels in the tree,
- Bayesian networks in reliability analysis,
- Bayesian methods for analyzing both failure count and failure time data collected from repairable systems and the assessment of various related performance criteria,
- Estimation of reliability using information contained in explanatory variables,
- Bayesian approaches for designing and analyzing reliability improvement experiments,
- Bayesian methods for modeling and analyzing nondestructive and destructive degradation data,
- Illustration of a Bayesian approach for the optimal design of reliability experiments, and a
- Bayesian hierarchical approach to reliability assurance testing.

Of course, we have not covered all topics in reliability. For example, we have chosen not to cover topics like nonparametric methods in reliability (including hazard function and proportional hazards modeling), software and structural reliability, and certain topics related to repairable systems, such as maintenance. We also do not discuss probability plotting as a means for identifying a sampling distribution, mainly because this topic is already well covered in other books.

Chapter 1 develops the main definitions of reliability and introduces reliability and lifetime data. In Chap. 2, we cover basic concepts common to all Bayesian analyses, including the definitions and specifications of prior distributions, likelihood functions and sampling distributions, posterior distributions, and predictive distributions. Chapter 3 introduces the primary numerical, simulation-based tool for estimating these distributions: MCMC algorithms. We provide detailed examples to illustrate the two most common types of MCMC algorithms, the Gibbs sampler, and Metropolis-Hastings algorithm. We then introduce the notions of hierarchical modeling and empirical Bayes methods.

Reliability models and lifetime analyses for component-level data are presented and developed in Chap. 4. In this first applications chapter, we discuss diagnostics for addressing model fit and describe hierarchical models that facilitate the joint analyses of data collected from similar components.

In Chap. 5, we extend the models for component-level data to the system level. This extension requires us to specify logical relationships between the components in a system and how the functioning of the complete system depends on the functioning of each of its components. Probability models developed in Chap. 5 account for both dependent and independent components and multilevel data.

Chapter 6 develops a Bayesian treatment of the classical models for repairable systems: renewal processes and homogeneous and nonhomogeneous Poisson processes. We also consider some alternative models as well as Bayesian hierarchical adaptations of these common models. Several real-data examples address the reliability of highly parallel supercomputers.

Bayesian estimation methods for the standard regression models used in reliability are considered in Chap. 7. In particular, we consider linear, nonlinear, logistic, and Poisson regression models. We also present Bayesian methods for accelerated life testing models. The chapter also contains methodology for analyzing reliability improvement experiments.

Chapter 8 extends Bayesian methods to degradation data models. In addition to a general model for degradation data, we consider models that include both continuous and discrete covariates. We compare reliability estimates based on degradation data to those based on lifetime data. We also consider models for destructive degradation data, as well as an alternative stochastic process-based degradation model.

Chapter 9 presents methods for the optimal design of reliability experiments. These designs attempt to allocate resources in the most efficient way to meet specified experimental goals. These goals usually involve the quality of the inferences that can be made using experimental data.

Finally, in Chap. 10, we apply these ideas to design tests that assure, at some level of confidence, that a reliability-related quantity exceeds a specified requirement. Within the framework of Bayesian hierarchical models, we derive test plans for binomial, Poisson, and Weibull sampling distributions.

We use several existing statistical software packages for solving examples and exercises. One is the software package WinBUGS, which is a Windows-based implementation of BUGS (**B**ayesian inference **U**sing **G**ibbs **S**ampling). The package contains flexible software for analyzing complex statistical models using MCMC methods. It is available for free download at <http://www.mrc-bsu.cam.ac.uk/bugs/>. This program is relatively simple to use, and detailed examples of its implementation accompany the package.

YADAS (**Y**et **A**nother **D**ata **A**nalysis **S**ystem) is another Bayesian software system for doing MCMC calculations that is based entirely on the Metropolis and Metropolis-Hastings algorithms. It is written in Java and provides tools to implement nonstandard models. In several examples, we found it to be easier to use than WinBUGS. A detailed description of YADAS is available at <http://yadas.lanl.gov>, and it is also available for free download.

In many of the examples, we used the statistical software package R. Although it does not directly support Bayesian MCMC calculations, R is a language and environment for general statistical computing and graphics. It runs on a wide variety of platforms, including UNIX, Windows, and Mac operating systems, and is also available for free download at <http://www.r-project.org>.

We provide a list of acronyms in Appendix A. For convenient reference, Appendix B contains an extensive list of probability distributions and their properties. For each distribution, we define a standard form used throughout this book. For example, $X \sim \text{Beta}(\alpha, \beta)$ means that the random variable X has a beta distribution with parameters α and β . If we need to precisely indicate which random variable we are considering, we sometimes include it in the notation. For example, $\text{Beta}(x|\alpha, \beta)$ indicates that X is a random variable having a $\text{Beta}(\alpha, \beta)$ distribution. Throughout the book we use $\mathbf{P}(A)$ to denote the probability of the event A .

We are indebted to several people for their valuable help. Val Johnson contributed substantially throughout the writing of the book. Valerie Riedel painstakingly edited the original manuscript; and Megan Wyman, a later draft. Hazel Kutac provided invaluable word processing and editing support. Todd Graves provided support for developing YADAS code, as well as help on several of the research topics considered in Chap. 5. Brian Weaver assisted in preparing the distribution appendix and solutions manual. Finally, we thank Sallie Keller-McNulty and David Higdon for providing support and encouragement by allocating time for us to write the book.

