



Version 16

Reliability and Survival Methods

*"The real voyage of discovery consists not in seeking new landscapes,
but in having new eyes."*

Marcel Proust

JMP, A Business Unit of SAS
SAS Campus Drive
Cary, NC 27513

16.0

The correct bibliographic citation for this manual is as follows: SAS Institute Inc. 2020–2021. *JMP® 16 Reliability and Survival Methods*. Cary, NC: SAS Institute Inc.

JMP® 16 Reliability and Survival Methods

Copyright © 2020–2021, SAS Institute Inc., Cary, NC, USA

All rights reserved. Produced in the United States of America.

U.S. Government License Rights; Restricted Rights: The Software and its documentation is commercial computer software developed at private expense and is provided with RESTRICTED RIGHTS to the United States Government. Use, duplication or disclosure of the Software by the United States Government is subject to the license terms of this Agreement pursuant to, as applicable, FAR 12.212, DFAR 227.7202-1(a), DFAR 227.7202-3(a) and DFAR 227.7202-4 and, to the extent required under U.S. federal law, the minimum restricted rights as set out in FAR 52.227-19 (DEC 2007). If FAR 52.227-19 is applicable, this provision serves as notice under clause (c) thereof and no other notice is required to be affixed to the Software or documentation. The Government's rights in Software and documentation shall be only those set forth in this Agreement.

SAS Institute Inc., SAS Campus Drive, Cary, North Carolina 27513-2414.

March 2021

SAS® and all other SAS Institute Inc. product or service names are registered trademarks or trademarks of SAS Institute Inc. in the USA and other countries. ® indicates USA registration.

Other brand and product names are trademarks of their respective companies.

SAS software may be provided with certain third-party software, including but not limited to open-source software, which is licensed under its applicable third-party software license agreement. For license information about third-party software distributed with SAS software, refer to <http://support.sas.com/thirdpartylicenses>.

Get the Most from JMP

Whether you are a first-time or a long-time user, there is always something to learn about JMP.

Visit JMP.com to find the following:

- live and recorded webcasts about how to get started with JMP
- video demos and webcasts of new features and advanced techniques
- details on registering for JMP training
- schedules for seminars being held in your area
- success stories showing how others use JMP
- the JMP user community, resources for users including examples of add-ins and scripts, a forum, blogs, conference information, and so on

<https://www.jmp.com/getstarted>

Contents

Reliability and Survival Methods

1	Learn about JMP	13
	Documentation and Additional Resources	
	Formatting Conventions in JMP Documentation	15
	JMP Help	16
	JMP Documentation Library	16
	Additional Resources for Learning JMP	22
	JMP Tutorials	22
	Sample Data Tables	22
	Learn about Statistical and JSL Terms	23
	Learn JMP Tips and Tricks	23
	JMP Tooltips	23
	JMP User Community	24
	Free Online Statistical Thinking Course	24
	JMP New User Welcome Kit	24
	Statistics Knowledge Portal	24
	JMP Training	24
	JMP Books by Users	25
	The JMP Starter Window	25
	JMP Technical Support	25
2	Introduction to Reliability and Survival	27
	Lifetime and Failure Analysis	
3	Life Distribution	29
	Fit Distributions to Lifetime Data	
	Overview of the Life Distribution Platform	31
	Example of the Life Distribution Platform	31
	Launch the Life Distribution Platform	34
	Life Distribution Report	37
	Event Plot	38
	Compare Distributions	41
	Life Distribution: Statistics	45
	Life Distribution Report Options	52

Mixture	54
Fit Competing Risk Mixture	58
Weibayes Report	59
Competing Cause Report	60
Competing Cause Workflow	60
Competing Cause Model	61
Cause Combination	61
Competing Cause: Statistics	62
Individual Causes	63
Competing Cause Report Options	64
Life Distribution - Compare Groups Report	66
Compare Groups: Statistics	66
Individual Group	67
Life Distribution - Compare Groups Report Options	67
Additional Examples of the Life Distribution Platform	69
Omit Competing Causes	69
Change the Scale	71
Examine the Same Distribution across Groups	73
Weibayes Estimates	75
Fit Mixture Example	77
Fit Competing Risk Mixture Example	80
Statistical Details for the Life Distribution Platform	82
Distributions	82
Competing Cause Details	97
Notes Regarding Median Rank Regression	102
4 Fit Life by X.....	105
Fit Single-Factor Models to Time-to Event Data	
Overview of the Fit Life by X Platform	107
Example of the Fit Life by X Platform	108
Launch the Fit Life by X Platform	110
The Fit Life by X Report	112
Summary of Data	113
Scatterplot	113
Nonparametric Overlay	115
Comparisons	116
Results	120
Custom Relationship	131
Fit Life by X Platform Options	132
Additional Examples of the Fit Life by X Platform	133

	Capacitor Example	134
	Custom Relationship Example	135
5	Cumulative Damage	139
	Model Product Deterioration with Variable Stress over Time	
	Overview of the Cumulative Damage Platform	141
	Example of the Cumulative Damage Platform	141
	Launch the Cumulative Damage Platform	145
	Time to Event	146
	Stress Pattern	147
	The Cumulative Damage Report	150
	Event Plot	150
	Stress Patterns Report	150
	Model List	151
	Model Results	151
	Cumulative Damage Platform Options	152
	Stress Patterns Options	152
	Additional Example of the Cumulative Damage Platform	154
	Example of Simulating New Data	154
6	Recurrence Analysis	157
	Model the Frequency or Cost of Recurrent Events over Time	
	Overview of the Recurrence Analysis Platform	159
	Example of the Recurrence Analysis Platform	159
	Launch the Recurrence Analysis Platform	162
	Recurrence Analysis Platform Options	163
	Fit Model	164
	Additional Examples of the Recurrence Analysis Platform	168
	Bladder Cancer Recurrences Example	169
	Diesel Ship Engines Example	172
7	Degradation	177
	Model Product Deterioration over Time	
	Overview of the Degradation Platform	179
	Example of the Degradation Platform	179
	Launch the Degradation Platform	181
	The Degradation Reports	183
	Model Specification	186
	Simple Linear Path	186
	Nonlinear Path	188
	Inverse Prediction	196

Prediction Graph	198
Degradation Platform Options	199
Model Reports	201
Model Lists	201
Reports	202
Destructive Degradation	203
Stability Analysis	208
8 Destructive Degradation	213
Model Product Deterioration over Time	
Example of the Destructive Degradation Platform	215
Launch the Destructive Degradation Platform	221
Specify Two Y Columns	222
The Destructive Degradation Plot Options and Models	223
Plot Options	224
Models	225
The Destructive Degradation Report	229
Model List	230
Model Outlines	230
Destructive Degradation Platform Options	233
Statistical Details for the Destructive Degradation Platform	234
9 Reliability Forecast	235
Forecast Product Failure Using Production and Failure Data	
Overview of the Reliability Forecast Platform	237
Example Using the Reliability Forecast Platform	237
Launch the Reliability Forecast Platform	241
The Reliability Forecast Report	244
Observed Data Report	244
Life Distribution Report	244
Forecast Report	245
Reliability Forecast Platform Options	249
Additional Example Using the Reliability Forecast Platform	250
10 Reliability Growth	253
Model System Reliability as Changes Are Implemented	
Overview of the Reliability Growth Platform	255
Example Using the Reliability Growth Platform	256
Launch the Reliability Growth Platform	260
Launch Window Roles	261
Data Table Structure	262

The Reliability Growth Report	265
Model Descriptions and Feasibilities	265
Observed Data Report	265
Reliability Growth Platform Options	268
Model Reports	269
Crow AMSAA	270
Crow AMSAA with Modified MLE	276
Fixed Parameter Crow AMSAA	277
Piecewise Weibull NHPP	278
Reinitialized Weibull NHPP	282
Piecewise Weibull NHPP Change Point Detection	285
Additional Examples of the Reliability Growth Platform	286
Piecewise NHPP Weibull Model Fitting with Interval-Censored Data	286
Piecewise Weibull NHPP Change Point Detection with Time in Dates Format	289
Statistical Details for the Reliability Growth Platform	291
Statistical Details for the Crow-AMSAA Report	291
Statistical Details for the Piecewise Weibull NHPP Change Point Detection Report ..	292
11 Reliability Block Diagram	293
Engineer System Reliabilities	
Overview of the Reliability Block Diagram Platform	295
Example Using the Reliability Block Diagram Platform	295
Add Components	297
Align Shapes	299
Connect Shapes	300
Configure Components	301
The Reliability Block Diagram Window	303
Preview Window	304
Reliability Block Diagram Platform Options	305
Workspace Options	308
Configuration Settings	309
Distribution Configurations	309
Specify a Nonparametric Distribution	311
Options for Design and Library Items	312
Profilers	313
Component Importance and Time to Failure	315
Component Plots	317
Print Algebraic Reliability Formula	320
Generate Algebraic Expression Data Table	321
Clone and Delete	321

12	Repairable Systems Simulation	323
	Estimate Outage Time for a Complex System	
	Overview of the Repairable Systems Simulation Platform	325
	Example Using the Repairable Systems Simulation Platform	325
	The Repairable Systems Simulation Window	327
	System Diagram	328
	Shapes Toolbar	329
	Configuration Panel	330
	Add Event Panel	333
	Add Action Panel	334
	Repairable Systems Simulation Platform Options	337
	Options for Block Items	338
	Distribution Options	338
	Specify a Nonparametric or Estimated Distribution	339
	Event Settings	340
	Action Settings	341
	Repairable Systems Simulation Results	342
	Results Table	342
	Results Explorer	343
	Repairable Systems Simulation Results Options	345
	RSS Results Explorer Point Estimation Profilers	346
13	Survival Analysis	347
	Analyze Survival Time Data	
	Overview of the Survival Analysis Platform	349
	Example of Survival Analysis	350
	Launch the Survival Platform	352
	The Survival Plot	353
	Survival Platform Options	354
	Exponential, Weibull, and Lognormal Plots and Fits	356
	Fitted Distribution Plots	360
	Competing Causes	361
	Additional Examples of the Survival Platform	362
	Example of Survival Analysis	362
	Example of Competing Causes	364
	Example of Interval Censoring	367
	Statistical Reports for Survival Analysis	368
14	Fit Parametric Survival	371
	Fit Survival Data Using Regression Models	
	Overview of the Fit Parametric Survival Platform	373

Example of Parametric Regression Survival Fitting	373
Launch the Fit Parametric Survival Platform	375
The Parametric Survival Fit Report	377
The Parametric Survival - All Distributions Report	378
Parametric Competing Cause Report	380
Fit Parametric Survival Options	380
Nonlinear Parametric Survival Models	382
Additional Examples of Fitting Parametric Survival	383
Arrhenius Accelerated Failure LogNormal Model	383
Interval-Censored Accelerated Failure Time Model	387
Analyze Censored Data Using the Nonlinear Platform	388
Left-Censored Data	390
Weibull Loss Function Using the Nonlinear Platform	391
Fitting Simple Survival Distributions Using the Nonlinear Platform	394
Statistical Details for the Fit Parametric Survival Platform	397
Loss Formulas for Survival Distributions	397
15 Fit Proportional Hazards	399
Fit Survival Data Using Semiparametric Regression Models	
Overview of the Fit Proportional Hazards Platforms	401
Example of the Fit Proportional Hazards Platform	401
Risk Ratios for One Nominal Effect with Two Levels	403
Launch the Fit Proportional Hazards Platform	404
The Fit Proportional Hazards Report	406
Fit Proportional Hazards Options	407
Example Using Multiple Effects and Multiple Levels	408
Risk Ratios for Multiple Effects and Multiple Levels	411
A References	413
B Technology License Notices	417

Chapter **1**

Learn about JMP

Documentation and Additional Resources

Learn about JMP documentation, such as book conventions, descriptions of each JMP document, the Help system, and where to find additional support.

Contents

Formatting Conventions in JMP Documentation.....	15
JMP Help	16
JMP Documentation Library	16
Additional Resources for Learning JMP	22
JMP Tutorials	22
Sample Data Tables	22
Learn about Statistical and JSL Terms	23
Learn JMP Tips and Tricks.....	23
JMP Tooltips.....	23
JMP User Community	24
Free Online Statistical Thinking Course	24
JMP New User Welcome Kit	24
Statistics Knowledge Portal.....	24
JMP Training	24
JMP Books by Users	25
The JMP Starter Window	25
JMP Technical Support.....	25

Formatting Conventions in JMP Documentation

These conventions help you relate written material to information that you see on your screen:

- Sample data table names, column names, pathnames, filenames, file extensions, and folders appear in Helvetica (or sans-serif online) font.
- Code appears in *Lucida Sans Typewriter* (or monospace online) font.
- Code output appears in *Lucida Sans Typewriter* italic (or monospace italic online) font and is indented farther than the preceding code.
- **Helvetica bold** formatting (or bold sans-serif online) indicates items that you select to complete a task:
 - buttons
 - check boxes
 - commands
 - list names that are selectable
 - menus
 - options
 - tab names
 - text boxes
- The following items appear in italics:
 - words or phrases that are important or have definitions specific to JMP
 - book titles
 - variables
- Features that are for JMP Pro only are noted with the JMP Pro icon . For an overview of JMP Pro features, visit <https://www.jmp.com/software/pro>.

Note: Special information and limitations appear within a Note.

Tip: Helpful information appears within a Tip.

JMP Help

JMP Help in the Help menu enables you to search for information about JMP features, statistical methods, and the JMP Scripting Language (or *JSL*). You can open JMP Help in several ways:

- Search and view JMP Help on Windows by selecting **Help > JMP Help**.
- On Windows, press the F1 key to open the Help system in the default browser.
- Get help on a specific part of a data table or report window. Select the Help tool  from the **Tools** menu and then click anywhere in a data table or report window to see the Help for that area.
- Within a JMP window, click the **Help** button.

Note: The JMP Help is available for users with Internet connections. Users without an Internet connection can search all books in a PDF file by selecting **Help > JMP Documentation Library**. See “[JMP Documentation Library](#)” on page 16 for more information.

JMP Documentation Library

The Help system content is also available in one PDF file called *JMP Documentation Library*. Select **Help > JMP Documentation Library** to open the file. If you prefer searching individual PDF files of each document in the JMP library, download the files from <https://www.jmp.com/documentation>.

The following table describes the purpose and content of each document in the JMP library.

Document Title	Document Purpose	Document Content
<i>Discovering JMP</i>	If you are not familiar with JMP, start here.	Introduces you to JMP and gets you started creating and analyzing data. Also learn how to share your results.
<i>Using JMP</i>	Learn about JMP data tables and how to perform basic operations.	Covers general JMP concepts and features that span across all of JMP, including importing data, modifying columns properties, sorting data, and connecting to SAS.

Document Title	Document Purpose	Document Content
<i>Basic Analysis</i>	Perform basic analysis using this document.	Describes these Analyze menu platforms: • Distribution • Fit Y by X • Tabulate • Text Explorer Covers how to perform bivariate, one-way ANOVA, and contingency analyses through Analyze > Fit Y by X. How to approximate sampling distributions using bootstrapping and how to perform parametric resampling with the Simulate platform are also included.
<i>Essential Graphing</i>	Find the ideal graph for your data.	Describes these Graph menu platforms: • Graph Builder • Scatterplot 3D • Contour Plot • Bubble Plot • Parallel Plot • Cell Plot • Scatterplot Matrix • Ternary Plot • Treemap • Chart • Overlay Plot The book also covers how to create background and custom maps.
<i>Profilers</i>	Learn how to use interactive profiling tools, which enable you to view cross-sections of any response surface.	Covers all profilers listed in the Graph menu. Analyzing noise factors is included along with running simulations using random inputs.

Document Title	Document Purpose	Document Content
<i>Design of Experiments Guide</i>	Learn how to design experiments and determine appropriate sample sizes.	Covers all topics in the DOE menu.
<i>Fitting Linear Models</i>	Learn about Fit Model platform and many of its personalities.	Describes these personalities, all available within the Analyze menu Fit Model platform: • Standard Least Squares • Stepwise • Generalized Regression • Mixed Model • MANOVA • Loglinear Variance • Nominal Logistic • Ordinal Logistic • Generalized Linear Model

Document Title	Document Purpose	Document Content
<i>Predictive and Specialized Modeling</i>	Learn about additional modeling techniques.	Describes these Analyze > Predictive Modeling menu platforms: • Neural • Partition • Bootstrap Forest • Boosted Tree • K Nearest Neighbors • Naive Bayes • Support Vector Machines • Model Comparison • Model Screening • Make Validation Column • Formula Depot Describes these Analyze > Specialized Modeling menu platforms: • Fit Curve • Nonlinear • Functional Data Explorer • Gaussian Process • Time Series • Matched Pairs Describes these Analyze > Screening menu platforms: • Modeling Utilities • Response Screening • Process Screening • Predictor Screening • Association Analysis • Process History Explorer

Document Title	Document Purpose	Document Content
<i>Multivariate Methods</i>	Read about techniques for analyzing several variables simultaneously.	Describes these Analyze > Multivariate Methods menu platforms: • Multivariate • Principal Components • Discriminant • Partial Least Squares • Multiple Correspondence Analysis • Structural Equation Models • Factor Analysis • Multidimensional Scaling • Item Analysis Describes these Analyze > Clustering menu platforms: • Hierarchical Cluster • K Means Cluster • Normal Mixtures • Latent Class Analysis • Cluster Variables
<i>Quality and Process Methods</i>	Read about tools for evaluating and improving processes.	Describes these Analyze > Quality and Process menu platforms: • Control Chart Builder and individual control charts • Measurement Systems Analysis • Variability / Attribute Gauge Charts • Process Capability • Model Driven Multivariate Control Chart • Legacy Control Charts • Pareto Plot • Diagram • Manage Spec Limits • OC Curves

Document Title	Document Purpose	Document Content
<i>Reliability and Survival Methods</i>	Learn to evaluate and improve reliability in a product or system and analyze survival data for people and products.	Describes these Analyze > Reliability and Survival menu platforms: • Life Distribution • Fit Life by X • Cumulative Damage • Recurrence Analysis • Degradation • Destructive Degradation • Reliability Forecast • Reliability Growth • Reliability Block Diagram • Repairable Systems Simulation • Survival • Fit Parametric Survival • Fit Proportional Hazards
<i>Consumer Research</i>	Learn about methods for studying consumer preferences and using that insight to create better products and services.	Describes these Analyze > Consumer Research menu platforms: • Categorical • Choice • MaxDiff • Uplift • Multiple Factor Analysis
<i>Scripting Guide</i>	Learn about taking advantage of the powerful JMP Scripting Language (JSL).	Covers a variety of topics, such as writing and debugging scripts, manipulating data tables, constructing display boxes, and creating JMP applications.
<i>JSL Syntax Reference</i>	Read about many JSL functions on functions and their arguments, and messages that you send to objects and display boxes.	Includes syntax, examples, and notes for JSL commands.