Los Alamos, NM
Los Alamos, NM
Provo, UT
Los Alamos, NM
February 2008

Michael Hamada
Harry Martz
Shane Reese
Alyson Wilson

Contents

Preface	VII
1 Reliability Concepts	1
1.1 Defining Reliability	1
1.2 Measures of Random Variation	2
1.3 Examples of Reliability Data	10
1.3.1 Bernoulli Success/Failure Data	10
1.3.2 Failure Count Data	10
1.3.3 Lifetime/Failure Time Data	11
1.3.4 Degradation Data	12
1.4 Censoring	13
1.5 Bayesian Reliability Analysis	15
1.6 Related Reading	18
1.7 Exercises for Chapter 1	19
2 Bayesian Inference	21
2.1 Introductory Concepts	21
2.1.1 Maximum Likelihood Estimation	24
2.1.2 Classical Point and Interval Estimation for a Proportion	26
2.2 Fundamentals of Bayesian Inference	27
2.2.1 The Prior Distribution	28
2.2.2 Combining Data with Prior Information	30
2.3 Prediction	35
2.4 The Marginal Distribution of the Data and Bayes' Factors	36
2.5 A Lognormal Example	39
2.6 More on Prior Distributions	46
2.6.1 Noninformative and Diffuse Prior Distributions	46
2.6.2 Conjugate Prior Distributions	47
2.6.3 Informative Prior Distributions	47
2.7 Related Reading	49
2.8 Exercises for Chapter 2	49

3	Advanced Bayesian Modeling and Computational Methods	51
3.1	Introduction to Markov Chain Monte Carlo (MCMC)	51
3.1.1	Metropolis-Hastings Algorithms	52
3.1.2	Gibbs Sampler	60
3.1.3	Output Analysis	64
3.2	Hierarchical Models	68
3.2.1	MCMC Estimation of Hierarchical Model Parameters	71
3.2.2	Inference for Launch Vehicle Probabilities	71
3.3	Empirical Bayes	73
3.4	Goodness of Fit	76
3.5	Related Reading	82
3.6	Exercises for Chapter 3	82
4	Component Reliability	85
4.1	Introduction	85
4.2	Discrete Data Models for Reliability	86
4.2.1	Success/Failure Data	86
4.2.2	Failure Count Data	87
4.3	Failure Time Data Models for Reliability	90
4.3.1	Exponential Failure Times	91
4.3.2	Weibull Failure Times	97
4.3.3	Lognormal Failure Times	102
4.3.4	Gamma Failure Times	104
4.3.5	Inverse Gaussian Failure Times	105
4.3.6	Normal Failure Times	106
4.4	Censored Data	107
4.5	Multiple Units and Hierarchical Modeling	111
4.6	Model Selection	116
4.6.1	Bayesian Information Criterion	116
4.6.2	Deviance Information Criterion	117
4.6.3	Akaike Information Criterion	120
4.7	Related Reading	120
4.8	Exercises for Chapter 4	120
5	System Reliability	125
5.1	System Structure	125
5.1.1	Reliability Block Diagrams	126
5.1.2	Structure Functions	126
5.1.3	Minimal Path and Cut Sets	129
5.1.4	Fault Trees	131
5.2	System Analysis	135
5.2.1	Calculating System Reliability	135
5.2.2	Prior Distributions for Systems	138
5.2.3	Fault Trees with Bernoulli Data	141

5.2.4	Fault Trees with Lifetime Data	145
5.2.5	Bayesian Network Models	147
5.2.6	Models for Dependence	155
5.3	Related Reading	158
5.4	Exercises for Chapter 5	159
6	Repairable System Reliability	161
6.1	Introduction	161
6.1.1	Types of Data	162
6.1.2	Characteristics of System Repairs	162
6.2	Renewal Processes	163
6.3	Poisson Processes	165
6.3.1	Homogeneous Poisson Processes (HPPs)	167
6.4	Nonhomogeneous Poisson Processes (NHPPs)	170
6.4.1	Power Law Processes (PLPs)	170
6.4.2	Log-Linear Processes	176
6.5	Alternatives to NHPPs	176
6.5.1	Modulated Power Law Processes (MPLPs)	176
6.5.2	Piecewise Exponential Model (PEXP)	179
6.6	Goodness of Fit and Model Selection	180
6.7	Current Reliability and Other Performance Criteria	181
6.7.1	Current Reliability	181
6.7.2	Other Performance Criteria	182
6.8	Multiple-Unit Systems and Hierarchical Modeling	183
6.9	Availability	192
6.9.1	Other Data Types for Availability	194
6.9.2	Complex System Availability	196
6.10	Related Reading	198
6.11	Exercises for Chapter 6	199
7	Regression Models in Reliability	203
7.1	Introduction	203
7.1.1	Covariate Types	204
7.1.2	Covariate Relationships	205
7.2	Logistic Regression Models for Binomial Data	205
7.3	Poisson Regression Models for Count Data	215
7.4	Regression Models for Lifetime Data	221
7.5	Model Selection	228
7.6	Residual Analysis	229
7.7	Accelerated Life Testing	235
7.7.1	Common Accelerating Variables and Relationships	237
7.8	Reliability Improvement Experiments	243
7.9	Other Regression Situations	258
7.10	Related Reading	259
7.11	Exercises for Chapter 7	259