Additional Resources for Learning JMP

In addition to reading JMP help, you can also learn about JMP using the following resources:

- [“JMP Tutorials”](#)
- [“Sample Data Tables”](#)
- [“Learn about Statistical and JSL Terms”](#)
- [“Learn JMP Tips and Tricks”](#)
- [“JMP Tooltips”](#)
- [“JMP User Community”](#)
- [“Free Online Statistical Thinking Course”](#)
- [“JMP New User Welcome Kit”](#)
- [“Statistics Knowledge Portal”](#)
- [“JMP Training”](#)
- [“JMP Books by Users”](#)
- [“The JMP Starter Window”](#)

JMP Tutorials

You can access JMP tutorials by selecting **Help > Tutorials**. The first item on the **Tutorials** menu is **Tutorials Directory**. This opens a new window with all the tutorials grouped by category.

If you are not familiar with JMP, start with the **Beginners Tutorial**. It steps you through the JMP interface and explains the basics of using JMP.

The rest of the tutorials help you with specific aspects of JMP, such as designing an experiment and comparing a sample mean to a constant.

Sample Data Tables

All of the examples in the JMP documentation suite use sample data. Select **Help > Sample Data Library** to open the sample data directory.

To view an alphabetized list of sample data tables or view sample data within categories, select **Help > Sample Data**.

Sample data tables are installed in the following directory:

On Windows: C:\Program Files\SAS\JMP\16\Samples\Data

On macOS: \Library\Application Support\JMP\16\Samples\Data

In JMP Pro, sample data is installed in the JMPPRO (rather than JMP) directory.

To view examples using sample data, select **Help > Sample Data** and navigate to the Teaching Resources section. To learn more about the teaching resources, visit <https://jmp.com/tools>.

Learn about Statistical and JSL Terms

For help with statistical terms, select Help > Statistics Index. For help with JSL scripting and examples, select **Help > Scripting Index**.

Statistics Index Provides definitions of statistical terms.

Scripting Index Lets you search for information about JSL functions, objects, and display boxes. You can also edit and run sample scripts from the Scripting Index and get help on the commands.

Learn JMP Tips and Tricks

When you first start JMP, you see the Tip of the Day window. This window provides tips for using JMP.

To turn off the Tip of the Day, clear the **Show tips at startup** check box. To view it again, select **Help > Tip of the Day**. Or, you can turn it off using the Preferences window.

JMP Tooltips

JMP provides descriptive tooltips (or *hover labels*) when you hover over items, such as the following:

- Menu or toolbar options
- Labels in graphs
- Text results in the report window (move your cursor in a circle to reveal)
- Files or windows in the Home Window
- Code in the Script Editor

Tip: On Windows, you can hide tooltips in the JMP Preferences. Select **File > Preferences > General** and then deselect **Show menu tips**. This option is not available on macOS.

JMP User Community

The JMP User Community provides a range of options to help you learn more about JMP and connect with other JMP users. The learning library of one-page guides, tutorials, and demos is a good place to start. And you can continue your education by registering for a variety of JMP training courses.

Other resources include a discussion forum, sample data and script file exchange, webcasts, and social networking groups.

To access JMP resources on the website, select **Help > JMP User Community** or visit <https://community.jmp.com>.

Free Online Statistical Thinking Course

Learn practical statistical skills in this free online course on topics such as exploratory data analysis, quality methods, and correlation and regression. The course consists of short videos, demonstrations, exercises, and more. Visit <https://www.jmp.com/statisticalthinking>.

JMP New User Welcome Kit

The JMP New User Welcome Kit is designed to help you quickly get comfortable with the basics of JMP. You'll complete its thirty short demo videos and activities, build your confidence in using the software, and connect with the largest online community of JMP users in the world. Visit <https://www.jmp.com/welcome>.

Statistics Knowledge Portal

The Statistics Knowledge Portal combines concise statistical explanations with illuminating examples and graphics to help visitors establish a firm foundation upon which to build statistical skills. Visit <https://www.jmp.com/skp>.

JMP Training

SAS offers training on a variety of topics led by a seasoned team of JMP experts. Public courses, live web courses, and on-site courses are available. You might also choose the online e-learning subscription to learn at your convenience. Visit <https://www.jmp.com/training>.

JMP Books by Users

Additional books about using JMP that are written by JMP users are available on the JMP website. Visit <https://www.jmp.com/books>.

The JMP Starter Window

The JMP Starter window is a good place to begin if you are not familiar with JMP or data analysis. Options are categorized and described, and you launch them by clicking a button. The JMP Starter window covers many of the options found in the Analyze, Graph, Tables, and File menus. The window also lists JMP Pro features and platforms.

- To open the JMP Starter window, select **View (Window on macOS) > JMP Starter**.
- To display the JMP Starter automatically when you open JMP on Windows, select **File > Preferences > General**, and then select **JMP Starter** from the Initial JMP Window list. On macOS, select **JMP > Preferences > Initial JMP Starter Window**.

JMP Technical Support

JMP technical support is provided by statisticians and engineers educated in SAS and JMP, many of whom have graduate degrees in statistics or other technical disciplines.

Many technical support options are provided at <https://www.jmp.com/support>, including the technical support phone number.

Introduction to Reliability and Survival

Lifetime and Failure Analysis

Reliability and Survival Methods describes a number of methods and tools that are available in JMP to help you evaluate and improve reliability in a product or system and analyze survival data for people and products:

- The Life Distribution platform enables you to analyze the lifespan of a product, component, or system to improve quality and reliability. This analysis helps you determine the best material and manufacturing process for the product, thereby increasing the quality and reliability of the product. See [Chapter 3, “Life Distribution”](#).
- The Fit Life by X platform helps you analyze lifetime events when only one factor is present. You can choose to model the relationship between the event and the factor using various transformations, or create a custom transformation of your data. See [Chapter 4, “Fit Life by X”](#).
- The Cumulative Damage platform enables you to analyze an accelerated life test where the stress levels might have changed over time. See [Chapter 5, “Cumulative Damage”](#).
- The Recurrence Analysis platform analyzes event times where the events can recur several times for each unit. Typically, these events occur when a unit breaks down, is repaired, and then put back into service after the repair. See [Chapter 6, “Recurrence Analysis”](#).
- The Degradation platform analyzes degradation data to predict pseudo failure times. These pseudo failure times can then be analyzed by other reliability platforms to estimate failure distributions. You can include an explanatory factor. You can perform stability analysis to set product expiration dates. You can also fit custom destructive degradation models. See [Chapter 7, “Degradation”](#).
- The Destructive Degradation platform models failure data for product characteristics whose measurement requires that the product be destroyed. This results in a single observation per product unit. You can also include an acceleration factor. A wide range of common degradation models is available. See [Chapter 8, “Destructive Degradation”](#).
- The Reliability Forecast platform helps you predict the number of future failures. The analysis estimates the parameters for a life distribution using production dates, failure dates, and production volume. See [Chapter 9, “Reliability Forecast”](#).
- The Reliability Growth platform models the change in reliability of a single repairable system over time as improvements are incorporated into its design. See [Chapter 10, “Reliability Growth”](#).

-  The Reliability Block Diagram platform displays the reliability relationship between a system's components and, if reliability distributions are given to the components, analytically obtains the reliability behavior. See [Chapter 11, "Reliability Block Diagram"](#).
-  The Repairable Systems Simulation platform enables you to interactively define the relationships between the components of a repairable system. It can also simulate the down time of the system. See [Chapter 12, "Repairable Systems Simulation"](#).
- The Survival platform computes survival estimates for one or more groups. It can be used as a complete analysis or is useful as an exploratory analysis to gain information for more complex model fitting. See [Chapter 13, "Survival Analysis"](#).
- The Fit Parametric Survival platform fits the time to event variable using linear regression models that can involve both location and scale effects. The fit is performed using several distributions. See [Chapter 14, "Fit Parametric Survival"](#).
- The Fit Proportional Hazards platform fits the Cox proportional hazards model, which assumes a multiplying relationship between predictors and the hazard function. See [Chapter 15, "Fit Proportional Hazards"](#).

Chapter 3

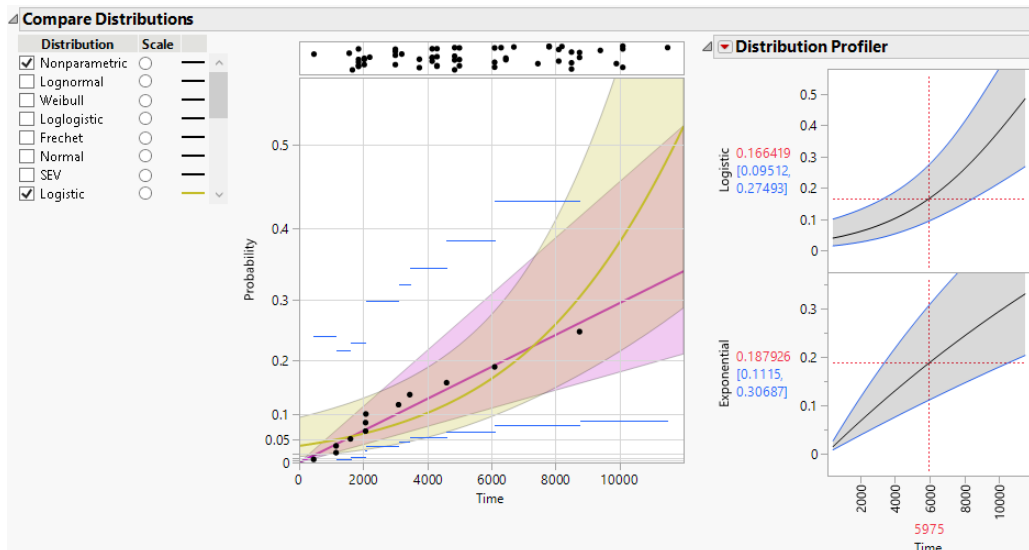
Life Distribution

Fit Distributions to Lifetime Data

The Life Distribution platform models time-to-event data. The platform accommodates both right-censored and interval-censored data. Use Life Distribution to do the following:

- Compare multiple distributional fits to determine which distribution best fits your data.
- Construct Bayesian fits.
- Model zero-failure data.
- Compare groups to analyze group differences.
- Analyze multiple causes of failure.
- Estimate the components of a mixture and estimate the probability that an observation comes from a given component.
- Estimate the components of a competing risk mixture and estimate the impact of a component on an observation.

Figure 3.1 Distributional Fits and Comparisons



Contents

Overview of the Life Distribution Platform	31
Example of the Life Distribution Platform	31
Launch the Life Distribution Platform	34
Life Distribution Report	37
Event Plot	38
Compare Distributions	41
Life Distribution: Statistics	45
Life Distribution Report Options	52
Mixture	54
Fit Competing Risk Mixture	58
Weibayes Report	59
Competing Cause Report	60
Competing Cause Workflow	60
Competing Cause Model	61
Cause Combination	61
Competing Cause: Statistics	62
Individual Causes	63
Competing Cause Report Options	64
Life Distribution - Compare Groups Report	66
Compare Groups: Statistics	66
Individual Group	67
Life Distribution - Compare Groups Report Options	67
Additional Examples of the Life Distribution Platform	69
Omit Competing Causes	69
Change the Scale	71
Examine the Same Distribution across Groups	73
Weibayes Estimates	75
Fit Mixture Example	77
Fit Competing Risk Mixture Example	80
Statistical Details for the Life Distribution Platform	82
Distributions	82
Competing Cause Details	97
Notes Regarding Median Rank Regression	102

Overview of the Life Distribution Platform

Life data analysis, or *life distribution analysis*, is the process of modeling the lifespan of a product, component, or system to predict lifetime or time to failure. For example, you can observe failure rates over time to predict when a computer component might fail. This technique enables you to compare materials and manufacturing processes for the product, enabling you to increase the quality and reliability of the product. Decisions on warranty periods and advertising claims can also be more accurate.

With the Life Distribution platform, you can analyze *censored* data in which some time observations are unknown. And if there are potentially multiple causes of failure, you can analyze the *competing causes* to estimate which cause is more influential.

You can use the **Reliability Test Plan** and **Reliability Demonstration** calculators to choose the appropriate sample sizes for reliability studies. These calculators are found at **DOE > Sample Size and Power**. See the *Design of Experiments Guide*.

Example of the Life Distribution Platform

Suppose you have failure times for 70 engine fans, with some of the failure times censored. You want to fit a distribution to the failure times and then estimate various measurements of reliability.

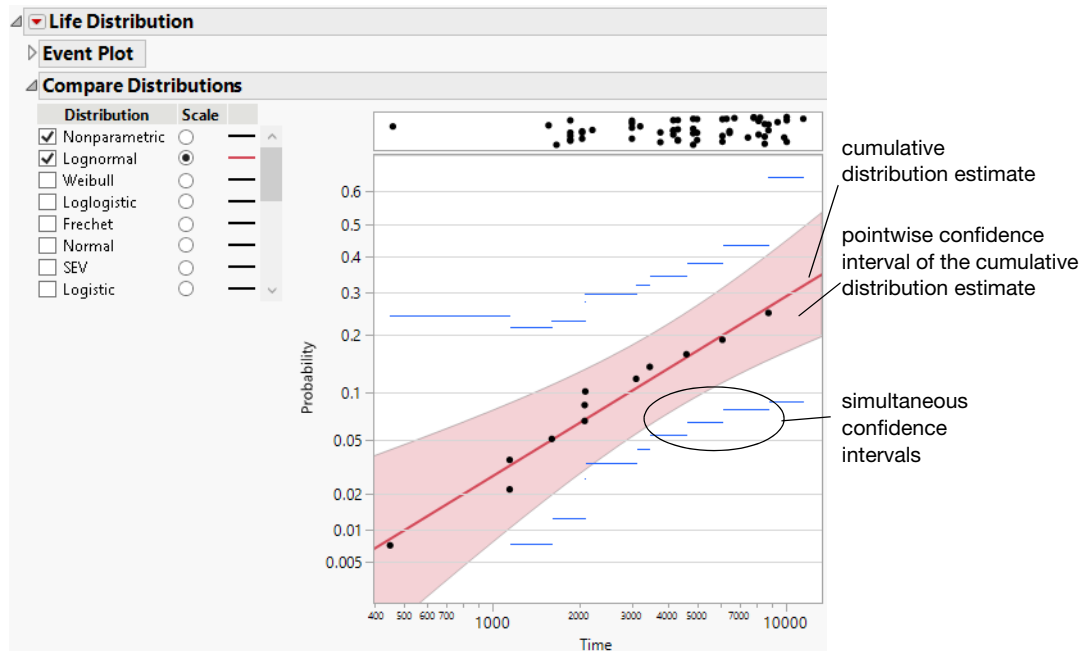
1. Select **Help > Sample Data Library** and open Reliability/Fan.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Time and click **Y, Time to Event**.
4. Select Censor and click **Censor**.
5. Click **OK**.

The Life Distribution report window appears.

6. In the Compare Distribution report, select **Lognormal** distribution and the corresponding **Scale** radio button.

A probability plot appears in the report window.

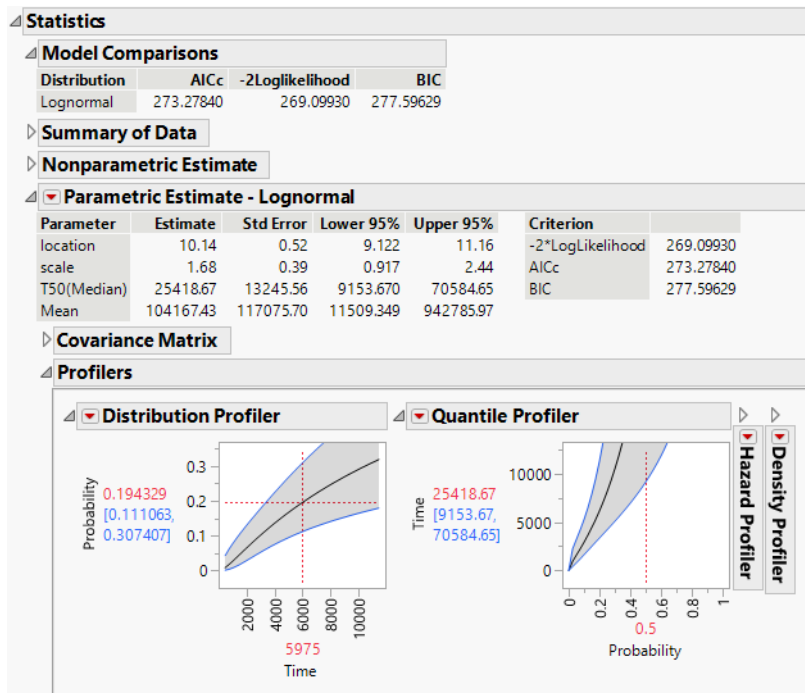
Figure 3.2 Probability Plot



In the probability plot, the data points generally fall along the red line, indicating that the lognormal fit is reasonable.

Below the Compare Distributions report, the Statistics report appears. This report provides a Model Comparison report, nonparametric and parametric estimates, profilers, and more.

Figure 3.3 Statistics Report

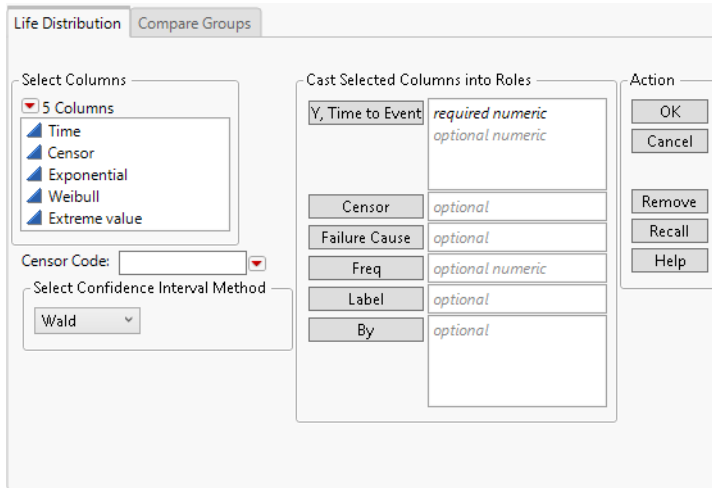


The parameter estimates for the lognormal distribution are provided. The profilers are useful for visualizing the fitted distribution and for estimating probabilities and quantiles. For example, the Quantile Profiler indicates that the estimated median time to failure is 25,418.67 hours.

Launch the Life Distribution Platform

Launch the Life Distribution platform by selecting **Analyze > Reliability and Survival > Life Distribution**. For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

Figure 3.4 The Life Distribution Launch Window



Launch Window Tabs

The launch window includes two tabs:

- The Life Distribution tab models ungrouped data. The following types of reports can result:
 - The Life Distribution report appears when you do not specify a Failure Cause role. From this report, you can compare common distributions and examine statistics. See [“Life Distribution Report”](#) on page 37.
 - The Weibayes report appears when you have zero failures in your data. See [“Weibayes Report”](#) on page 59.
 - The Competing Cause report appears when you specify a Failure Cause role. In addition to the features in the Life Distribution report, you can also compare individual failure causes. See [“Competing Cause Report”](#) on page 60.

Note: You can examine Fixed Parameter and Bayesian models in the Life Distribution and Competing Cause reports.

- The Compare Groups tab enables you to specify a Grouping variable. The Compare Groups report compares different groups using a single specified distribution. For

example, you might compare Weibull fits for components grouped by supplier. In contrast, the Life Distribution tab compares several fitted distributions for a single group. See [“Life Distribution - Compare Groups Report”](#) on page 66.

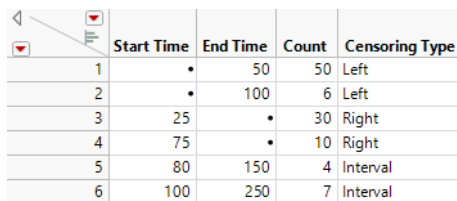
Launch Window Options

The launch window contains the following options:

Y, Time to Event One or more response columns. The number of response columns specified depends on the censoring structure in the data table.

- If one variable is specified, it is interpreted as the time to event (such as the time to failure) or time to right censoring. Use the Censor column to indicate right-censored responses. For more information about right censoring, see [“Single Time to Event Column”](#) on page 39.
- If two variables are specified, they are interpreted as interval-censored observations. The first Y variable gives the lower limit and the second Y variable gives the upper limit for each unit. For an example of using two response columns to represent various types of censoring, see Figure 3.5. For more information about censoring with two response columns, see [“Two Time to Event Columns”](#) on page 40.

Figure 3.5 Censored Data Types for Two Response Variables



	Start Time	End Time	Count	Censoring Type
1	•	50	50	Left
2	•	100	6	Left
3	25	•	30	Right
4	75	•	10	Right
5	80	150	4	Interval
6	100	250	7	Interval

- If three or more variables are specified, the report contains a separate analysis computed using each specified variable as time to event data.

Grouping (Appears only in the Compare Groups tab.) A column containing the groups that you want to compare. For an example, see [“Examine the Same Distribution across Groups”](#) on page 73.

Censor A column that identifies right-censored observations. Select the value that identifies right-censored observations from the Censor Code menu beneath the Select Columns list. The Censor column is used only when one Y is entered.

Failure Cause A column that contains multiple failure causes. If a Failure Cause column is selected, then a section is added to the window. This section contains check boxes that allow the failure mode to use ZI distributions, TH distributions, DS distributions, fixed parameter models, or Bayesian models for the analysis. The following options are also available:

Distribution Specifies the initial distribution to fit for each failure cause. Select one distribution to fit for all causes; select **Individual Best** to let the platform automatically choose the best fit for each cause; or select **Manual Pick** to manually choose the distribution to fit for each failure cause after JMP creates the Life Distribution report. You can also change the distribution fits in the Life Distribution report itself.

Comparison Criterion (Appears only when you choose the **Individual Best** distribution fit.) Specifies the method by which JMP chooses the best distribution: Corrected Akaike Information Criterion (AICc), Bayesian Information Criterion (BIC), or twice the negative log-likelihood (-2Loglikelihood). See *Fitting Linear Models*. You can change the method later in the Model Comparisons report. See [“Model Comparisons”](#) on page 45.

Censor Indicator in Failure Cause Column Identifies the value used in the **Failure Cause** column to indicate that an observation did not fail. To specify such an indicator, select this option and then enter the indicator in the box that appears. You can also select a value from the list to the right of the box.

See Meeker and Escobar (1998, ch. 15) for a discussion of multiple failure causes. [“Omit Competing Causes”](#) on page 69 illustrates how to analyze multiple causes.

Freq A column that contains frequencies or observation counts when the information in a row represents multiple units. If the value in a row is 0 or a positive number, then the value represents the frequencies or counts of observations for that row.

Label A column that contains identifiers other than the row number. These labels appear on the Y axis in the event plot.

By An optional variable whose levels define rows used to create separate models.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

Select Confidence Interval Method Defines the method used for computing confidence intervals for the parameters. The default is Wald, but you can select Likelihood instead. However, all confidence intervals provided in the profilers are based on the Wald method. This is done to reduce computation time. See [“Estimation and Confidence Intervals”](#) on page 83.

Failure Distribution by Cause (Appears only in the Life Distribution tab when a Cause is specified.) Specify which families of distributions should be available to model the life distributions for individual causes. Select an initial distribution, Individual Best, or Manual Pick from the Distribution menu. See [“Failure Cause”](#) on page 36.

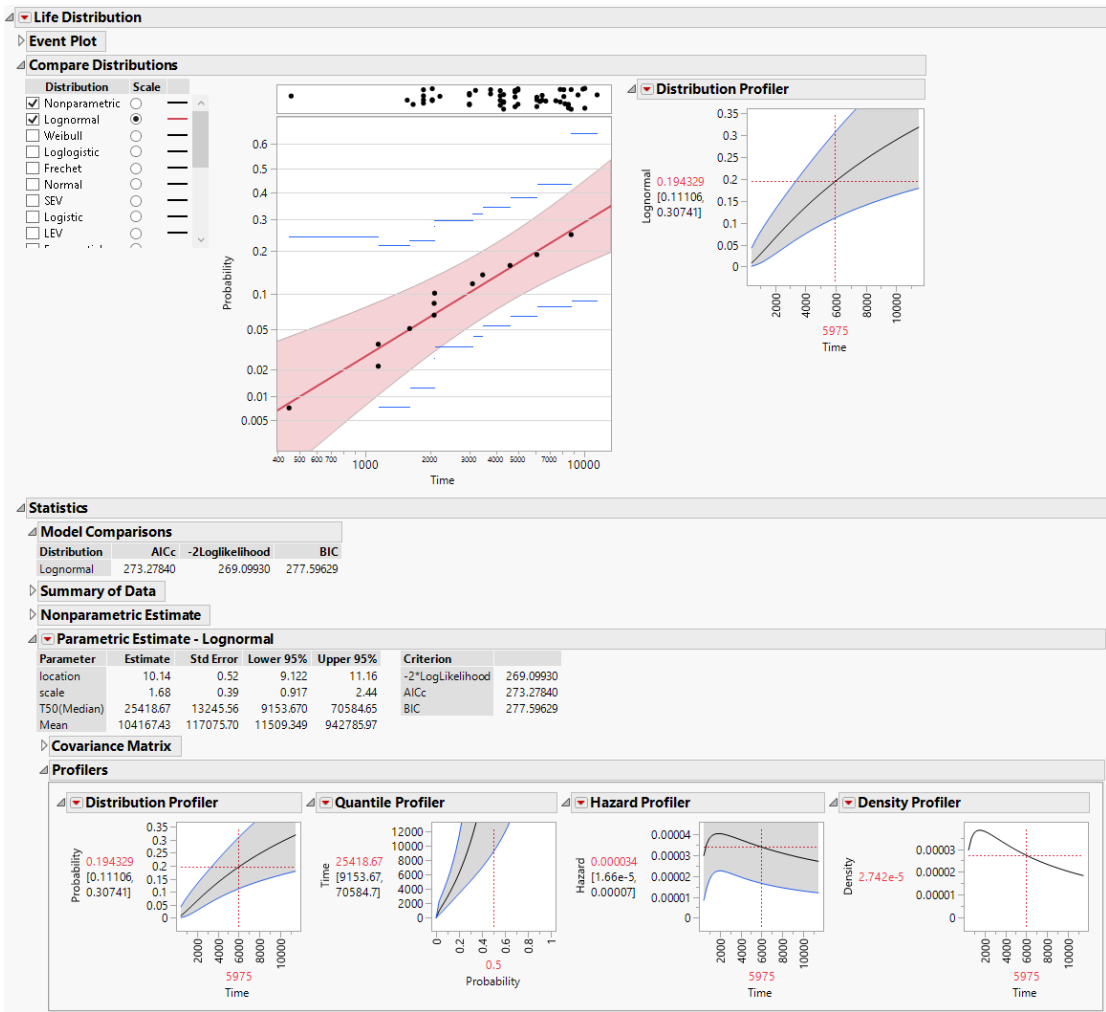
Life Distribution Report

Tip: If you find that the report window is too long, select **Tabbed Report** from the Life Distribution red triangle menu.

If you have not selected a Failure Cause in the launch window, the Life Distribution report appears. Use this report to analyze lifetime data where some time observations might be censored. The Life Distribution report contains the following content and options:

- [“Event Plot”](#) on page 38
- [“Compare Distributions”](#) on page 41
- [“Life Distribution: Statistics”](#) on page 45
- [“Life Distribution Report Options”](#) on page 52

Figure 3.6 Example of the Life Distribution Report for Fan.jmp



Event Plot

Click the Event Plot disclosure icon to see a plot of the failure or censoring times. For each row in the data table, the Event Plot shows a horizontal line indicating whether the units in the row have been censored. When units have been censored, the line indicates the nature of the censoring.

- The time period when the units in the row are known to be functioning is indicated with a solid line.
- The time period when it is *not* known if the units in the row are functioning is indicated with a dashed line.

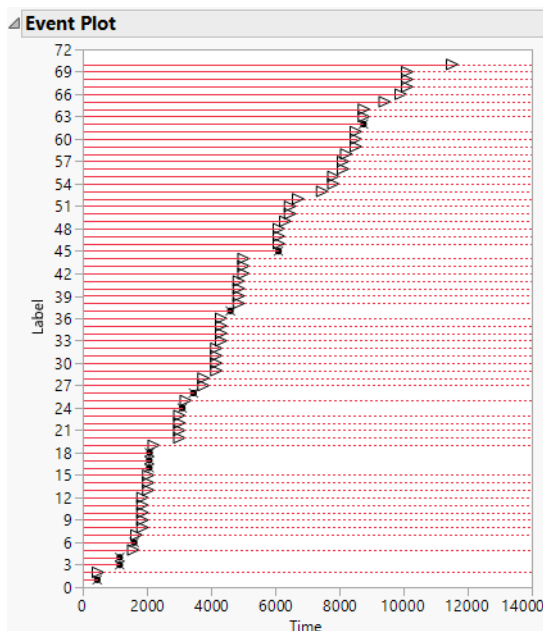
- The line terminates once the units in the row are known to have failed.

Single Time to Event Column

In the Fan.jmp sample data table there is a single Time column indicating failure time. When the failure time is unknown, the value Censored is recorded in the Censor column. All censored units are assumed to be right-censored. Figure 3.7 shows the Event Plot for this data.

Note: To construct the plot in Figure 3.7, select **Help > Sample Data Library** and open Reliability/Fan.jmp. Click the green triangle next to the **Life Distribution - Exponential** script. Click the **Event Plot** disclosure icon.

Figure 3.7 Event Plot for Right-Censored Data



The unit in row 3 failed at Time 1150. Its lifetime is represented by a solid horizontal line that ends at Time 1150. The failure time is marked with an "x".

The unit in row 5 is right censored. It was last known to be functioning at Time 1560. The time period during which the unit is known to be functioning is represented by a solid horizontal line that ends at Time 1560. At Time 1560, a right arrow is plotted. The line continues as a dashed line, indicating that the failure time is unknown, but greater than 1560.

Two Time to Event Columns

In the *Censor Labels.jmp* sample data table, there are two columns, *Start Time* and *End Time*. *Start Time* indicates when units in a row were last known to be functioning. *End Time* indicates when units in that row were first known to have failed.

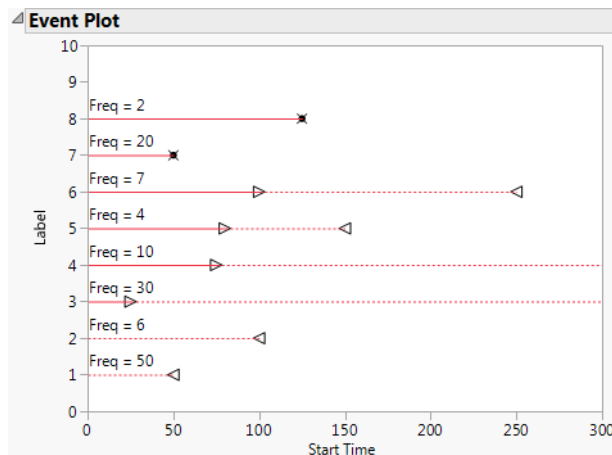
The *Start Time* and *End Time* values indicate the following about the units:

- Units in rows 1 and 2 are left censored. They were known to fail before the time in the *End Time* column, but their exact failure times are unknown.
- Units in rows 3 and 4 are right censored. They were known to be last functioning at the time in the *Start Time* column, but their failure times are unknown.
- Units in rows 5 and 6 are interval censored. They were known to fail within the interval defined by the *Start Time* and *End Time*.
- Units in rows 7 and 8 are not censored. Their failure times are given by the values in the *Start Time* and *End Time* columns, which are identical.

Figure 3.8 shows the Event Plot for this data.

Note: To construct the plot in Figure 3.8, select **Help > Sample Data Library** and open *Censor Labels.jmp*. Click the green triangle next to the **Life Distribution** script.

Figure 3.8 Event Plot for Mixed-Censored Data



The line patterns in the Event Plot represent the various types of censoring:

- The pattern  indicates right censoring. The unit failed after its last inspection.
- The pattern  indicates left censoring. The unit failed after being put on test and prior to the indicated time, but it is not known when it was last functioning.

- The pattern  indicates that the unit failed during the time interval marked by the two arrow heads.
- The pattern  indicates no censoring. The unit failed at the time marked by the x.

Compare Distributions

The Compare Distributions report lets you fit and compare different failure time distributions. Two lists appear:

Distribution Select a distribution for the response. Different distributions appear based on characteristics of the data. For more information about which distributions are available, see [“Available Parametric Distributions”](#) on page 42.

Scale Select a scale for the probability axis. The probability scale corresponds to the distribution listed to the left of the Scale button. Using this scale, the fitted model is represented by a line. Suppose that you fit a given distribution and then scale the axis using that distribution. If the points generally fall along a line, this indicates that the distribution provides a reasonable fit.

The default plot shows the nonparametric estimates (Kaplan-Meier-Turnbull) for the uncensored data values and their confidence intervals. The confidence intervals are indicated by horizontal blue lines.

Tip: To customize the plot of the nonparametric estimates, select **File > Preferences > Platforms > Life Distribution** and select one or more of the following preferences: **Show Shaded Pointwise Intervals**, **Show Shaded Simultaneous Intervals**, or **Show Staircase Style Function**.

By default, there is a panel at the top of the plot that displays times of right-censored observations.

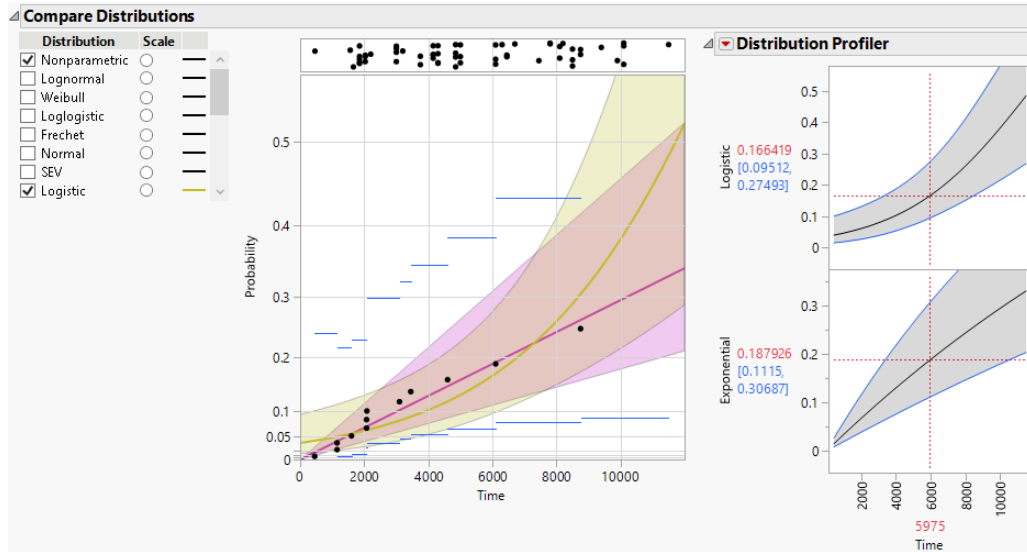
Tip: To hide the panel that contains markers for right-censored observations, select **File > Preferences > Platforms > Life Distribution** and uncheck **Show Markers for Right Censored Observations**.

For each distribution that you select, the Compare Distributions report is updated to show the following:

- the estimated cumulative distribution curve, which appears on the probability plot
- a shaded region that indicates confidence intervals for the cumulative distribution
- a Distribution Profiler that shows the cumulative probability of failure for a given period of time

Figure 3.9 shows an example of the Compare Distributions report. The Logistic (yellow) and Exponential (magenta) distributions are shown. The plot is scaled using the Exponential distribution.

Figure 3.9 Compare Distributions Report and Distribution Profiler



Available Parametric Distributions

This section addresses the distributions available in the Compare Distributions report.

Note: Distributions for the Competing Cause report are covered in [“Available Distributions for Competing Cause Compare Distributions Reports”](#) on page 64.