8	Using Degradation Data to Assess Reliability	271
8.1	Introduction	271
8.1.1	Comparison with Lifetime Data	278
8.2	More Complex Degradation Data Models	279
8.2.1	Reliability Function	281
8.3	Diagnostics for Degradation Data Models	283
8.4	Incorporating Covariates	287
8.4.1	Accelerated Degradation Testing	288
8.4.2	Improving Reliability Using Designed Experiments	295
8.5	Destructive Degradation Data	298
8.6	An Alternative Degradation Data Model Using Stochastic Processes	306
8.7	Related Reading	309
8.8	Exercises for Chapter 8	310
9	Planning for Reliability Data Collection	319
9.1	Introduction	319
9.2	Planning Criteria, Optimization, and Implementation	320
9.2.1	Optimization in Planning	321
9.2.2	Implementing the Simulation-Based Framework	323
9.3	Planning for Binomial Data	324
9.4	Planning for Lifetime Data	327
9.5	Planning Accelerated Life Tests	328
9.6	Planning for Degradation Data	330
9.7	Planning for System Reliability Data	331
9.8	Related Reading	339
9.9	Exercises for Chapter 9	339
10	Assurance Testing	343
10.1	Introduction	343
10.1.1	Classical Risk Criteria	345
10.1.2	Average Risk Criteria	345
10.1.3	Posterior Risk Criteria	346
10.2	Binomial Testing	348
10.2.1	Binomial Posterior Consumer's and Producer's Risks	349
10.2.2	Hybrid Risk Criterion	353
10.3	Poisson Testing	354
10.4	Weibull Testing	358
10.4.1	Single Weibull Population Testing	360
10.4.2	Combined Weibull Accelerated/Assurance Testing	364
10.5	Related Reading	368
10.6	Exercises for Chapter 10	369
A	Acronyms and Abbreviations	375

B	Special Functions and Probability Distributions	377
B.1	Greek Alphabet	377
B.2	Special Functions	377
B.2.1	Beta Function	377
B.2.2	Binomial Coefficient	378
B.2.3	Determinant	378
B.2.4	Factorial	378
B.2.5	Gamma Function	378
B.2.6	Incomplete Beta Function	378
B.2.7	Incomplete Beta Function Ratio	378
B.2.8	Indicator Function	379
B.2.9	Logarithm	379
B.2.10	Lower Incomplete Gamma Function	379
B.2.11	Standard Normal Cumulative Density Function	379
B.2.12	Standard Normal Probability Density Function	379
B.2.13	Trace	379
B.2.14	Upper Incomplete Gamma Function	379
B.3	Probability Distributions	380
B.3.1	Bernoulli	380
B.3.2	Beta	380
B.3.3	Binomial	382
B.3.4	Bivariate Exponential	382
B.3.5	Chi-squared	383
B.3.6	Dirichlet	383
B.3.7	Exponential	386
B.3.8	Extreme Value	386
B.3.9	Gamma	389
B.3.10	Inverse Chi-squared	389
B.3.11	Inverse Gamma	392
B.3.12	Inverse Gaussian	392
B.3.13	Inverse Wishart	392
B.3.14	Logistic	396
B.3.15	Lognormal	396
B.3.16	Multinomial	399
B.3.17	Multivariate Normal	399
B.3.18	Negative Binomial	399
B.3.19	Negative Log-Gamma	401
B.3.20	Normal	403
B.3.21	Pareto	403
B.3.22	Poisson	403
B.3.23	Poly-Weibull	403
B.3.24	Student's t	406
B.3.25	Uniform	408
B.3.26	Weibull	408
B.3.27	Wishart	411

References 413

Author Index 427

Subject Index 431

Reliability Concepts

This chapter introduces the fundamental definitions of reliability and gives examples of common types of reliability data.

1.1 Defining Reliability

There are many ways to define *reliability*. Colloquially, reliability is the property that a thing works when we want to use it. By necessity, more formal definitions of reliability must account for whether or not an item performs at or above a specified standard, how long it is able to perform at that standard, and the conditions under which it is operated. The reliability of an electrical switch, for example, may be defined as the probability that it successfully functions under a specified load and at a particular temperature. In contrast, reliability may be expressed as an explicit function of time. Defining the reliability of a pump in a nuclear power plant depends both on its environment and on its ability to provide a specified capacity over time.

As these examples illustrate, an operational definition of reliability must be specific enough to permit a clear distinction between items that are reliable and those that are not, but also must be sufficiently general to account for the complexities that arise in making this determination. In an effort to achieve both aims, the International Organization for Standardization (ISO) defines reliability as “the ability of an item to perform a required function, under given environmental and operating conditions and for a stated period of time” (ISO, 1986).

From this definition of reliability, we see that reliability analyses often involve the analysis of binary outcomes (i.e., success/failure data). However, in practice, it is often as important to analyze the time periods over which items or systems function. Such analyses are called *lifetime* or *failure time* analyses.

Lifetime analyses involve the analysis of positive, continuous-valued quantities (e.g., the length of time an item functions), and so require different

statistical models than analyses based on success/failure data. There are advantages and disadvantages to each type of analysis. Much of the information contained in experimental data may be lost when it is distilled into a success/failure format. For example, failures at 1 and 99 hours are regarded as equivalent if a “reliable” component must operate for 100 hours. On the other hand, accounting for the distribution of times when items fail requires additional assumptions in the statistical models, and these additional assumptions may be difficult to validate.

The notion of randomness is inherent to both types of analyses. For example, suppose that we would like to predict whether electrical switches from a particular production lot can successfully complete 10,000 on/off cycles. We choose 100 switches from the lot and test them to see whether they complete 10,000 cycles or not. This testing will generally not allow us to predict whether any particular electrical switch from the lot will complete 10,000 cycles or not. Instead, we want to estimate the probability that a switch selected at random from the lot will complete 10,000 cycles. Similarly, if we view this problem from a lifetime analysis perspective, we want to estimate the probability that a randomly selected switch fails on or before a particular cycle. That is, we might want to estimate the probability that a switch fails before its i th cycle, for $i = 1, \dots, 10,000$. In fact, there are a number of summaries that we might be interested in; some of these are described in Sect. 1.2.

The random nature of item failures requires us to choose a philosophy for performing statistical inference about an item’s reliability or lifetime. Throughout this book, we have chosen to adopt a Bayesian approach to statistical inference. In our view, the Bayesian approach toward inference offers many advantages. Among these, it allows us to pool information obtained from related experiments into the joint estimation of quantities of interest from each, and it allows us to incorporate expert opinion and subject matter expertise into the analysis of an experiment in a coherent way. Perhaps as importantly, it provides a remarkable degree of flexibility in modeling the phenomena that contribute to reliability and lifetime. With the advent of Markov chain Monte Carlo (MCMC) algorithms, fitting complicated statistical models to data and evaluating the uncertainty in fitted values is now almost routine.

To study these ideas in greater depth, we first need to establish some definitions for describing the properties of random quantities. We then use this expanded vocabulary to discuss a variety of simple examples that illustrate the range of experiments and types of data that can be analyzed using the techniques described in this book.

1.2 Measures of Random Variation

We define a *random variable* to be a function that maps the outcome of an experiment to a real number. For example, if we cycle a switch until it fails, the number of cycles before failure is a random variable.

The *sample space* S of an experiment is the set of all possible outcomes of the experiment. In a simple electrical switch experiment, the sample space for a single cycle is the set containing the events “operates” and “fails,” and a random variable for this experiment could be defined as 0 if the switch fails and 1 if it operates. Alternatively, we might perform an experiment that tests an item until it fails. The sample space of that experiment is any positive time, and we can define a random variable T as the lifetime of the item.

Much of reliability and lifetime analysis focuses on modeling the failure time distribution for an item — in other words, modeling the properties of a random variable like T . We can specify the properties of a random variable using different representations, all of which contain equivalent information. Each representation is useful in specific contexts. These representations include the *probability density* or *probability mass function*, the *reliability function*, the *cumulative distribution function*, and the *hazard function*.

For a discrete random variable X with sample space S , the *probability mass function* is a function, $m(x)$, that satisfies

$$m(x) \geq 0, \quad x \in S,$$

and

$$\sum_{x \in S} m(x) = 1.$$

For a continuous random variable T taking values on the real line, the *probability density function* is a function, $f(t)$, that satisfies

$$f(t) \geq 0, \quad -\infty < t < \infty,$$

and

$$\int_{-\infty}^{\infty} f(t) dt = 1.$$

Thus, any nonnegative function that integrates to 1 over the real line is a probability density function. For simplicity, we refer to both probability mass functions and probability density functions as probability density functions throughout the book.

Example 1.1 Exponential probability density function. Exponential random variables are widely used for modeling lifetimes. We say that the random variable T has an exponential distribution (or is an exponential random variable), and we write $T \sim \text{Exponential}(\lambda)$ if the probability density function for T is

$$\begin{aligned} f(t) &= \lambda e^{-\lambda t}, & t > 0, \quad \lambda > 0, \\ &= 0, & t \leq 0. \end{aligned} \tag{1.1}$$

This function meets the requirements for a probability density function because $f(t) \geq 0$ for all t , and $\int_{-\infty}^{\infty} f(t) dt = 1$ for any value of $\lambda > 0$. Figure 1.1

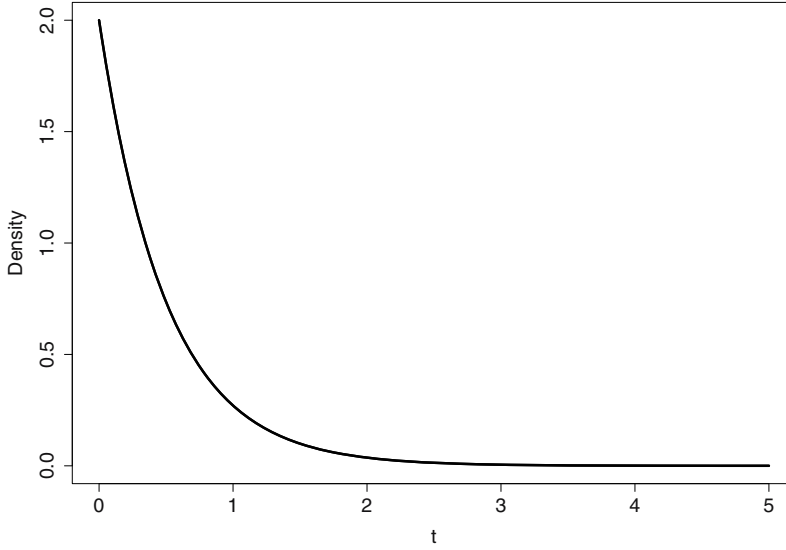


Fig. 1.1. The probability density function for an exponential random variable with $\lambda = 2$.

is a plot of the probability density function for an exponential random variable with $\lambda = 2$.

A second way to specify the properties of a random variable is through its *reliability function*, also known as the *survival function*. We define the reliability function as $R(t) = \mathbf{P}(T > t) = \int_t^\infty f(s)ds$, where $f(t)$ is a probability density function. Notice that the reliability function takes values in $[0, 1]$.

A third way to specify the properties of T is the *cumulative distribution function*. The cumulative distribution function defines the probability that a random variable takes on a value less than or equal to t . The cumulative distribution function is the complement of the reliability function, so it is also called the *unreliability function*. Mathematically,

$$F(t) = \mathbf{P}(T \leq t) = \int_{-\infty}^t f(s)ds.$$

Example 1.2 Exponential reliability and cumulative distribution functions. The reliability function for an exponential random variable is

$$R(t) = \mathbf{P}(T > t) = \int_t^\infty f(s)ds$$

$$\begin{aligned}
&= \int_t^{\infty} \lambda e^{-\lambda s} ds \\
&= e^{-\lambda t}.
\end{aligned}$$

Figure 1.2 is a plot of the reliability function for an exponential random variable with $\lambda = 2$.