The available distributions are listed and described in detail in [“Parametric Distributions”](#) on page 84. There are four major groupings of parametric distributions:

- [“Basic Failure-Time Distributions”](#) on page 43
- [“Threshold Distributions”](#) on page 43
- [“Defective Subpopulation Distributions”](#) on page 44
- [“Zero-Inflated Distributions”](#) on page 44

Tip: To restrict which distributions are available by default, select **File > Preferences > Platforms > Life Distribution** and uncheck the distributions that you do not want to appear. The distributions listed include the Threshold, Defective Subpopulation, Zero-Inflated, LogGenGamma, and GenGamma distributions. By default, all of the distributions are checked and available.

The rules that determine which distributions appear in the Compare Distributions panel depend on the particular implementation. As a general guide, distributions are available if they are not disabled in Preferences and if they are appropriate in the given situation.

Basic Failure-Time Distributions

The basic failure-time distributions are available whenever all failure times are positive. They include the following:

- Lognormal
- Weibull
- Loglogistic
- Fréchet
- Normal
- SEV
- Logistic
- LEV
- Exponential
- LogGenGamma
- GenGamma

Note: When there are negative or zero failure times, only the Normal, SEV, Logistic, LEV, and LogGenGamma are available.

Threshold Distributions

The threshold (TH) distributions are always available. Threshold distributions are log-location-scale distributions with threshold parameters. The threshold parameter shifts the distribution away from 0. These distributions assume that all units survive until the threshold value. Threshold distributions are useful for fitting moderate to heavily shifted distributions. The threshold distributions are the following:

- TH Lognormal
- TH Weibull
- TH Loglogistic
- TH Fréchet

Defective Subpopulation Distributions

The defective subpopulation (DS) distributions are available when all failure times are positive. These distributions are useful when only a fraction of the population has a particular defect leading to failure. Use the DS distribution options to model failures that occur on only a subpopulation. The DS distributions are the following:

- DS Lognormal
- DS Weibull
- DS Loglogistic
- DS Fréchet

Zero-Inflated Distributions

When the time-to-event data contain zero as the minimum value in the Life Distribution platform, the following zero-inflated distributions are available:

- Zero-Inflated Lognormal (ZI Lognormal)
- Zero-Inflated Weibull (ZI Weibull)
- Zero-Inflated Loglogistic (ZI Loglogistic)
- Zero-Inflated Fréchet (ZI Fréchet)

Zero-inflated distributions are used when some proportion of units fails at time zero. When the data contain more zeros than expected by a standard model, the number of zeros is inflated.

Zero-Failure Data

In the case of zero-failure data, none of the above distributions are available by default. To obtain Bayesian fits for those distributions where the Bayesian Estimate option is available, select **File > Preferences > Platforms > Life Distribution** and uncheck **Weibayes Only for Zero Failure Data**. See [“Weibayes Only for Zero Failure Data”](#) on page 51.

Parametric Distributions That Allow Bayesian Estimation

Bayesian estimation is available for the following parametric distributions:

- Lognormal
- Weibull
- Loglogistic
- Fréchet
- Normal
- SEV

- Logistic
- LEV

A list of distributions that are available as priors for hyperparameters of these distributions is given in [“Prior Distributions for Bayesian Estimation”](#) on page 96.

Life Distribution: Statistics

The Statistics report includes the following sub-reports:

- [“Model Comparisons”](#) on page 45
- [“Summary of Data”](#) on page 45
- [“Nonparametric Estimate”](#) on page 45
- [“Parametric Estimate - <Distribution Name>”](#) on page 46 (one report appears for each distribution that you select in the Compare Distributions report)

Model Comparisons

The Model Comparisons report provides the AICc, -2Loglikelihood, and BIC statistics for each fitted distribution. Smaller values of each of these statistics indicate a better fit. For more information about these statistics, see *Fitting Linear Models*.

Initially, the rows are sorted by AICc. To change the statistic used to sort the report, click the Life Distribution red triangle and select **Comparison Criterion**. See [“Life Distribution Report Options”](#) on page 52 for more information about this option.

Summary of Data

The Summary of Data report shows the total number of units observed, the number of uncensored units, and the numbers of right-censored, left-censored, and interval-censored units.

Nonparametric Estimate

The Nonparametric Estimate report shows nonparametric estimates for each observation. For right-censored data specified as a single Time to Event column, the report gives the following:

Midpoint Estimate Midpoint-adjusted Kaplan-Meier estimates.

Lower 95%, Upper 95% Pointwise 95% confidence intervals. You can change the confidence level by selecting Change Confidence Level from the report options.

Simultaneous Lower 95% (Nair), Simultaneous Upper 95% (Nair) Simultaneous 95% confidence intervals. You can change the confidence level by selecting Change Confidence Level from the report options. See Nair (1984) and Meeker and Escobar (1998).

Kaplan-Meier Estimate Standard Kaplan-Meier estimates.

If failure times are represented by two Time to Event columns, the report gives Turnbull estimates (in a column called Estimate), pointwise confidence intervals, and simultaneous confidence intervals (Nair).

See “Nonparametric Fit” on page 83 for more information about nonparametric estimates.

Parametric Estimate - <Distribution Name>

A report called Parametric Estimate - <Distribution Name> appears for each distribution that is fit. The report gives the distribution's parameter estimates, their standard errors, and confidence intervals. The criteria that appear in the Model Comparisons report are shown under Criterion.

Note: Whenever an estimate of the mean is provided, its confidence interval is computed as a Wald interval even if you select Likelihood as the Confidence Interval Method in the launch window. In this case, the notation Mean (Wald CI) appears in the Parameter column to indicate that the confidence interval for the mean is a Wald interval.

For more information about how the distributions are parametrized, see “Parametric Distributions” on page 84.

The Parametric Estimate report contains the following reports:

- “Covariance Matrix” on page 46
- “Profilers” on page 46
- Additional reports can be added by selecting report options from the Parametric Estimate red triangle menu. These include the Fix Parameter, Bayesian Estimates, Custom Estimation (Estimate Probability, Estimate Quantile), and Mean Remaining Life reports. See “Parametric Estimate Options” on page 47.

Covariance Matrix

For each distribution, the Covariance Matrix report shows the covariance matrix for the estimates.

Profilers

Four types of profilers appear for each distribution:

- The Distribution Profiler shows cumulative failure probability as a function of time.

- The Quantile Profiler shows failure time as a function of cumulative probability.
- The Hazard Profiler shows the hazard rate as a function of time.
- The Density Profiler shows the density function for the distribution.

The profilers contain the following red triangle menu options:

Confidence Intervals The Distribution, Quantile, and Hazard profilers show Wald-based confidence curves for the plotted functions. This option shows or hides the confidence curves.

Reset Factor Grid Displays a window for each factor enabling you to enter a specific value for the factor's current setting, to lock that setting, and to control aspects of the grid. See *Profilers*.

Factor Settings Provides a menu that consists of several options. See *Profilers*.

Note: The confidence intervals provided in the profilers are based on the Wald method even if the Likelihood Confidence Interval Method is selected in the launch window. This is done to reduce computation time.

Parametric Estimate Options

The Parametric Estimate red triangle menu has the following options:

Save Probability Estimates Saves the estimated failure probabilities and confidence intervals to the data table.

Save Quantile Estimates Saves the estimated quantiles and confidence intervals to the data table.

Save Hazard Estimates Saves the estimated hazard values and confidence intervals to the data table.

Show Likelihood Contour Shows or hides a contour plot of the log-likelihood function. If you have selected the Weibull distribution, a second contour plot appears for the alpha-beta parameterization. This option is available only for distributions with two parameters.

Show Likelihood Profiler Shows or hides a profiler of the log-likelihood function. This option is not available for the threshold (TH) distributions.

Fix Parameter Opens a report where you can specify the value of parameters. Enter your fixed parameter values, select the appropriate check box, and then click **Update**. JMP re-estimates the other parameters, covariances, and profilers based on the new parameters, and shows them in the Fix Parameter report. A distribution profiler of the unconstrained model is shown below the distribution profiler for the fixed parameter

model. For an example in a competing cause situation, see [“Specify a Fixed Parameter Model as a Distribution for a Cause”](#) on page 97.

For the Weibull distribution, the **Fix Parameter** option lets you select the Weibayes method. For an example, see [“Weibayes Estimates”](#) on page 75. The Weibayes option is not available for interval-censored data.

Bayesian Estimates Performs Bayesian estimation of parameters for certain distributions based on three methods of specifying prior distributions (Location and Scale Priors, Quantile and Parameter Priors, and Failure Probability Priors). See [“Bayesian Estimation - <Distribution Name>”](#) on page 48. This option is available only for the following distributions: Lognormal, Weibull, Loglogistic, Fréchet, Normal, SEV, Logistic, LEV.

Custom Estimation Provides calculators that enable you to predict failure probabilities, survival probabilities, and quantiles for specific time and failure probability values. Each calculated quantity includes confidence intervals, which can be two-sided or one-sided (in either direction). Two reports appear: Estimate Probability and Estimate Quantile. See [“Custom Estimation”](#) on page 51.

Mean Remaining Life Provides a calculator that enables you to estimate the mean remaining life of a unit. In the Mean Remaining Life Calculator, enter a Time and press **Enter** to see the estimate. Click the plus sign to enter additional times. This calculator is available for the following distributions: Lognormal, Weibull, Loglogistic, Fréchet, Normal, SEV, Logistic, LEV, and Exponential.

Bayesian Estimation - <Distribution Name>

For certain distributions, you can fit Bayesian models. This is done using rejection sampling or a Markov Chain Monte Carlo (MCMC) algorithm. More specifically, the platform attempts a basic rejection sampler. If the rejection sampler produces valid results, these results are reported. If the rejection sampler cannot produce valid results, the platform uses a random walk Metropolis-Hastings algorithm and adds a note to the top of the Bayesian Estimation report. See Robert and Casella (2004).

From the Parametric Estimate - <Distribution Name> report outline, select **Bayesian Estimates**. This opens an outline called Bayesian Estimation - <Distribution Name>. The initial report is a control panel where you can specify the parameters for the priors and control aspects of the simulation.

The following steps describe the workflow:

- Select a prior specification method from the Bayesian Estimation red triangle menu and set values for the parameters of the priors. See [“Bayesian Estimation Red Triangle Options”](#) on page 49.
- Specify the simulation options. See [“Bayesian Estimates - Result <N>”](#) on page 50.
- Select Fit Model to fit a model. See [“Bayesian Estimates - Result <N>”](#) on page 50.

Bayesian Estimation Red Triangle Options

You can choose from the following prior specification methods in the Bayesian Estimation red triangle menu:

Location and Scale Priors Enables you to specify hyperparameters for prior distributions on generic parameters (location and scale parameters). Select the Prior Distribution red triangle menu to select a distribution for each parameter. You can enter new values for the hyperparameters of the priors. The initial values that are provided are estimates consistent with the MLEs. See [“Prior Distributions for Bayesian Estimation”](#) on page 96.

Quantile and Parameter Priors Enables you to specify prior information about a quantile and the scale parameter (or Weibull β if the parametric fit is Weibull). The quantile is defined by the value next to Probability. The default Probability value is 0.10, but you can specify a value that corresponds to the quantile of interest. Specify information about the prior information in terms of Lower and Upper 99% limits on the range of each prior distribution. See Meeker and Escobar (1998). The initial values that are provided are estimates consistent with the MLEs. See [“Prior Distributions for Bayesian Estimation”](#) on page 96.

Failure Probability Priors Enables you to specify prior information about failure probabilities at two distinct time points. You can specify the two time points. The prior distribution for each time point is Beta. You can specify the prior distributions using either of two synchronized approaches:

1. Specify failure probability by estimates and error percentages. The prior information for each Beta prior distribution can be specified using a probability estimate and an estimate error. See Kaminskiy and Krivtsov (2005).
2. Specify failure probability estimate ranges. You can specify the 99% range for the two Beta distributions in the following ways:
 - For each failure time, enter an initial value for the Lower and Upper 99% Limits.
 - Click the vertical line segments in the graph and drag them to your two time points. Adjust the vertical spread of each marker to specify the 99% limits.

Simulation Options

For any of the prior specification methods that you select in the Bayesian Estimation red triangle menu, the following options appear at the bottom of the panel:

Number of Monte Carlo Iterations Controls the sample size that will be drawn from the posterior distribution after a burn-in procedure.

Random Seed Sets the initial state of the simulation. By default, it is the clock time. The number should be a positive integer greater than 1. If you specify 1, the current clock time is used.

Show Prior Scatter Plot Select this option to draw random samples from the prior distributions and to plot results on a scatter plot. After you select Fit Model, the scatter plot appears in an outline entitled Prior Scatter Plot in the Bayesian Estimates - Results <N> report.

Overlay Likelihood Contour Overlays likelihood-based contours on scatter plots in the Bayesian Estimates Results report.

Fit Model Estimates the posterior lifetime distribution based on prior distributions that JMP fits using the values that you specified. Adds a report entitled Bayesian Estimates - Results <N>, where N is an integer that consecutively numbers the Bayesian Results reports.

Bayesian Estimates - Result <N>

Once you have specified priors using one of the red triangle menu options, select Fit Model. A Bayesian Estimates - Result <N> report is provided for each selection of priors. This report contains these headings:

Priors Documents the specifications that you entered in the Bayesian Estimation report to fit the Bayesian model. The Priors report also specifies the random seed.

Posterior Estimates Shows five marginal statistics that describe the posterior distribution of the generic parameters (location and scale parameters). The marginal statistics are the median, 0.025 quantile (Lower Bound), 0.975 quantile (Upper Bound), mean, and standard deviation computed from the Monte Carlo samples. When the posterior estimates are generated using the Quantile and Parameter Priors specification, this table also includes the posterior estimate of the quantile and Slope β (for the Weibull distribution).

To compute statistics for other derived variables based on the posterior estimates of the generic parameters, click the Export Monte Carlo Samples link.

Prior Scatter Plot Appears when you select Show Prior Scatter Plot before clicking Fit Model. Shows prior scatter plots of parameters or equivalent quantities associated with the prior specification method for the distribution.

Posterior Scatter Plot Shows posterior scatter plots of parameters or equivalent quantities associated with the prior specification method for the distribution.

Profilers Shows two profilers based on samples from the posterior distribution.

The values shown in the Distribution Profiler, at a given time t , are calculated as follows:

- For each set of sampled parameter values from the posterior distribution, the value of the cumulative distribution function at time t is calculated.
- The predicted value is the median of these calculated values.

- The upper and lower confidence limits are the 0.025 and 0.975 quantiles of these calculated values.

The plot and confidence limits shown in the Quantile Profiler are obtained in a similar fashion. For a given Probability value p , the quantiles corresponding to p are calculated from the distributions associated with the posterior parameter values.

Weibayes Only for Zero Failure Data

In a zero-failure situation, no units fail. All observations are right censored. If you have zero-failure data, it is possible to conduct either Bayesian estimation or Weibayes inference. See [“Weibayes Report”](#) on page 59.

Note: By default, zero-failure data is analyzed using the Weibayes method. If you want to conduct a broader Bayesian analysis on zero-failure data, select **File > Preferences > Platforms > Life Distribution** and uncheck **Weibayes Only for Zero Failure Data**.

Custom Estimation

The Custom Estimation option produces two reports: Estimate Probability and Estimate Quantile. The Estimate Probability report contains a calculator that enables you to predict failure and survival probabilities for specific time values. The Estimate Quantile report contains a calculator that enables you to predict quantiles for specific failure probability values. Both Wald-based and likelihood-based confidence intervals appear for each estimated quantity. The confidence level for these intervals is determined by the Change Confidence Level option in the Life Distribution red triangle menu.

Estimate Probability

In the Estimate Probability calculator, enter a value for Time. Press **Enter** to see the estimates of failure probabilities, survival probabilities, and corresponding confidence intervals. To calculate multiple probability estimates, click the plus sign, enter another Time value in the box, and press **Enter**. Click the minus sign to remove the last entry.

The Estimate Probability calculator contains an option, Side, that enables you to change the form of the intervals. Select one of the following suboptions:

Two Sided Provides two-sided confidence intervals for failure probability and survival probability.

Upper Failure Probability Provides one-sided confidence intervals that contain upper limits for the failure probability and lower limits for the survival probability.

Lower Failure Probability Provides one-sided confidence intervals that contain lower limits for the failure probability and upper limits for the survival probability.

Estimate Quantile

In the Estimate Quantile report, enter a value for Failure Probability. Press **Enter** to see the quantile estimates and corresponding confidence intervals. To calculate multiple quantile estimates, click the plus sign, enter another Failure Probability value in the box, and press **Enter**. Click the minus sign to remove the last entry.

The Estimate Quantile calculator contains an option, Side, that enables you to change the form of the intervals. Select one of the following suboptions:

Two Sided Provides two-sided confidence intervals for quantiles.

Lower Provides one-sided confidence intervals that contain lower limits for quantiles.

Upper Provides one-sided confidence intervals that contain upper limits for quantiles.

Life Distribution Report Options

The Life Distribution red triangle menu contains the following options:

Fit All Distributions Fits all distributions other than the threshold (TH) distributions. The distributions are compared in the Model Comparisons report. See [“Compare Distributions”](#) on page 41.

Tip: Select the **Comparison Criterion** option to change the criterion for finding the best distribution.

Fit All Nonnegative Fits all nonnegative distributions (Exponential, Lognormal, Loglogistic, Fréchet, Weibull, and Generalized Gamma). The distributions are compared in the Model Comparisons report. See [“Compare Distributions”](#) on page 41. Note the following:

- The option does not fit DS or TH distributions.
- If the data have negative values, then the option produces no results.
- If the data have zeros, the option fits the four zero-inflated (ZI) distributions: ZI Lognormal, ZI Weibull, ZI Loglogistic, and ZI Fréchet. For more information about zero-inflated distributions, see [“Zero-Inflated Distributions”](#) on page 95.

Fit All DS Distributions Fits all defective subpopulation (DS) distributions: DS Lognormal, DS Weibull, DS Loglogistic, and DS Fréchet. For more information about defective subpopulation distributions, see [“Distributions for Defective Subpopulations”](#) on page 94.

Fit Mixture Fits a distribution that is a mixture of the distributions other than the threshold (TH) distributions. See [“Mixture”](#) on page 54.

Fit Competing Risk Mixture Fits a competing risk mixture distribution to the data. See [“Fit Competing Risk Mixture”](#) on page 58.