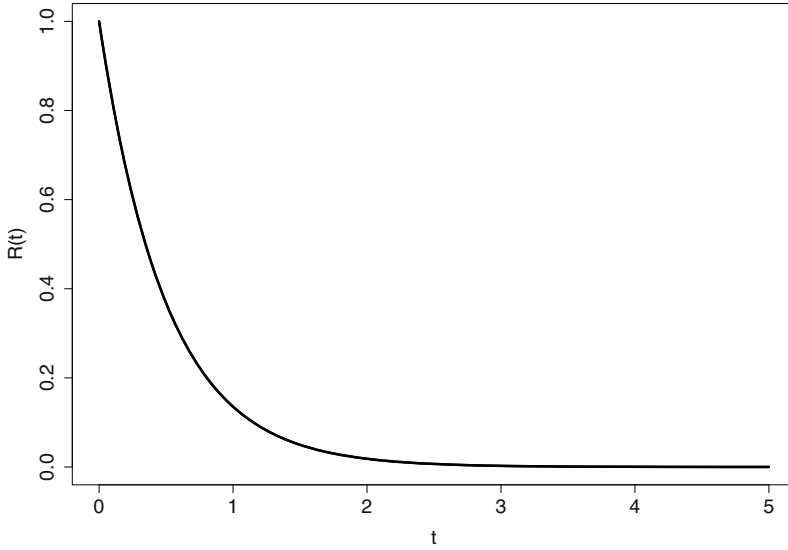


Fig. 1.2. The reliability function for an exponential random variable with $\lambda = 2$.

The cumulative distribution function for an exponential random variable is

$$\begin{aligned}
F(t) = \mathbf{P}(T \leq t) &= \int_{-\infty}^t f(s) ds \\
&= \int_0^t \lambda e^{-\lambda s} ds \\
&= 1 - e^{-\lambda t},
\end{aligned}$$

where $f(t)$ is the probability density function for an exponential random variable. Figure 1.3 is a plot of the cumulative distribution function for an exponential random variable with $\lambda = 2$.

Another way to specify the properties of a random variable is the *hazard function*, also called the *instantaneous failure rate function*. Suppose that we

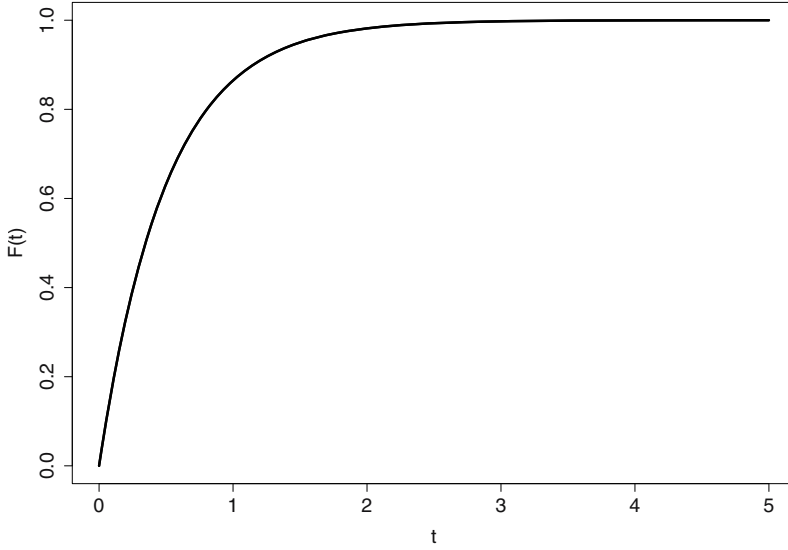


Fig. 1.3. The cumulative distribution function for an exponential random variable with $\lambda = 2$.

are interested in the probability that an item will fail in the time interval $[t, t + \Delta t]$ when we know that item is working at time t . Let $\mathbf{P}(A|B)$ denote the conditional probability of an event A , given that event B has occurred. From elementary probability, we know that $\mathbf{P}(A|B) = \mathbf{P}(A \cap B)/\mathbf{P}(B)$. We can write the probability that an item will fail in the time interval $[t, t + \Delta t]$, given that the item is working at time t , as

$$\mathbf{P}(t < T \leq t + \Delta t | T > t) = \frac{\mathbf{P}(t < T \leq t + \Delta t)}{P(T > t)} = \frac{F(t + \Delta t) - F(t)}{R(t)}.$$

If we want to know the failure rate, we divide by the length of the interval, Δt , and let $\Delta t \rightarrow 0$. This gives

$$\begin{aligned} h(t) &= \lim_{\Delta t \rightarrow 0} \frac{\mathbf{P}(t < T \leq t + \Delta t | T > t)}{\Delta t} \\ &= \lim_{\Delta t \rightarrow 0} \frac{F(t + \Delta t) - F(t)}{\Delta t} \frac{1}{R(t)}. \end{aligned}$$

The first term on the right-hand side is the derivative of the cumulative distribution function, $F(t)$, which is the probability density function, $f(t)$. Therefore,

$$h(t) = \frac{f(t)}{R(t)}.$$

We call $h(t)$ the *hazard function*. We can think of the hazard function as an item's propensity to fail in the next short interval of time, given that the item has survived to time t .

Figure 1.4 shows four of the most common types of hazard functions. These include:

1. *Increasing failure rate (IFR)*: the instantaneous failure rate (hazard) increases as a function of time. We expect to see an increasing number of failures for a given period of time.
2. *Decreasing failure rate (DFR)*: the instantaneous failure rate decreases as a function of time. We expect to see a decreasing number of failures for a given period of time.
3. *Bathtub failure rate (BFR)*: the instantaneous failure rate begins high because of early failures (“infant mortality” or “burn-in” failures), levels off for a period of time (“useful life”), and then increases (“wearout” or “aging” failures).
4. *Constant failure rate (CFR)*: the instantaneous failure rate is constant for the observed lifetime. We expect to see a relatively constant number of failures for a given period of time.

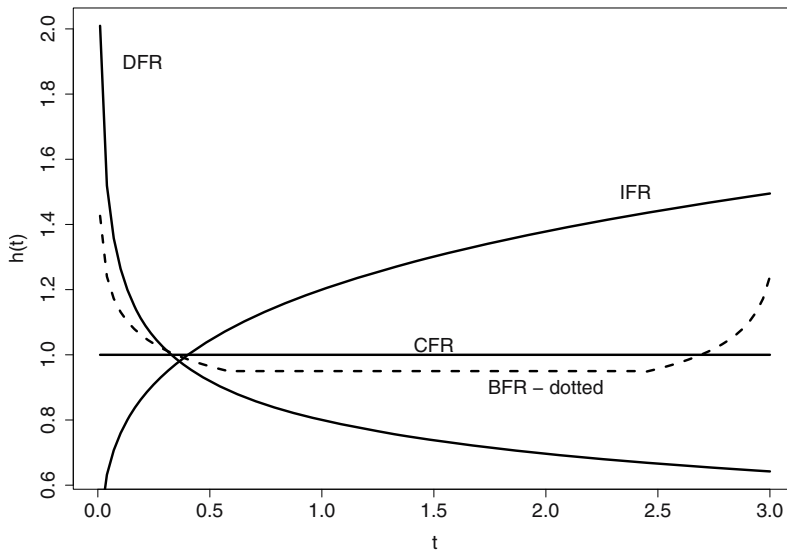


Fig. 1.4. Four different classifications of hazard functions. The dotted line represents the bathtub hazard function.

The *cumulative hazard function* is defined as $H(t) = \int_{-\infty}^t h(s)ds$, where $h(t)$ is the hazard function. The *average hazard rate* (AHR) between times t_1 and t_2 is defined as

$$AHR(t_1, t_2) = \frac{H(t_2) - H(t_1)}{t_2 - t_1}.$$

Example 1.3 Exponential hazard and cumulative hazard functions. The hazard function for an exponential random variable is

$$h(t) = \frac{f(t)}{R(t)} = \frac{\lambda e^{-\lambda t}}{e^{-\lambda t}} = \lambda.$$

Notice that this hazard function is constant, which implies that an item's propensity to fail in the next small unit of time does not change as the item ages.

The cumulative hazard function for an exponential random variable is

$$H(t) = \int_{-\infty}^t h(s)ds = \int_0^t \lambda ds = \lambda t.$$

For positive random variables, Table 1.1 summarizes the mathematical relationships between the probability density function $[f(t)]$, cumulative distribution function $[F(t)]$, reliability function $[R(t)]$, hazard function $[h(t)]$, and cumulative hazard function $[H(t)]$.

In applications, less complete descriptions of a random variable are also reported. Such descriptions usually involve the report of a mean, median, or variance of a random variable without specifying its complete distribution. For example, one such summary is the *mean time to failure* (MTTF), which is defined as

$$MTTF = E(T) = \int_{-\infty}^{\infty} tf(t)dt,$$

where $E(T)$ is the *expected value* of T . The MTTF is also called the *expected life*.

Example 1.4 MTTF for an exponential random variable. The MTTF for an exponential random variable is

$$\begin{aligned} MTTF &= \int_{-\infty}^{\infty} tf(t)dt \\ &= \int_0^{\infty} t\lambda e^{-\lambda t}dt \\ &= \frac{1}{\lambda}. \end{aligned}$$

Table 1.1. Relationships between the probability density function, cumulative distribution function, reliability function, hazard function, and cumulative hazard function, assuming $f(t) = 0$ for $t < 0$

	$f(t)$	$F(t)$	$R(t)$	$h(t)$	$H(t)$
$f(t)$	$f(t)$	$\frac{d}{dt} F(t)$	$-\frac{d}{dt} R(t)$	$h(t) \exp[-\int_0^t h(s)ds]$	$[\frac{d}{dt} H(t)] \exp[-H(t)]$
$F(t)$	$\int_0^t f(s)ds$	$F(t)$	$1 - R(t)$	$1 - \exp[-\int_0^t h(s)ds]$	$1 - \exp[-H(t)]$
$R(t)$	$\int_t^\infty f(s)ds$	$1 - F(t)$	$R(t)$	$\exp[-\int_0^t h(s)ds]$	$\exp[-H(t)]$
$h(t)$	$f(t) / \int_t^\infty f(s)ds$	$\frac{d}{dt} F(t) / [1 - F(t)]$	$-\frac{d}{dt} \log R(t)$	$h(t)$	$\frac{d}{dt} H(t)$
$H(t)$	$-\log[1 - \int_0^t f(s)ds]$	$-\log[1 - F(t)]$	$-\log[R(t)]$	$\int_0^t h(s)ds$	$H(t)$

Other summaries of the properties of a random variable are *quantiles* and *mean residual life*. A quantile is the inverse of the cumulative distribution function; it is the time by which a specified proportion of the population fails. Mathematically, the quantile q at which a proportion p of the population fails is the value q such that $F(q) = p$, or equivalently, $q = F^{-1}(p)$. If the cumulative distribution function is flat over some interval, we define q as the earliest time for which $F(q) = p$.

The *reliable life* is the time for which $100R\%$ of a population will survive, where R is a specified proportion between 0 and 1. For example, with $R = 0.5$, the reliable life is median of the lifetime distribution.

The *mean residual life* is the expected time to failure of a device that has survived to time t . We define the mean residual life as

$$M(t) = \frac{1}{R(t)} \int_t^{\infty} sf(t+s)ds,$$

where $R(t)$ is the reliability function, and $f(t)$ is a probability density function. Notice that $M(0) = E(T)$.