Show Points Shows or hides data points in the probability plot. The Life Distribution platform uses the midpoint estimates of the step function to construct probability plots. When you deselect **Show Points**, the midpoint estimates are replaced by Kaplan-Meier estimates.

Show Event Plot Frequency Label (Appears only if you have specified a Freq variable.) Shows or hides the Frequency label in the Event Plot.

Show Survival Curve Switches between the failure probability and the survival curve on the Compare Distributions probability plot and the Distribution Profiler plots.

Show Quantile Functions Shows or hides a Quantile Profiler that overlays the plots for the selected distributions. The Quantile plot also shows points plotted at time values. The plot appears beneath the Compare Distributions report. If you select distributions in any of the Compare Distributions, Quantile Profiler, and Hazard Profiler plots, they appear in the other two plots.

Show Hazard Functions Shows or hides a Hazard Profiler that overlays the plots for the selected distributions. The plot appears above the Statistics report. If you select distributions in any of the Compare Distributions, Quantile Profiler, and Hazard Profiler plots, they appear in the other two plots.

Show Statistics Shows or hides the Statistics report. See [“Life Distribution: Statistics”](#) on page 45.

Tabbed Report Shows graphs and data on individual tabs rather than in the default outline style.

Show Confidence Area Shows or hides the shaded confidence regions in the plots.

Interval Type Determines the type of confidence interval shown for the Nonparametric fit in the Compare Distributions plot. Select either pointwise or simultaneous confidence intervals.

Change Confidence Level Enables you to change the confidence level for the entire platform. All plots and reports update accordingly.

Comparison Criterion Enables you to select the criterion used to rank models in the Model Comparison report. For all three criteria, smaller values indicate better fit. Burnham and Anderson (2002) and Akaike (1974) discuss using AICc and BIC for model selection. See *Fitting Linear Models*.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

If you have specified a By variable, separate Life Distribution reports appear for each level of the By variable.

Save By Group Results For each By group, saves the estimates that appear in all Parametric Estimate reports for that group as a separate row in a new table.

Do Same Analyses for All Groups Applies all of the selected options for the current group to all other By groups.

Mixture

The Fit Mixture option adds the Mixture outline to the report where you can fit a mixture distribution to the data. For an example, see [“Fit Mixture Example”](#) on page 77.

The mixture distribution's probability function $F(x)$ is defined as follows:

$$F(x) = \sum_{i=1}^k w_i F_i(x)$$

where $F_i(x)$ is one of the supported distributions, k is the number of components in the mixture, and the w_i are positive weights that sum to 1. The Fit Mixture option attempts to identify clusters of observations that are drawn from each of the component distributions, $F_i(x)$. It estimates the parameters of the mixture and the probability that an observation is drawn from any given component.

Model Fit and Mixture Starting Value Methods

The fitting methodology is based on assumptions about the underlying clusters, called the Starting Value Method. Suppose that you designate k distributions. There are three Starting Value Methods:

- Single Cluster assumes that all observations are affected by all of the ingredient distributions to some extent. None of the densities stand out as affecting only a portion of the observations.
- Separable Clusters assumes that the ingredient distributions affect some observations more profoundly than others. For separable clusters, each of the k densities has an identifiable mode and defines a cluster.
- Overlapping Clusters assumes a situation that is intermediate between Single Cluster and Separable Clusters. Some densities stand out, but others jointly affect a portion of the observations. In this case, there are m clusters in the data, where m is less than k , the total number of densities.

The fitting process consists of these steps:

1. Clusters of observations are defined.
2. Assignment of clusters to densities is based on the Starting Value Method:
 - For Separable Clusters, the highest likelihood assignment of clusters to the specified ingredient densities is determined by examining the possible permutations.
 - For Overlapping Clusters, the highest likelihood assignment of clusters to the specified ingredient densities is determined by examining the possible permutations of clusters and combinations of observations.

Note: Suppose that you fit a model using a given Starting Value Method and then select another Starting Value Method. If a better fit based on the likelihood value cannot be achieved, no new model is added.

Mixture Control Panel

The control panel consists of these items:

Ingredient Lists distributions that you can use as components of the fitted mixture distribution.

Quantity Select the number of components in the mixture distribution that have the given distribution. The sum of the Quantity values is k , the number of densities in the mixture.

Starting Value Methods Select a method that reflects your assumptions about the mixture. See [“Model Fit and Mixture Starting Value Methods”](#) on page 55.

Overlay Shows the nonparametric estimates (Kaplan-Meier-Turnbull) for the uncensored data values. When you fit a mixture, the plot is updated to show the model and 95% level confidence bands. The confidence level for these bands is determined by the Change Confidence Level option in the Life Distribution red triangle menu. A Legend appears to the right of the plot.

Go Click **Go** to fit the desired mixture. The Model List is updated with the model that you fit, and a report with the name of the mixture model is added.

Fit Mixture Reports

Model List

The Model List report lists the mixture distributions that you fit. The report provides the number of parameters, the number of actual observations, and the AICc, -2Loglikelihood, and BIC statistics for each mixture distribution. For more information about these statistics, see *Fitting Linear Models*.

Note the following:

- Smaller values of each of these statistics indicate a better fit.
- The rows are sorted by AICc.
- The **Comparison Criterion** red triangle option does not affect the order of models in the Model List.
- The AICc, -2Loglikelihood, and BIC statistics also appear in the Model Comparisons table. This enables you to compare mixture distribution to other distributions for your data. See [“Model Comparisons”](#) on page 45.

Mixture Reports

The Model List report is followed by reports for each of the mixture distributions that you have fit. The title of each report describes the corresponding mixture using the specified ingredients and their quantities. The report lists the parameters, their estimates, standard errors, and 95% Wald confidence intervals. These intervals are not affected by the selection of Likelihood as the Confidence Interval Method in the launch window.

Parameter estimates are given for each distribution in the mixture. The Parameter column also includes parameters called Portion $\langle i \rangle$, where $i = 1, 2, \dots, k-1$. These are estimates of the weights w_i for the mixture. Since the weights sum to 1, the k^{th} weight can be computed from the first $k - 1$ weights.

Density Overlay Plot

The Density Overlay plot shows estimates of the density functions for each of the components in the mixture. A legend to the right of the plot enables you to select which density functions appear.

Mixture Report Options

The red triangle menu contains the following options:

Remove Removes the model report and the entry for the model in the Model List.

Show Profilers Shows four types of profilers for the combined mixture distribution F . See [“Mixture Profiler Options”](#) on page 57 for a description of their red triangle options.

- The Distribution Profiler shows cumulative failure probability as a function of time.
- The Quantile Profiler shows failure time as a function of cumulative probability.
- The Hazard Profiler shows the hazard rate as a function of time.
- The Density Profiler shows the density function for the distribution.

Save Predictions For each mixture density, saves a column to the data table containing the probability that an observation belongs to that density. For the formulas used in the calculation, see [“Fit Mixture Save Predictions Formulas”](#) on page 101.

Mixture Profiler Options

The profilers for each mixture report contain the following red triangle options:

Confidence Intervals The Distribution, Quantile, and Hazard profilers show 95% Wald-based confidence curves for the plotted functions. This option shows or hides the confidence curves. The confidence level for these curves is determined by the Change Confidence Level option in the Life Distribution red triangle menu.

Note: To reduce computation time, the confidence intervals provided in the profilers are based on the Wald method, even if the Likelihood Confidence Interval Method is selected in the launch window.

Reset Factor Grid Displays a window for each factor enabling you to enter a specific value for the factor's current setting, to lock that setting, and to control aspects of the grid. See *Profilers*.

Factor Settings Provides a menu that consists of options relating to profiler settings, scripts, and linking profilers. See *Profilers*.

Fit Competing Risk Mixture

The Fit Competing Risk Mixture option enables you to fit competing risk mixture distributions. It estimates the probability that a given observation fails due to the cause represented by each of the component mixture distributions. For an example, see [“Fit Competing Risk Mixture Example”](#) on page 80.

The competing risk mixture probability distribution function $F(x)$ is defined as follows:

$$F(x) = 1 - \prod_{i=1}^k (1 - F_i(x))$$

where $F_i(x)$ represents the cumulative failure distributions for the i^{th} risk and k is the number of components (or risks) in the mixture. The Fit Competing Risk Mixture option attempts to identify clusters of observations that are drawn from each of the component distributions, $F_i(x)$. It estimates the parameters of the mixture and the probability that an observation is drawn from any given component.

The Competing Risk Mixture report is structured in a fashion that mirrors the Mixture report. See [“Mixture”](#) on page 54. However, the Fit Competing Risk Mixture reports do not show a Density Overlay plot. They show a Distribution Overlay plot instead.

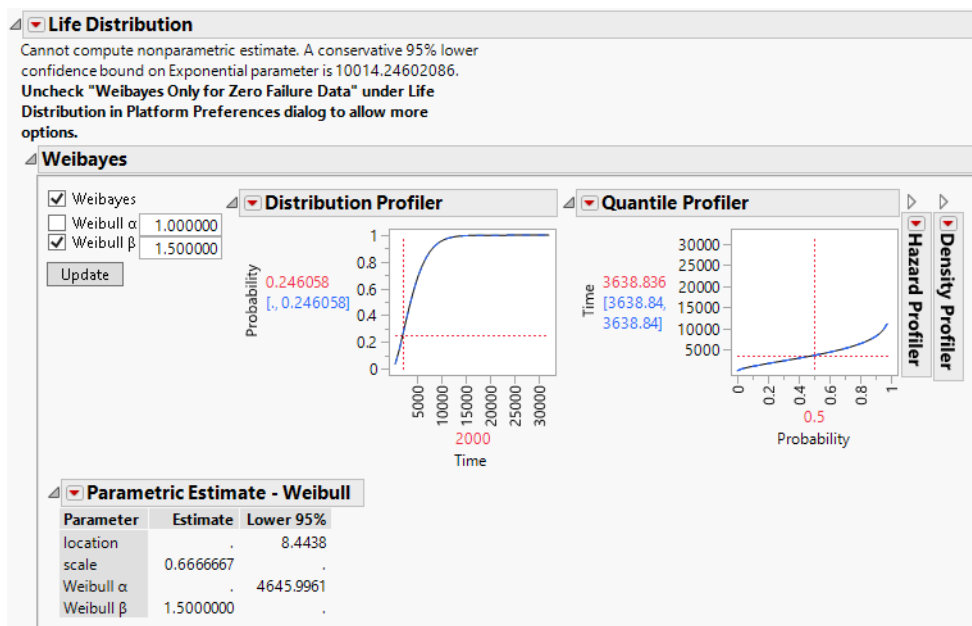
Distribution Overlay Plot

The Distribution Overlay plot shows the cumulative distribution functions for each of the mixture components and for the combined mixture (Aggregated). A legend to the right of the plot enables you to select which cumulative distribution functions appear.

Weibayes Report

If you have data with zero failures (right-centered) and you have not turned off the preference **Weibayes Only for Zero Failure Data**, a special Weibayes report appears. Figure 3.10 shows the Weibayes report for the Weibayes No Failures.jmp sample data table, found in the Reliability folder.

Figure 3.10 Weibayes Report



The analysis begins by assuming an exponential model. The note at the top of the report provides a lower confidence bound for the parameter of an exponential distribution. This lower confidence bound is computed using the method described in Meeker and Escobar (1998, sec. 7.7).

The Weibayes section of the report conducts the Weibayes analysis. For a description of the procedure, see Nelson (1985).

To obtain Weibayes estimates, make sure that Weibayes and Weibull β options are selected. Change the Weibull β value and click Update. The estimates and profilers are updated. The values shown in the profilers use the conservative confidence bound. For an example, see “Weibayes Example for Data with No Failures” on page 75.

Note: If you deselect the Weibayes option, JMP shows the fixed parameter MLE.

When your data table contains at least one failure, a full Life Distribution report appears, but the maximum likelihood inference might not produce useful results. In this instance, a Weibayes analysis might be preferable. For an example, see [“Weibayes Example for Data with One Failure”](#) on page 76.

Note: To conduct Bayesian inference using the usual Life Distribution options when you have zero-failure data, select **File > Preferences > Platforms > Life Distribution** and deselect **Weibayes Only for Zero Failure Data** before launching the analysis.

Competing Cause Report

Tip: If you find that the report window is too long, select **Tabbed Report** from the Competing Cause red triangle menu.

The Competing Cause report appears if you have assigned a Failure Cause column in the launch window. Use this report to analyze the competing causes to determine which causes are influential. For an example of this report, see [“Omit Competing Causes”](#) on page 69. For technical details, see [“Competing Cause Details”](#) on page 97.

The Competing Cause report contains the following content and options:

- [“Cause Combination”](#) on page 61
- [“Competing Cause: Statistics”](#) on page 62
- [“Individual Causes”](#) on page 63
- [“Competing Cause Report Options”](#) on page 64

Competing Cause Workflow

Follow these steps to facilitate your use of the Competing Cause report:

1. For convenience, click the Competing Cause red triangle and select the Tabbed Report option.
2. Select the Individual Causes tab. For each failure cause, use the options in its individual Life Distribution outline to select a distributional fit. See [“Individual Causes”](#) on page 63.
3. Select the Cause Combination tab. For each failure cause, specify the desired distribution in each **Distribution** list. See [“Cause Combination”](#) on page 61.
4. Click **Update Model**.
5. Select the Statistics tab to explore and save results for the model. See [“Competing Cause: Statistics”](#) on page 62.

Note: Customizations made to the competing cause model report in the Statistics outline might be lost if you change the model and click **Update Model** again.

Competing Cause Model

In a competing cause situation, the aggregated failure function can be written as follows:

$$F(x) = 1 - \prod_{i=1}^k [1 - F_i(x)]$$

where $F_i(x)$ is the cumulative failure distribution for the i^{th} cause and k is the total number of causes. The function $F_i(x)$ is cause-specific. It reflects the probability of failure due to cause i alone and does not account for other causes of failure.

An alternative formulation is defined as follows:

$$F(x) = \sum_{i=1}^k \tilde{F}_i(x)$$

where each $\tilde{F}_i$ is a monotone increasing function with values in the interval $[0, 1]$. The function $\tilde{F}_i$ is called a subdistribution. This form of the aggregated failure distribution is used to predict proportions of failures that are associated with individual causes, accounting for failure due to other causes.

Cause Combination

The Cause Combination report lets you fit and compare different failure time distributions ($F_i(x)$) for the various causes. Different distributions are available based on selections that you have made in the launch window. Negative and zero-failure times are not allowed.

The default plot is based on a Linear scale and shows the following:

- Nonparametric estimates (Kaplan-Meier-Turnbull) for the uncensored data values and their confidence intervals. The confidence intervals are represented by horizontal blue lines.
- Fits of cumulative failure distributions ($F_i(x)$) to each of the causes. The initial distribution is the one that you selected in the launch window. If you selected Individual Best, the best distribution is computed for each group and these fits are shown. (This selection can be time-intensive.) If you selected Manual Pick, the initial Distribution is Nonparametric for all groups and nonparametric fits are shown. A legend appears to the right of the plot.

- The Aggregated cumulative failure distribution, $F(x)$, represented by a black line. This function is computed based on the selected cause distribution. If a nonparametric distribution is specified for a cause, the aggregated cumulative failure distribution extends only as far as the final time observation for that cause.

As you interact with the report, statistics for the aggregated model are re-evaluated.

The Cause Combination report contains these elements:

Scale Specifies the probability scale for the plot's vertical axis. "[Change the Scale](#)" on page 71 illustrates how changing the scale affects the distribution fit.

Omit Check a box to remove the fit for the corresponding cause. Use this when a particular cause has been fixed. The Aggregated model updates to reflect the removal of the failures due to that cause. "[Omit Competing Causes](#)" on page 69 illustrates the effect of omitting causes.

Cause Lists the causes in the Cause column.

Distribution Lists the available distributions for each cause. To change the distribution for a specific failure cause, select the distribution from the **Distribution** list. Click **Update Model** to show the new distribution fit on the plot. The Cause Summary report is also updated.

Count Gives the number of observations with the given failure cause.

Update Model Shows the selected distributions in the plot; updates the Cause Summary report with the selected models; adds the selected model as the most recent one in the Individual Causes report.

Competing Cause: Statistics

The Statistics report for Competing Cause contains the following reports:

- "[Cause Summary Report](#)" on page 62
- "[Profilers](#)" on page 63

Cause Summary Report

The Cause Summary report shows information for the fit defined by the current selections under Distribution in the Cause Summary report. The report shows the number of failures for each cause and the parameter estimates for the distribution fit to each failure cause. When you change distribution fits in the Cause Combination report and click Update Model, the Cause Summary report is updated.

The following information is provided:

- The Cause column shows either labels of causes or the censor code.

- The Counts column lists the number of failures for each cause.
- A numerical entry in the Counts column indicates that the cause has enough failure events to consider. A cause with fewer than two events is considered right censored. The column also identifies missing causes.
- The Distribution column specifies the selected distribution for each cause.
- Depending on the selected distributions, various columns display the parameters of the distributions:
 - The location column specifies location parameters for various distributions.
 - The scale column specifies scale parameters for various distributions.
 - The Weibull α and Weibull β columns show Weibull estimates of alpha and beta.
 - Other columns show parameters of other selected distributions.
 - A Convergence column appears if there are convergence issues.

The Cause Summary red triangle menu options enable you to save the probability, quantile, hazard, and density estimates for the aggregated failure distribution to the data table.

Profilers

The Distribution, Quantile, Hazard, and Density profilers help you visualize the aggregated failure distribution. The Distribution, Quantile, and Hazard profilers show 95% level confidence bands. See [“Profilers”](#) on page 46.

Note: If a nonparametric distribution is specified for a cause, the Hazard and Density profilers are not provided. Also, confidence limits are not provided in the Distribution and Quantile profilers.

Individual Causes

The Individual Causes report contains Life Distribution - Failure Cause: <Distribution Name> reports for each of the individual causes. Each Life Distribution - Failure Cause: <Distribution Name> report shows plots and distributional fit statistics for the individual failure cause indicated in the report title.

Whenever you click Update Model in the Cause Combination report, any new cause distribution that you select is added to the Life Distribution - Failure Cause: <Distribution Name> report for that cause. In the Life Distribution - Failure Cause <Distribution Name> report, the following occur:

- The distribution is selected in the Compare Distributions report.
- The distribution is added to the Model Comparisons report.
- A Parametric report is added for that distribution.

Each Life Distribution - Failure Cause <Distribution Name> report is a Life Distribution: Statistics report as described in [“Life Distribution: Statistics”](#) on page 45. However, all confidence intervals are Wald intervals. These intervals are not affected by the selection of Likelihood as the Confidence Interval Method in the launch window.

Available Distributions for Competing Cause Compare Distributions Reports

If you specify a Failure Cause in the launch window, you can specify which groupings of distributions and models you want to allow in the resulting Compare Distributions reports for Individual Causes. You can select ZI (Zero-Inflated), TH (Threshold), and DS (Defective Subpopulation) distributions. You can also select fixed parameter and Bayesian models.

Note: If you have disallowed any distributions in Preferences, these do not appear. Also, rules that govern which distributions appear for Life Distribution apply. See [“Available Parametric Distributions”](#) on page 42.

Competing Cause Report Options

The Competing Cause red triangle menu contains the following options:

Tabbed Report Shows graphs and data on individual tabs rather than in the default outline style.

Tabbed Report for Individual Causes Shows the Life Distribution - Failure Cause: <Distribution Name> reports in tabs, rather than as a stack of Life Distribution reports.

Show Points Shows or hides data points in the Cause Combination plot. The Life Distribution platform uses the midpoint estimates of the step function to construct probability plots. When you deselect **Show Points**, the midpoint estimates are replaced by Kaplan-Meier estimates.

Show Subdistributions Shows the profiler for each individual cause subdistribution $\tilde{F}_i$. See [“Individual Subdistribution Profiler for Cause”](#) on page 65.

Show Remaining Life Distribution Shows the remaining life distribution through the Distribution Profiler. This is conditional upon the unit surviving through a given time.

Mean Remaining Life Estimates the mean remaining life of a unit, given a survival time. In the Mean Remaining Life Calculator, enter a time value and press **Enter** to see the estimate and confidence limits. Click the plus sign to enter additional times. Click the minus sign to remove the most recent entry.

To obtain a confidence interval, select the **Configuration** option from the Mean Remaining Life Calculator. Check **Use bootstrap to construct confidence intervals**. Enter appropriate

values, keeping in mind that computation can be time-intensive. See [“Mean Remaining Life Calculator”](#) on page 101.

Export Bootstrap Results Appears only when a Bayesian model is selected and when **Update Model** has been applied.

Bootstrap Sample Size When Bayesian Estimates or Weibayes results are used for any cause, the confidence limits for aggregated functions that appear in the Distribution Profiler must be simulated using parametric bootstrap. Use this option to specify the number of samples to be used in the bootstrap. See [“Specify a Bayesian Model for a Cause”](#) on page 100.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Individual Subdistribution Profiler for Cause

For a given cause, the Individual Subdistribution Profiler for Cause shows the estimated probability of failure, $\tilde{F}_i$, from that cause at time t . The estimated probability takes into account failures from competing causes. See [“Competing Cause Model”](#) on page 61.

To show a profiler of the subdistribution for each cause, click the Competing Cause red triangle and select **Show Subdistributions**. The Individual Subdistribution Profiler for Cause report appears beneath the other profilers. It consists of a profiler and a calculator.

Note: When you select **Show Subdistributions**, the Cause Combination plot is updated to show the subdistribution functions for all causes.

Select a Cause from the list to the right of the profiler to see its profiler. Options that apply to the profiler are provided in the Individual Subdistribution Profiler for Cause red triangle menu. See [“Profilers”](#) on page 46.

Use the Calculator panel to find values of the subdistribution functions for all causes at one or more Time to Event values. Enter a value for the Time to Event variable. When you press Enter (or click outside the text box), values for each of the causes are updated. To add a Time to Event value, click the plus sign. To remove the most recent value, click the minus sign.

Life Distribution - Compare Groups Report

Tip: If you find that the report window is getting too long, select **Tabbed Report** from the red triangle menu next to Life Distribution - Compare Groups.

If you selected the Compare Groups tab in the launch window, the Life Distribution - Compare Groups report appears. This report compares different groups using a single specified distribution. For example, you might compare Weibull fits for components grouped by supplier. In contrast, the Life Distribution tab compares several fitted distributions for a single group.

You can compare the CDF, Quantile, Hazard, and Density functions. You can also consolidate the probability and quantile predictions of all groups.

For an example of this report, see [“Examine the Same Distribution across Groups”](#) on page 73.

The Compare Groups report can contain the following content and options:

- [“Compare Distributions”](#) on page 41 (no Distribution Profiler)
- [“Compare Groups: Statistics”](#) on page 66
- [“Individual Group”](#) on page 67
- [“Life Distribution - Compare Groups Report Options”](#) on page 67

Compare Groups: Statistics

The Statistics report for Compare Groups contains the following reports:

- [“Summary”](#) on page 67
- [“Group Homogeneity Tests”](#) on page 67
- [“Model Comparisons”](#) on page 45
- [“Parameter Estimates”](#) on page 67

Summary

The Summary report contains a row for each group and for the combined data. Each row shows the number of units that failed and the numbers that were left, interval, or right censored. Each row also gives the corresponding mean and standard error. For more information about how the mean and standard error are computed, see [“Statistical Reports for Survival Analysis”](#) on page 368 in the “Survival Analysis” chapter.

Group Homogeneity Tests

The Group Homogeneity Tests report contains results of three tests for equality of the hazard functions across groups. The tests differ in how the observations are weighted over time.

Log-Rank The Log-Rank test uses equal weights for all time points.

Peto-Peto The Peto-Peto test uses the survival proportion at each time point for weights.

Wilcoxon The generalized Wilcoxon test uses the number at risk at each time point for weights. This test is referred to as the Gehan test in Klein and Moeschberger (1997).

For each of the above tests, the report includes the chi-square approximation, the associated degrees of freedom, and the p -value for each test. A small p -value suggests that the groups differ. For more information about these tests and comparisons of failure curves, see Klein and Moeschberger (1997, ch. 7).

Parameter Estimates

The Parameter Estimates outline contains reports entitled Parametric Estimate - <Distribution Name> for each distribution that is fit. For each group, the Parametric Estimate - <Distribution Name> report gives the distribution's parameter estimates and their 95% level confidence intervals. The confidence level for these intervals is determined by the Change Confidence Level option in the Life Distribution - Compare Groups menu.

Individual Group

The tabs within the Individual Group report contain Life Distribution reports for each individual group. For more information about these reports, see [“Life Distribution Report”](#) on page 37 and [“Life Distribution Report Options”](#) on page 52.

Life Distribution - Compare Groups Report Options

Many of the options in the Compare Groups red triangle menu can also be found in the Life Distribution red triangle menu. See [“Life Distribution Report Options”](#) on page 52.

However, the following options are specific to Compare Groups:

Show Quantile Functions Shows or hides the Compare Quantile report. Select a distribution. For each group, a curve is plotted showing the estimated quantiles for the time variable. Confidence bands are displayed. A legend is shown to the right of the plot. Only one distribution can be specified at a time.

Show Hazard Functions Shows or hides the Compare Hazard report. Select a distribution. For each group, a curve is plotted showing the hazard function. Confidence bands are displayed. A legend is shown to the right of the plot. Only one distribution can be specified at a time.

Show Density Functions Shows or hides the Compare Density report. Select a distribution. For each group, the density function and confidence bands are displayed. A legend is shown to the right of the plot. Only one distribution can be specified at a time.

Estimate Probability Adds an Estimate Probability report corresponding to the most recently selected distribution under Compare Distribution. Enter a value for the Time to Event variable in the text box and press **Enter**. To add a Time to Event value, click the plus sign. To remove the most recent value, click the minus sign. See [“Estimate Probability Report”](#) on page 69.

Estimate Quantile (Appears only if Compare Quantile is selected.) Adds an Estimate Quantile report corresponding to the most recently selected distribution under Compare Quantile. Enter a value for the probability of interest in the text box and press **Enter**. To add a Prob value, click the plus sign. To remove the most recent value, click the minus sign. For each group and probability value, the Time to Event quantile and 95% Wald and Likelihood confidence intervals are shown.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Estimate Probability Report

For each group and Time to Event value, this report provides the following:

Midpoint Estimate Midpoint-adjusted Kaplan-Meier estimate of failure by the specified time.

Lower 95%, Upper 95% Pointwise 95% confidence intervals for the probability of failure by the specified time.

Simultaneous Lower 95% (Nair), Simultaneous Upper 95% (Nair) Simultaneous 95% confidence intervals for the probability of failure by the specified time. See Nair (1984) and Meeker and Escobar (1998).

Survival Probability Midpoint-adjusted estimate of survival beyond the specified time.

Survival Probability Lower 95%, Upper 95% Pointwise 95% confidence intervals for the probability of survival beyond the specified time.

Survival Probability Simultaneous Lower 95% (Nair), Simultaneous Upper 95% (Nair)

Simultaneous 95% confidence intervals for the probability of survival beyond the specified time. See Nair (1984) and Meeker and Escobar (1998).

Additional Examples of the Life Distribution Platform

- [“Omit Competing Causes”](#)
- [“Change the Scale”](#)
- [“Examine the Same Distribution across Groups”](#)
- [“Weibayes Estimates”](#)
- [“Fit Mixture Example”](#)
- [“Fit Competing Risk Mixture Example”](#)

Omit Competing Causes

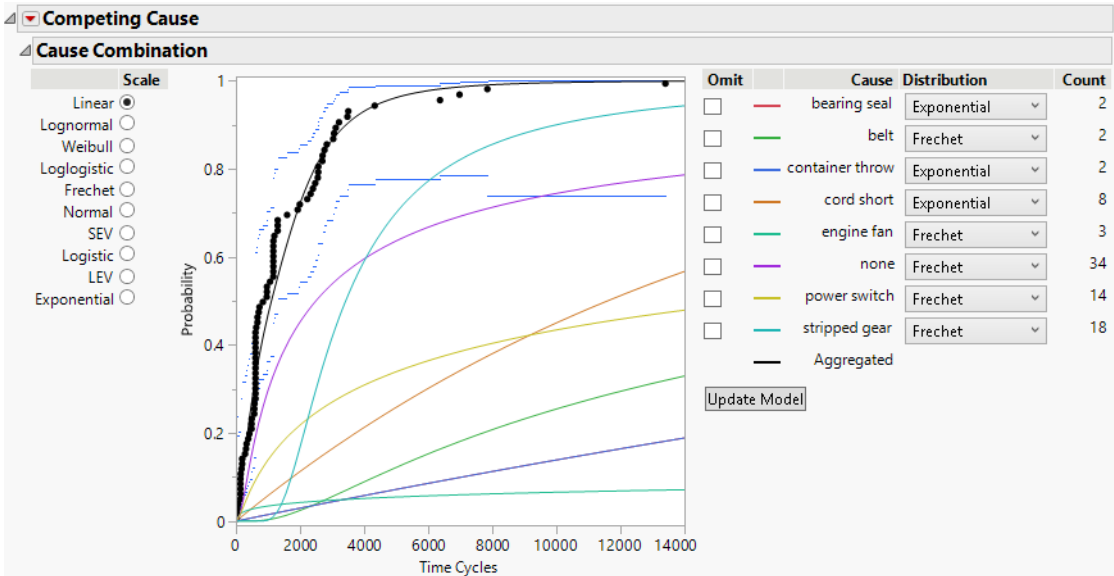
This example illustrates how to decide on the best fit for competing causes.

1. Select **Help > Sample Data Library** and open Reliability/Blenders.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Time Cycles and click **Y, Time to Event**.
4. Select Causes and click **Failure Cause**.
5. Select Censor and click **Censor**.

- 6. Select **Individual Best** as the Distribution.
- 7. Make sure that **AICc** is the Comparison Criterion.
- 8. Click **OK**.

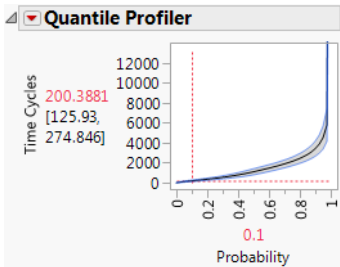
On the Competing Cause report, JMP shows the best distribution fit for each failure cause.

Figure 3.11 Initial Competing Cause Report



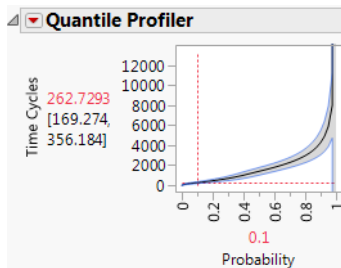
- 9. In the Quantile Profiler, type 0.1 for the probability.
- The estimated time by which 10% of the failures occur is 200.

Figure 3.12 Estimated Failure Time for 10% of the Units



- 10. Select **Omit** for bearing seal, belt, container throw, cord short, and engine fan (the causes with the fewest failures).
- The estimated time by which 10% of the failures occur is now 263.

Figure 3.13 Updated Failure Time



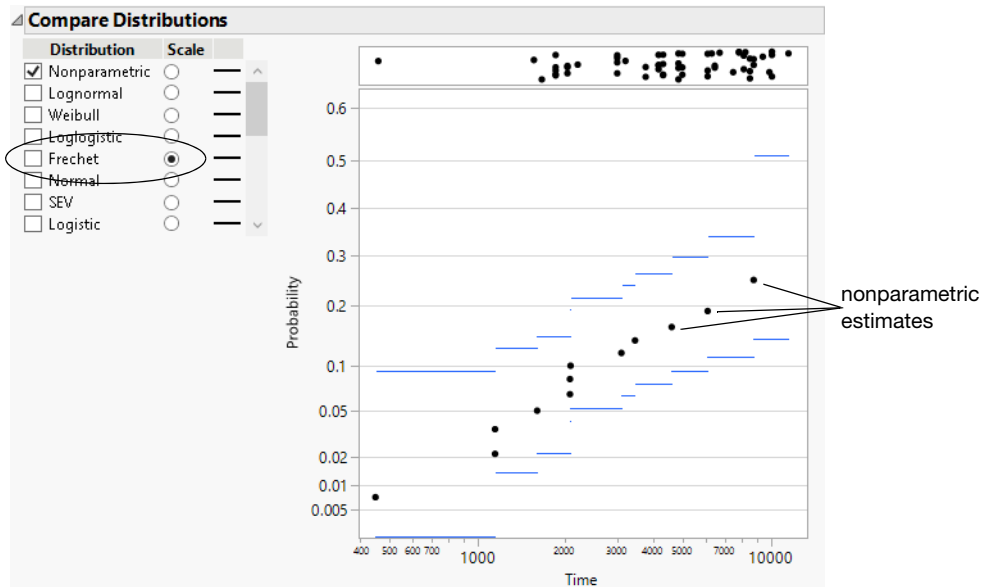
When power switch and stripped gear are the only causes of failure, the estimated time by which 10% of the failures occur increases by approximately 31%.

Change the Scale

In the initial Compare Distributions report, the probability and time axes are linear. But suppose that you want to see distribution estimates on a Fréchet scale.¹

1. Follow step 1 through step 5 in “[Example of the Life Distribution Platform](#)” on page 31.
2. In the Compare Distributions report, select **Fréchet** in the Scale column.
3. Click the Life Distribution red triangle and select **Interval Type > Pointwise**.

1. Using different scales is sometimes referred to as drawing the distribution on different types of *probability paper*.

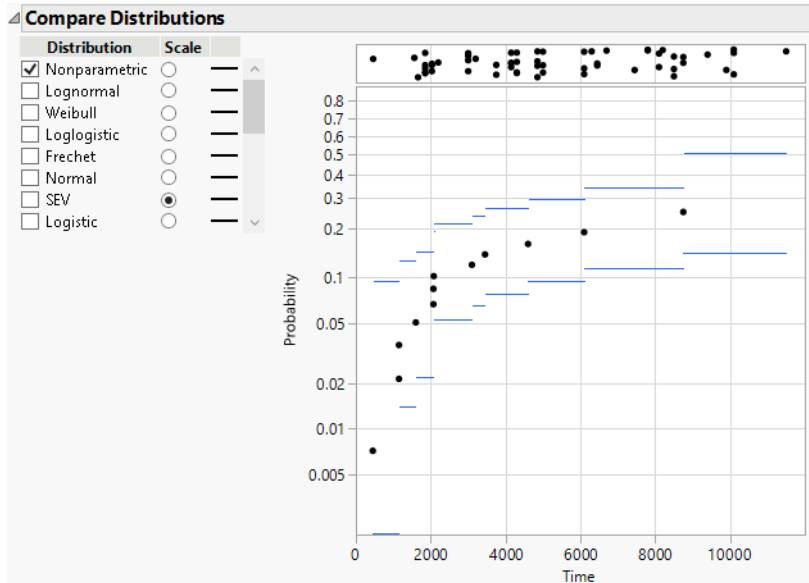
Figure 3.14 Nonparametric Estimates with a Fréchet Probability Scale

Using a Fréchet scale, the nonparametric estimates approximate a straight line, meaning that a Fréchet fit might be reasonable.

4. Select **SEV** in the Scale column.

The nonparametric estimates no longer approximate a straight line. You now know that the SEV distribution is not appropriate.

Figure 3.15 Nonparametric Estimates with a SEV Probability Scale

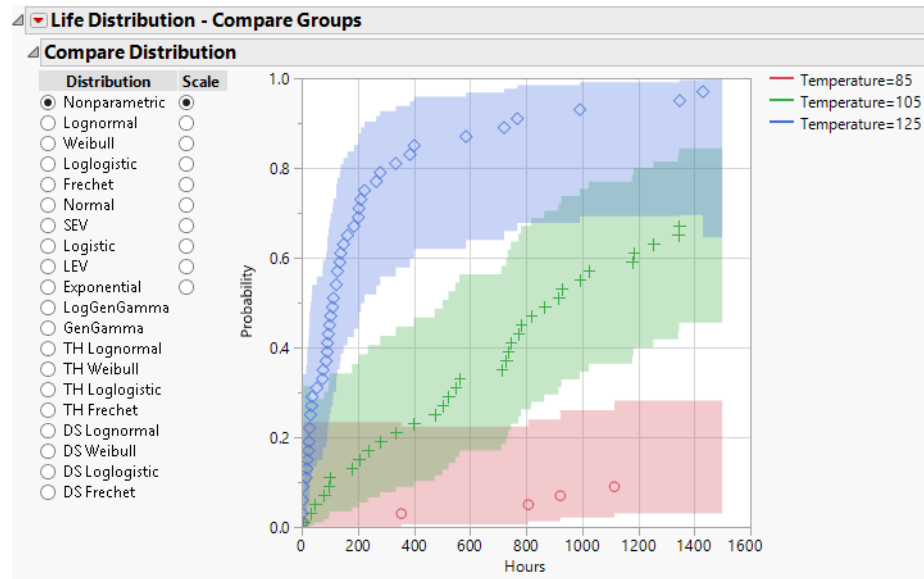


Examine the Same Distribution across Groups

Suppose you want to compare the same distribution across different groups. You want to examine estimates of failure probabilities for a single type of capacitor operating at three different temperatures.