1.3 Examples of Reliability Data

In Sects. 1.1 and 1.2, we defined reliability and discussed various ways to summarize the distribution of the lifetime of an item. In this section, we give a few simple examples of the kinds of reliability data that are often encountered in practice. These include pass/fail, failure count, failure time, and degradation data.

1.3.1 Bernoulli Success/Failure Data

The simplest form of reliability data is “pass/fail” or *Bernoulli trial* data. This data can arise from simple “pass/fail” testing. In addition, it can be derived from lifetime data by letting the random variable $X(t) = 1$ (“pass” or “success”) if the item is functioning at time t , and $X(t) = 0$ (“fail”) if the item has failed by time t .

Table 1.2 contains the outcomes from a set of Bernoulli trials. These data are the launch outcomes of new aerospace vehicles conducted by “new” companies during the period 1980–2002. A total of 11 launches occurred; 3 were successes and 8 were failures. Reliability is the probability of a successful launch. We analyze these data in Chap. 2.

1.3.2 Failure Count Data

Bernoulli data can also be recorded as a function of time. *Failure count data* represent the number of failures that occur over a period of time. For example,

Table 1.2. New launch vehicle outcomes (Johnson et al., 2005)

Vehicle	Outcome
Pegasus	Success
Percheron	Failure
AMROC	Failure
Conestoga	Failure
Ariane 1	Success
India SLV-3	Failure
India ASLV	Failure
India PSLV	Failure
Shavit	Success
Taepodong	Failure
Brazil VLS	Failure

Gaver and O’Muircheartaigh (1987) provides the data shown in Table 1.3 on the number of pump failures x_i observed in t_i thousands of operating hours for 10 different systems at the Farley 1 United States commercial nuclear power plant. The random variable is the number of pump failures, X_i , and the reliability is the probability that no pump failures occur in a given period of time.

Table 1.3. Pump failure count data from Farley 1 U.S. nuclear power plant (number of failures x in t thousands of operating hours) (Gaver and O’Muircheartaigh, 1987)

System	x_i (failures)	t_i (thousand hours)
1	5	94.320
2	1	15.720
3	5	62.880
4	14	125.760
5	3	5.240
6	19	31.440
7	1	1.048
8	1	1.048
9	4	2.096
10	22	10.480

We consider this dataset in Chap. 10; we present a general discussion of failure count data in Chap. 4.

1.3.3 Lifetime/Failure Time Data

Table 1.4 presents an example of *lifetime data*. This dataset comprises 11 observed failure times and 55 times when a test was suspended (censored) before item failure occurred (*) for a 4.5 roller bearing in a set of J-52 engines

from EA-6B Prowler aircraft (Muller, 2003). Degradation of the 4.5 roller bearing in the J-52 engine has caused in-flight engine failures. The random variable is the time to failure T , and reliability is the probability that a roller bearing does not fail before time t . We analyze these data in Chap. 4.

Table 1.4. Roller bearing lifetime data (in operating hours) for the Prowler attack aircraft (Muller, 2003)

Failure Times (operating hours)		
1,085*	1,795*	100*
1,500*	1,890	1,628
1,390*	1,145*	759*
152*	1,380*	246*
971*	61*	861*
966*	1,165*	462*
997*	437*	1,079*
887*	1,152*	1,199*
977*	159*	424*
1,022*	3,428*	763*
2,087*	555*	1,297*
646	727*	2,238*
820*	2,294*	1,388
897	663*	1,153*
810*	1,427*	2,892*
80*	951	2,153*
1,167	767*	853*
711	546*	911*
1,203	736*	2,181
85*	917*	1,042*
1,070*	2,871*	799*
719*	1,231*	750

1.3.4 Degradation Data

In some applications, it is useful to measure the *degradation* of an item rather than its lifetime. This is particularly true when items have relatively long lifetimes, which makes experimentation difficult. Many item failures can be traced to an underlying degradation mechanism; for example, when degradation reaches a certain critical threshold, the item fails. Studying degradation mechanisms can often tell us how to improve the reliability of an item.

Table 1.5 presents an example of degradation data adapted from Chow and Shao (1991). One aspect of developing a new drug is to determine its shelf life. Since the potency of a drug degrades over time, its lifetime or shelf life is defined to be the time when its potency reaches 90% of its stated potency.

The random variable is the potency of the drug, and reliability at time t is the probability that the drug has not degraded past 90% of its stated potency by t . We analyze these data in Chap. 8.

Table 1.5. Drug potency degradation data (in percent of stated potency)

Batch	Time (months)				Batch	Time (months)			
	0	12	24	36		0	12	24	36
1	99.9	98.9	95.9	92.9	13	99.8	98.8	93.8	89.8
2	101.1	97.1	94.1	91.1	14	100.1	99.1	93.1	90.1
3	100.3	98.3	95.3	92.3	15	100.7	98.7	93.7	91.7
4	100.8	96.8	94.8	90.8	16	100.3	98.3	96.3	93.3
5	100.0	98.0	96.0	92.0	17	100.2	98.2	97.2	94.2
6	100.1	98.1	98.1	95.1	18	99.8	97.8	95.8	90.8
7	99.6	98.6	96.6	92.6	19	100.8	98.8	95.8	94.8
8	100.4	99.4	96.4	95.4	20	100.0	98.0	96.0	92.0
9	100.9	98.9	96.9	96.9	21	99.6	99.6	92.6	88.6
10	100.5	99.5	94.5	93.5	22	100.2	98.2	97.2	94.2
11	101.1	98.1	93.1	91.1	23	99.8	97.8	95.8	90.8
12	100.9	97.9	95.9	93.9	24	100.0	99.0	95.0	92.0

1.4 Censoring

One of the features of reliability data is the presence of *censoring*. Lifetime data are censored when the exact failure time for a specific item is unknown. There are several types of censoring, including left, right, interval, time, and failure censoring.

Left censoring occurs when an item fails before the first inspection. For example, suppose that an experiment tests the lifetime of a new battery. A set of 500 batteries are tested at 8:00 a.m. every day for 90 days to determine whether each battery is still usable. Suppose the test starts at midnight. The data for any battery that fails before 8:00 a.m. on the first day are *left censored*.

Right censoring occurs when an item has not failed by the last inspection. Consider our battery example. Data for any battery that has not failed by 8 a.m. on the 90th day are *right censored*.

Both left and right censoring are special cases of *interval censoring*. Interval censoring occurs when an item's failure time is only known to be in an interval, (t_i, t_{i+1}) . If an observation is left censored at t , then its failure time is in $(0, t)$. If an observation is right censored at t , then its failure time is in (t, ∞) . In our battery example, the failure times are *interval censored* because they can only be determined to within a 24-hour interval.

Other categories of censoring describe the cause of the censoring. *Type I censoring* or *time censoring* occurs when we remove unfailed items from

testing at a prespecified time; in other words, the test ends after a fixed amount of time. In our battery example, data from any battery that has not failed before 8 a.m. on the 90th day are *time censored*.

Type II censoring (also called *failure censoring* or *item censoring*) occurs when a test ends after a specific number of failures have occurred. Suppose that we are testing a set of 150 air conditioners. We decide that the test will end after 100 failures. The data for the 50 air conditioners that do not fail are *failure censored*.

Type III censoring combines Type I (time) and Type II (failure) censoring. Type III censoring occurs when we set both time and failure criteria and end the experiment when either the time or failure criterion (whichever comes first) has been met. For example, suppose that the test of our 150 air conditioners must end after 1 year or 100 failures, whichever comes first. Data for any air conditioners that have not failed when the experiment ends are *Type III censored*.

Systematic multiple censoring (also called *Type IV censoring*) occurs when items enter into an experiment over a period of time. For example, suppose that our air conditioners are entered into the study as they come off the production line over a six-month period. Suppose that the study ends after 50 air conditioners have failed. The data from any air conditioner still working after the trial ends are subject to *systematic multiple censoring*.

Random right censoring occurs when we remove an item from a test because of a failure that is not of interest. For example, suppose that we are testing our air conditioners for wearout failures, but partway through the test, one falls off the test stand and breaks. The air conditioner's datum for the time to failure is *random right censored*.

We usually assume that the censoring and survival times are independent. This is known as *independent censoring* or *noninformative censoring*. For example, if an air conditioner is removed from our study because it is sold to a customer, then the censoring of its failure time is independent of its survival time.

Because Bayesian modeling uses the observed lifetime in its analyses, censoring mechanisms are easily addressed. Suppose that we observe that an item has failed before time t_L , and its lifetime data are left censored. We know that its lifetime is in $[0, t_L]$. The probability of observing a failure in this interval is

$$\begin{aligned}\mathbf{P}(T \leq t_L) &= \int_0^{t_L} f(t)dt \\ &= F(t_L).\end{aligned}$$

As we will see later, $F(t_L)$ represents this item's contribution to the likelihood function for estimating the parameters of $f(\cdot)$, and the cause of the censoring does not matter. Similar derivations lead to the expressions displayed in Table 1.6. These probabilities are central to Bayesian and likelihood-based analyses and represent all information provided by the censored data.

Table 1.6. The probability of observing a failure in censored and uncensored data

Type of Observation	Failure Time	Contribution
Uncensored	$T = t$	$f(t)$
Left censored	$T \leq t_L$	$F(t_L)$
Interval censored	$t_L < T \leq t_R$	$F(t_R) - F(t_L)$
Right censored	$T > t_R$	$1 - F(t_R)$

1.5 Bayesian Reliability Analysis

The acceptance and applicability of Bayesian methods have increased in recent years. Today, with advances in computation and methodology, researchers are using Bayesian methods to solve an increasing variety of complex problems. In many applications, Bayesian methods provide important computational and methodological advantages over classical techniques.

This book focuses on Bayesian reliability analysis, which includes the topics of modeling, computation, sensitivity analysis, and model checking. While reading the examples in the book, it is important to keep the following in mind:

Every attempt to use mathematics to study some real phenomena must begin with building a mathematical model of these phenomena. Of necessity, the model simplifies matters to a greater or lesser extent and a number of details are ignored. The success depends on whether or not the details ignored are really unimportant in the development of the phenomena studied. The solution of the mathematical problem may be correct and yet it may be in violent conflict with realities simply because the original assumptions of the mathematical model diverge essentially from the conditions of the practical problem considered. Beforehand, it is impossible to predict with certainty whether or not a given mathematical model is adequate. To find this out, it is necessary to deduce a number of consequences of the model and to compare them with observation. (Neyman, 1949, p. 22)

In statistical analyses, we must make trade-offs between selecting a model that is sufficiently simple to be readily interpretable and selecting one that is sufficiently rich to capture the essential features of the problem. In Bayesian reliability analysis, the statistical model consists of two parts: the *likelihood function* and the *prior distribution*. The likelihood function is typically constructed from the *sampling distribution* of the data, defined by the probability density function assumed for the data. For example, Eq. 1.1 could be a sampling distribution for lifetime data. The sampling distribution usually contains unknown parameters, such as λ in Eq. 1.1. Once we perform the experiment and observe its outcome, we regard the sampling distribution as a function of the unknown parameters. This function (or any function proportional to it) is called the *likelihood function*. Bayesian inference is the only framework for