1. Select **Help > Sample Data Library** and open Reliability/Capacitor ALT.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Click the Compare Groups tab.
4. Select Hours and click **Y, Time to Event**.
5. Select Temperature and click **Grouping**.
6. Select Censor and click **Censor**.
7. Select Freq and click **Freq**.
8. Click **OK**.

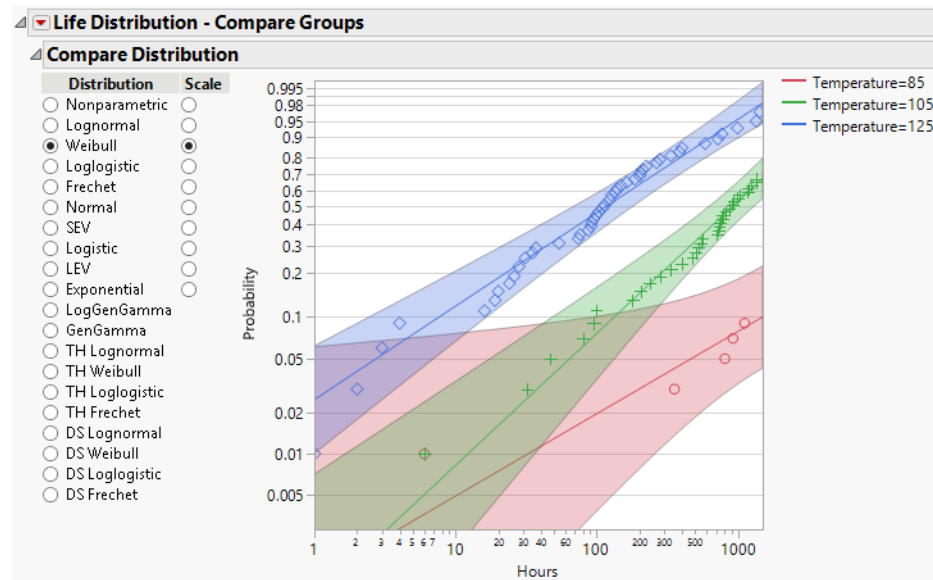
Figure 3.16 Compare Distribution for Groups



The default graph shows the nonparametric estimates. At a higher temperature, the capacitor has a higher probability of failure. You want to try fitting a parametric distribution.

9. Select **Weibull** for Distribution and Scale.

Figure 3.17 Compare Weibull Distribution for Groups



When plotted against a Weibull probability scale, the points come close to following three lines. This indicates that a Weibull distribution provides a reasonable fit for each of the Temperature groups.

Weibayes Estimates

There are two possible ways to obtain a Weibayes analysis:

- You have no failures (all observations are right-censored) and the preference **Weibayes Only for Zero Failure Data** is checked. Then the Weibayes report appears. See [“Weibayes Example for Data with No Failures”](#) on page 75.
- You have few failures. A full Life Distribution report is presented. Fit a Weibull distribution. In the Parametric Estimate - Weibull report, select the Fix Parameter option. Then select the Weibayes option in the Fixed Parameter report. See [“Weibayes Example for Data with One Failure”](#) on page 76.

Weibayes Example for Data with No Failures

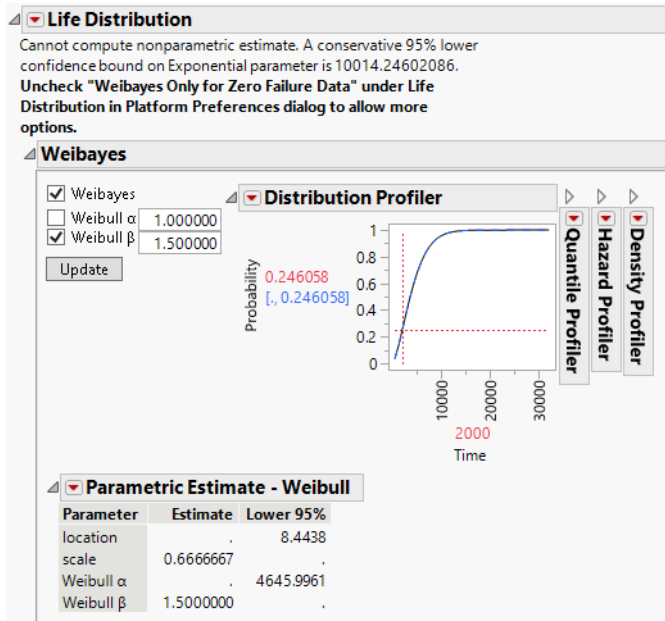
You have data for a product that is mostly reliable. Thirty were tested for 1,000 hours with no failures occurring. You want to predict the failure probability at 2,000 hours.

1. Select **Help > Sample Data Library** and open Reliability/Weibayes No Failures.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Time and click **Y, Time to Event**.
4. Select Censor and click **Censor**.
5. Select Freq and click **Freq**.
6. Select **Likelihood** as the Confidence Interval Method.
7. Click **OK**.

A special Life Distribution report appears. **Weibayes** and **Weibull beta** should be selected.

8. Type 1.5 as the known Weibull β value.
The value 1.5 is considered appropriate for this example.
9. Click **Update**.
10. In the Distribution Profiler, type 2000 for Time.

Figure 3.18 Life Distribution Report for Zero Failures



From the Distribution Profiler, you can see that at 2,000 hours, the conservative probability is 24.6058%. That means that the one-tailed conservative 95% confidence limit for the failure probability is 24.6058%.

Weibayes Example for Data with One Failure

Suppose you have the same data, but this time, one failure occurred at 800 hours. Again, you want to predict the failure probability at 2,000 hours.

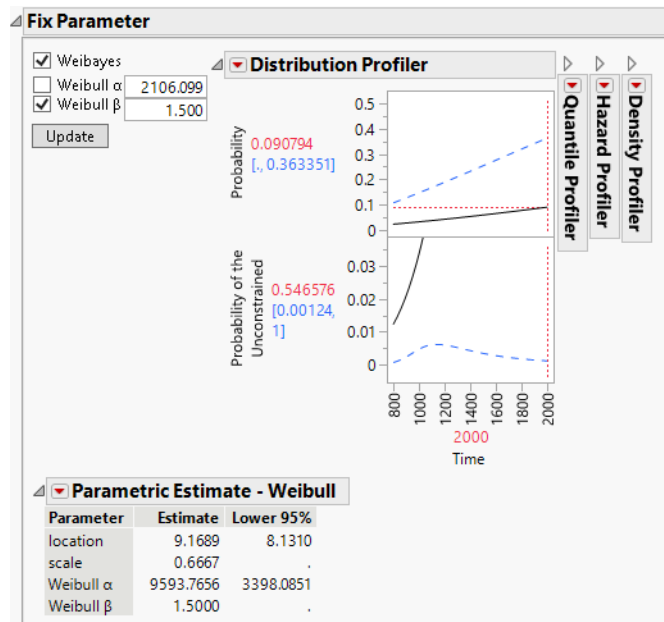
1. Select **Help > Sample Data Library** and open Reliability/Weibayes One Failure.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Time and click **Y, Time to Event**.
4. Select Censor and click **Censor**.
5. Select Freq and click **Freq**.
6. Select **Likelihood** as the Confidence Interval Method.
7. Click **OK**.

The Life Distribution report appears.

8. Select the **Weibull** distribution in the Compare Distributions plot.
9. Click the red triangle next to Parametric Estimate - Weibull and select **Fix Parameter**.
10. Select **Weibayes** and **Weibull beta** in the Fix Parameter report.

11. Type 1.5 as the known Weibull β value.
12. Click **Update**.
13. In the Distribution Profiler, type 2000 for Time.
14. Hover over the top of the Y axis. The cursor becomes a hand. Drag the axis downward until it reaches 0.5 as the top number.

Figure 3.19 Life Distribution Report for One Failure



In the Distribution Profiler, the solid line shows the MLE. The dashed line shows the Weibayes conservative limit. You can see that at 2,000 hours, the conservative probability is 36.3351%. That means that the one-tailed conservative 95% confidence limit for the failure probability is 36.3351%.

Fit Mixture Example

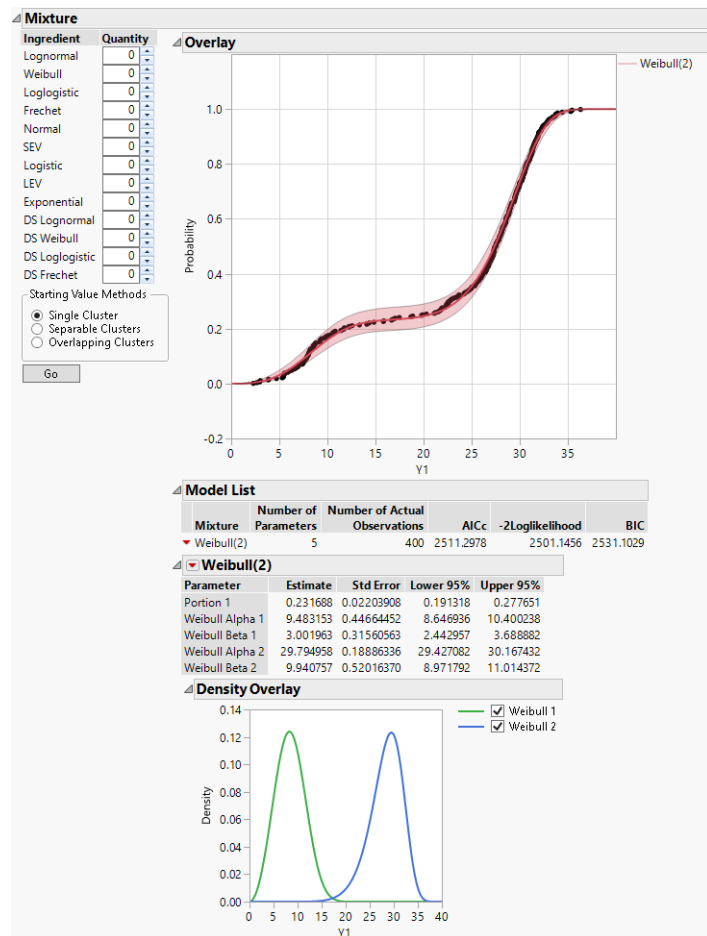
In this example, you fit two mixture distributions and then identify observations belonging to one of the clusters for the second mixture.

Fit Two Mixture Distributions

1. Select **Help > Sample Data Library** and open Reliability/Mixture Demo.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Y1 and click **Y, Time to Event**.

4. Click **OK**.
5. Click the Life Distribution red triangle and select **Fit Mixture**.
6. Type 2 in the **Quantity** box next to **Weibull**.
7. Select **Separable Clusters** in the Starting Value Methods panel.
8. Click **Go**.

Figure 3.20 Fit Mixture for Weibull (2)

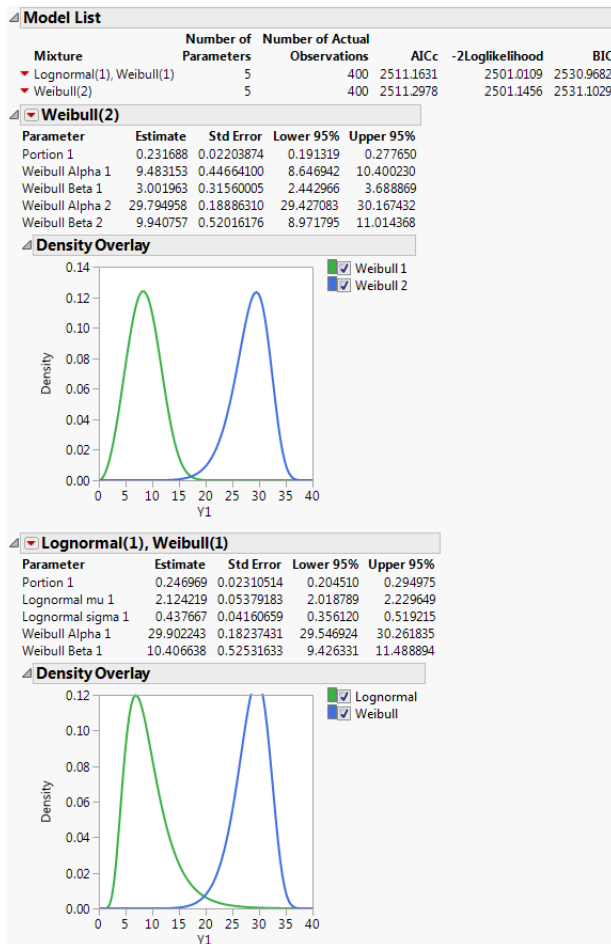


JMP fits a mixture model consisting of two Weibull components. Portion 1 is estimated as 0.231688, indicating that approximately 23% of observations have the Weibull distribution with $\alpha = 9.483153$ and $\beta = 3.001963$. The remaining 77% are estimated to come from the second Weibull distribution.

To compare this model to another, you can change the Ingredient selections and the Quantity of components.

9. Type 1 next to **Lognormal** and 1 next to **Weibull**.
10. Click **Go**.

Figure 3.21 Fit Mixture for Lognormal(1), Weibull(1)



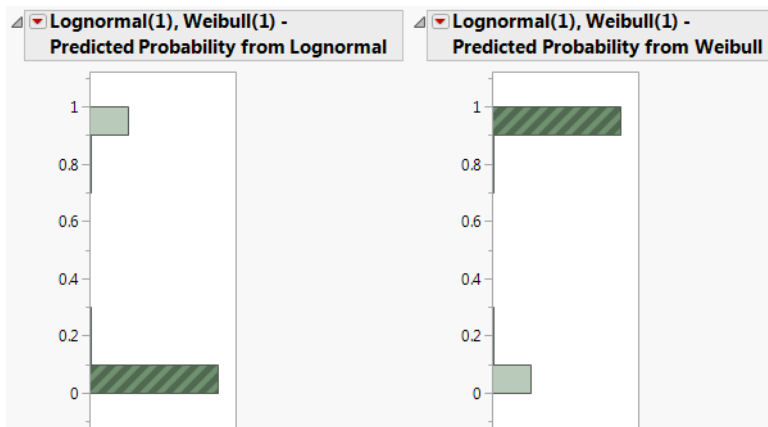
The Overlay plot is updated to show both mixture models. The plots and statistics in the Model List indicate that the Lognormal(1), Weibull(1) mixture seems to give a fit that is very similar to the Weibull(2) mixture.

Identify Observations Belonging to a Cluster

1. Click the red triangle next to Lognormal(1), Weibull(1) and select **Save Predictions**.
Two columns are added to the data table:
 - Lognormal(1), Weibull(1) - Predicted Probability from Lognormal

- Lognormal(1), Weibull(1) - Predicted Probability from Weibull
- 2. Select **Analyze > Distribution**.
- 3. Select the two new columns from the Select Columns list and click **Y, Columns**.
- 4. Check **Histograms Only**.
- 5. Click **OK**.
- 6. In the histogram for Lognormal(1), Weibull(1) - Predicted Probability from Weibull, click in the bar corresponding to the value near 1.

Figure 3.22 Histograms for Mixture Probabilities



In the data table, the 297 corresponding rows are selected. These are the observations that are likely to have come from the Weibull distribution with parameters $\alpha = 29.90$ and $\beta = 10.41$.

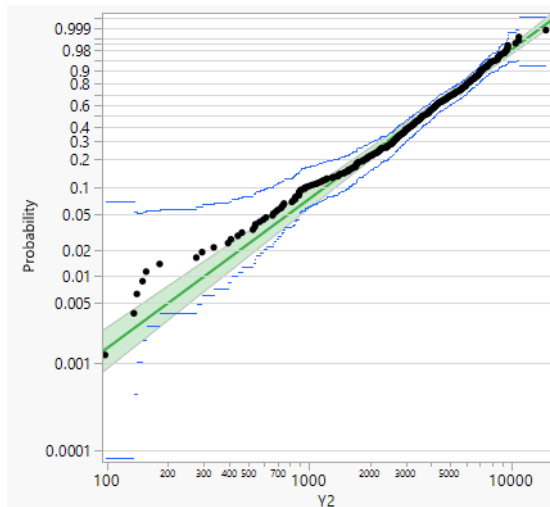
Fit Competing Risk Mixture Example

In this example, you compare the fit of a single Weibull distribution to a competing risk mixture distribution with two Weibull fits.

1. Select **Help > Sample Data Library** and open Reliability/Mixture Demo.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Y2 and click **Y, Time to Event**.
4. Click **OK**.
5. In the Compare Distributions report, select **Weibull** distribution and the corresponding **Scale** radio button.

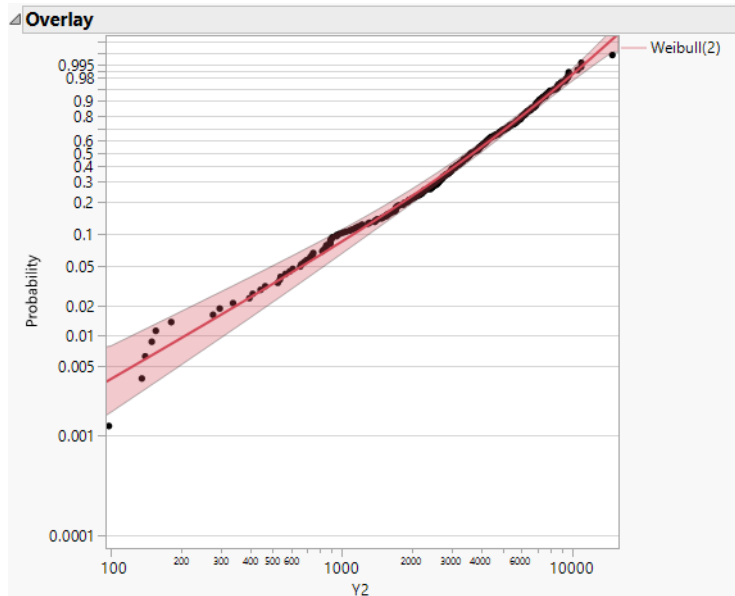
A probability plot for a single Weibull distribution fit appears. Note that the fit is not very good in the lower part of the range of Y2.

Figure 3.23 Weibull Distribution Fit



6. Click the Life Distribution red triangle and select **Fit Competing Risk Mixture**.
7. Scroll down to the Competing Risk Mixture report. Next to **Weibull**, type 2 in the **Quantity** box.
8. Click **Go**.
9. Scroll up to the Compare Distributions report. In the probability plot, right-click the vertical axis and select **Edit > Copy Axis Settings**.
10. Scroll down to the Competing Risk Mixture report. In the overlay plot, right-click the vertical axis and select **Edit > Paste Axis Settings**.
11. Do the same for the horizontal axis. In the probability plot, right-click the horizontal axis and select **Edit > Copy Axis Settings**.
12. In the overlay plot, right-click the horizontal axis and select **Edit > Paste Axis Settings**.

The probability plot for the Weibull(2) distribution fit appears. Note that the mixture of two Weibull distributions helps better capture the distribution in the lower part of the range of Y2.

Figure 3.24 Weibull(2) Competing Risk Mixture Distribution Fit

Statistical Details for the Life Distribution Platform

This section covers the following topics:

- The distributions used in the Life Distribution platform. See [“Distributions”](#) on page 82.
- Technical information about the Competing Cause report. See [“Competing Cause Details”](#) on page 97.
- Technical information about median rank regression. See [“Notes Regarding Median Rank Regression”](#) on page 102.

Distributions

This section provides details for the distributional fits in the Life Distribution platform.

Meeker and Escobar (1998, ch. 2-5) is an excellent source of theory, application, and discussion for both the nonparametric and parametric details that follow.

Estimation and Confidence Intervals

The parameters of all distributions, unless otherwise noted, are estimated using maximum likelihood estimates (MLEs). The only exceptions are the threshold distributions. If the smallest observation is an exact failure, then this observation is treated as interval-censored with a small interval. The parameter estimates are the MLEs estimated from this slightly modified data set. Without this modification, the likelihood can be unbounded, so an MLE might not exist. This approach is similar to that described in Meeker and Escobar (1998, p. 275), except that only the smallest exact failure is censored. This is the minimal change to the data that guarantees boundedness of the likelihood function.

The Life Distribution platform offers two methods for calculating confidence intervals for the distribution parameters. These methods are labeled as Wald or Likelihood and can be selected in the launch window for the Life Distribution platform. Wald confidence intervals are used as the default setting. The computations for the confidence intervals for the cumulative distribution function (cdf) start with Wald confidence intervals on the standardized variable. Next, the intervals are transformed to the cdf scale (Nelson 1982, pp. 332–333 and pp. 346–347). The confidence intervals given in the other graphs and profilers are transformed Wald intervals (Meeker and Escobar 1998, ch. 7). Joint confidence intervals for the parameters of a two-parameter distribution are shown in the log-likelihood contour plots. They are based on approximate likelihood ratios for the parameters (Meeker and Escobar 1998, ch. 8).

Nonparametric Fit

A nonparametric fit describes the basic curve of a distribution. For data with no censoring (failures only) and for data where the observations consist of both failures and right-censoring, JMP uses Kaplan-Meier estimates. For mixed, interval, or left censoring, JMP uses Turnbull estimates. When your data set contains only right-censored data, the Nonparametric Estimate report indicates that the nonparametric estimate cannot be calculated.

The Life Distribution platform uses the midpoint estimates of the step function to construct probability plots. The midpoint estimate is halfway between (or the average of) the current and previous Kaplan-Meier estimates.

Parametric Distributions

Parametric distributions provide a more concise distribution fit than nonparametric distributions. The estimates of failure-time distributions are also smoother. Parametric models are also useful for extrapolation (in time) to the lower or upper tails of a distribution.

Note: Many distributions in the Life Distribution platform are parameterized by location and scale. For lognormal fits, the median is also provided. A threshold parameter is also included in threshold distributions. Location corresponds to μ , scale corresponds to σ , and threshold corresponds to γ .

Lognormal

Lognormal distributions are used commonly for failure times when the range of the data is several powers of 10. This distribution is often considered as the multiplicative product of many small positive identically independently distributed random variables. It is reasonable when the log of the data values appears normally distributed. Examples of data appropriately modeled by the lognormal distribution include hospital cost data, metal fatigue crack growth, and the survival time of bacteria subjected to disinfectants. The pdf curve is usually characterized by strong right-skewness. The lognormal pdf and cdf are:

$$f(x; \mu, \sigma) = \frac{1}{x\sigma} \phi_{\text{nor}}\left[\frac{\log(x) - \mu}{\sigma}\right], \quad x > 0$$

$$F(x; \mu, \sigma) = \Phi_{\text{nor}}\left[\frac{\log(x) - \mu}{\sigma}\right]$$

where

$$\phi_{\text{nor}}(z) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right)$$

and

$$\Phi_{\text{nor}}(z) = \int_{-\infty}^z \phi_{\text{nor}}(w) dw$$

are the pdf and cdf, respectively, for the standardized normal, or $\text{nor}(\mu = 0, \sigma = 1)$ distribution.

Weibull

The Weibull distribution can be used to model failure time data with either an increasing or a decreasing hazard rate. It is used frequently in reliability analysis because of its tremendous flexibility in modeling many different types of data, based on the values of the shape parameter, β . This distribution has been successfully used for describing the failure of electronic components, roller bearings, capacitors, and ceramics. Various shapes of the Weibull distribution can be revealed by changing the scale parameter, α , and the shape parameter, β . The Weibull pdf and cdf are commonly represented as follows:

$$f(x;\alpha,\beta) = \frac{\beta}{\alpha} x^{(\beta-1)} \exp\left[-\left(\frac{x}{\alpha}\right)^\beta\right]; \quad x > 0, \alpha > 0, \beta > 0$$

$$F(x;\alpha,\beta) = 1 - \exp\left[-\left(\frac{x}{\alpha}\right)^\beta\right]$$

where α is a scale parameter, and β is a shape parameter. The Weibull distribution is particularly versatile because it reduces to an exponential distribution when $\beta = 1$. An alternative parameterization commonly used in the literature and in JMP is to use σ as the scale parameter and μ as the location parameter. These are easily converted to an α and β parameterization using the following definitions:

$$\alpha = \exp(\mu)$$

and

$$\beta = \frac{1}{\sigma}$$

The pdf and the cdf of the Weibull distribution are also expressed as a log-transformed smallest extreme value distribution (SEV). This uses a location scale parameterization, with $\mu = \log(\alpha)$ and $\sigma = 1/\beta$,

$$f(x;\mu,\sigma) = \frac{1}{x\sigma} \phi_{\text{sev}}\left[\frac{\log(x) - \mu}{\sigma}\right], \quad x > 0, \sigma > 0$$

$$F(x;\mu,\sigma) = \Phi_{\text{sev}}\left[\frac{\log(x) - \mu}{\sigma}\right]$$

where

$$\phi_{\text{sev}}(z) = \exp[z - \exp(z)]$$

and

$$\Phi_{\text{sev}}(z) = 1 - \exp[-\exp(z)]$$

are the pdf and cdf, respectively, for the standardized smallest extreme value ($\mu = 0$, $\sigma = 1$) distribution.

Loglogistic

The pdf of the loglogistic distribution is similar in shape to the lognormal distribution but has heavier tails. It is often used to model data exhibiting non-monotonic hazard functions, such as cancer mortality and financial wealth. The loglogistic pdf and cdf are:

$$f(x; \mu, \sigma) = \frac{1}{x\sigma} \phi_{\text{logis}} \left[\frac{\log(x) - \mu}{\sigma} \right]$$

$$F(x; \mu, \sigma) = \Phi_{\text{logis}} \left[\frac{\log(x) - \mu}{\sigma} \right]$$

where

$$\phi_{\text{logis}}(z) = \frac{\exp(z)}{[1 + \exp(z)]^2}$$

and

$$\Phi_{\text{logis}}(z) = \frac{\exp(z)}{[1 + \exp(z)]} = \frac{1}{1 + \exp(-z)}$$

are the pdf and cdf, respectively, for the standardized logistic or logis distribution ($\mu = 0$, $\sigma = 1$).

Fréchet

The Fréchet distribution is known as a log-largest extreme value distribution or sometimes as a Fréchet distribution of maxima when it is parameterized as the reciprocal of a Weibull distribution. This distribution is commonly used for financial data. The pdf and cdf are:

$$f(x; \mu, \sigma) = \exp \left[-\exp \left(-\frac{\log(x) - \mu}{\sigma} \right) \right] \exp \left(-\frac{\log(x) - \mu}{\sigma} \right) \frac{1}{x\sigma}$$

$$F(x; \mu, \sigma) = \exp \left[-\exp \left(-\frac{\log(x) - \mu}{\sigma} \right) \right]$$

and are more generally parameterized as follows:

$$f(x;\mu, \sigma) = \frac{1}{x\sigma} \phi_{\text{lev}}\left[\frac{\log(x) - \mu}{\sigma}\right]$$

$$F(x;\mu, \sigma) = \Phi_{\text{lev}}\left[\frac{\log(x) - \mu}{\sigma}\right]$$

where

$$\phi_{\text{lev}}(z) = \exp[-z - \exp(-z)]$$

and

$$\Phi_{\text{lev}}(z) = \exp[-\exp(-z)]$$

are the pdf and cdf, respectively, for the standardized largest extreme value LEV($\mu = 0$, $\sigma = 1$) distribution.

Normal

The normal distribution is the most widely used distribution in most areas of statistics because of its relative simplicity and the ease of applying the central limit theorem. However, it is rarely used in reliability. It is most useful for data where $\mu > 0$ and the coefficient of variation (σ / μ) is small. Because the hazard function increases with no upper bound, it is particularly useful for data exhibiting wear-out failure. Examples include incandescent light bulbs, toaster heating elements, and mechanical strength of wires. The pdf and cdf are:

$$f(x;\mu, \sigma) = \frac{1}{\sigma} \phi_{\text{nor}}\left(\frac{x - \mu}{\sigma}\right), \quad -\infty < x < \infty$$

$$F(x;\mu, \sigma) = \Phi_{\text{nor}}\left(\frac{x - \mu}{\sigma}\right)$$

where

$$\phi_{\text{nor}}(z) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right)$$

and

$$\Phi_{\text{nor}}(z) = \int_{-\infty}^z \phi_{\text{nor}}(w) dw$$

are the pdf and cdf, respectively, for the standardized normal, or nor($\mu = 0$, $\sigma = 1$) distribution.

Smallest Extreme Value (SEV)

This non-symmetric (left-skewed) distribution is useful in two cases. The first case is when the data indicate a small number of weak units in the lower tail of the distribution (the data indicate the smallest number of many observations). The second case is when σ is small relative to μ , because probabilities of being less than zero, when using the SEV distribution, are small. The smallest extreme value distribution is useful to describe data whose hazard rate becomes larger as the unit becomes older. Examples include human mortality of the aged and rainfall amounts during a drought. This distribution is sometimes referred to as a Gumbel distribution. The pdf and cdf are:

$$f(x;\mu,\sigma) = \frac{1}{\sigma} \phi_{\text{sev}}\left(\frac{x-\mu}{\sigma}\right), \quad -\infty < \mu < \infty, \quad \sigma > 0$$

$$F(x;\mu,\sigma) = \Phi_{\text{sev}}\left(\frac{x-\mu}{\sigma}\right)$$

where

$$\phi_{\text{sev}}(z) = \exp[z - \exp(z)]$$

and

$$\Phi_{\text{sev}}(z) = 1 - \exp[-\exp(z)]$$

are the pdf and cdf, respectively, for the standardized smallest extreme value, $\text{SEV}(\mu = 0, \sigma = 1)$ distribution.

Logistic

The logistic distribution has a shape similar to the normal distribution, but with longer tails. The logistic distribution is often used to model life data when negative failure times are not an issue. Logistic regression models for a binary or ordinal response assume the logistic distribution as the latent distribution. The pdf and cdf are:

$$f(x;\mu,\sigma) = \frac{1}{\sigma} \phi_{\text{logis}}\left(\frac{x-\mu}{\sigma}\right), \quad -\infty < \mu < \infty \text{ and } \sigma > 0$$

$$F(x;\mu,\sigma) = \Phi_{\text{logis}}\left(\frac{x-\mu}{\sigma}\right)$$

where

$$\phi_{\text{logis}}(z) = \frac{\exp(z)}{[1 + \exp(z)]^2}$$

and

$$\Phi_{\text{logis}}(z) = \frac{\exp(z)}{[1 + \exp(z)]} = \frac{1}{1 + \exp(-z)}$$

are the pdf and cdf, respectively, for the standardized logistic or logis distribution ($\mu = 0$, $\sigma = 1$).

Largest Extreme Value (LEV)

This right-skewed distribution can be used to model failure times if σ is small relative to $\mu > 0$. This distribution is not commonly used in reliability but is useful for estimating natural extreme phenomena, such as a catastrophic flood heights or extreme wind velocities. The pdf and cdf are:

$$f(x; \mu, \sigma) = \frac{1}{\sigma} \phi_{\text{lev}}\left(\frac{x - \mu}{\sigma}\right), \quad -\infty < \mu < \infty \text{ and } \sigma > 0$$

$$F(x; \mu, \sigma) = \Phi_{\text{lev}}\left(\frac{x - \mu}{\sigma}\right)$$

where

$$\phi_{\text{lev}}(z) = \exp[-z - \exp(-z)]$$

and

$$\Phi_{\text{lev}}(z) = \exp[-\exp(-z)]$$

are the pdf and cdf, respectively, for the standardized largest extreme value LEV($\mu = 0$, $\sigma = 1$) distribution.

Exponential

Both one- and two-parameter exponential distributions are used in reliability. The pdf and cdf for the two-parameter exponential distribution are:

$$f(x; \theta, \gamma) = \frac{1}{\theta} \exp\left(-\frac{x - \gamma}{\theta}\right), \quad \theta > 0$$

$$F(x; \theta, \gamma) = 1 - \exp\left(-\frac{x - \gamma}{\theta}\right)$$

where θ is a scale parameter and γ is both the threshold and the location parameter. Reliability analysis frequently uses the one-parameter exponential distribution, with $\gamma = 0$. The exponential distribution is useful for describing failure times of components exhibiting wear-out far beyond their expected lifetimes. This distribution has a constant failure rate, which means that for small time increments, failure of a unit is independent of the unit's age. The exponential distribution should not be used for describing the life of mechanical components that can be exposed to fatigue, corrosion, or short-term wear. This distribution is, however, appropriate for modeling certain types of robust electronic components. It has been used successfully to describe the life of insulating oils and dielectric fluids (Nelson 1990, p. 53).

Log Generalized Gamma (LogGenGamma)

The log generalized gamma distribution contains the SEV, LEV, and Normal. The pdf and cdf are:

$$f(x; \mu, \sigma, \lambda) = \begin{cases} \frac{|\lambda|}{\sigma} \phi_{lg}[\lambda \omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda \neq 0 \\ \frac{1}{\sigma} \phi_{nor}(\omega) & \text{if } \lambda = 0 \end{cases}$$

$$F(x; \mu, \sigma, \lambda) = \begin{cases} \Phi_{lg}[\lambda \omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda > 0 \\ \Phi_{nor}(\omega) & \text{if } \lambda = 0 \\ 1 - \Phi_{lg}[\lambda \omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda < 0 \end{cases}$$

where $-\infty < x < \infty$, $\omega = [x - \mu]/\sigma$, and

$$-\infty < \mu < \infty, \quad -12 < \lambda < 12, \quad \text{and } \sigma > 0.$$

Note that

$$\phi_{lg}(z; \kappa) = \frac{1}{\Gamma(\kappa)} \exp[\kappa z - \exp(z)]$$

$$\Phi_{lg}(z; \kappa) = \Gamma_I[\exp(z); \kappa]$$

are the pdf and cdf, respectively, for the log-gamma variable and $\kappa > 0$ is a shape parameter. The standardized distributions above are dependent upon the shape parameter κ .

Note: In JMP, the shape parameter, λ , for the generalized gamma distribution is bounded between [-12,12] to provide numerical stability.

Extended Generalized Gamma (GenGamma)

The extended generalized gamma distribution can include many other distributions as special cases, such as the generalized gamma, Weibull, lognormal, Fréchet, gamma, and exponential. It is particularly useful for cases with little or no censoring. This distribution has been successfully modeled for human cancer prognosis. The pdf and cdf are:

$$f(x; \mu, \sigma, \lambda) = \begin{cases} \frac{|\lambda|}{x\sigma} \phi_{lg}[\lambda\omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda \neq 0 \\ \frac{1}{x\sigma} \phi_{nor}(\omega) & \text{if } \lambda = 0 \end{cases}$$

$$F(x; \mu, \sigma, \lambda) = \begin{cases} \Phi_{lg}[\lambda\omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda > 0 \\ \Phi_{nor}(\omega) & \text{if } \lambda = 0 \\ 1 - \Phi_{lg}[\lambda\omega + \log(\lambda^{-2}); \lambda^{-2}] & \text{if } \lambda < 0 \end{cases}$$

where $x > 0$, $\omega = [\log(x) - \mu]/\sigma$, and

$$-\infty < \mu < \infty, \quad -12 < \lambda < 12, \quad \text{and } \sigma > 0.$$

Note that

$$\phi_{lg}(z; \kappa) = \frac{1}{\Gamma(\kappa)} \exp[\kappa z - \exp(z)]$$

$$\Phi_{lg}(z; \kappa) = \Gamma_I[\exp(z); \kappa]$$

are the pdf and cdf, respectively, for the standardized log-gamma variable and $\kappa > 0$ is a shape parameter.

The standardized distributions above are dependent upon the shape parameter κ . Meeker and Escobar (1998, ch. 5) give a detailed explanation of the extended generalized gamma distribution.

Note: In JMP, the shape parameter, λ , for the generalized gamma distribution is bounded between $[-12, 12]$ to provide numerical stability.

Distributions with Threshold Parameters

Threshold Distributions are log-location-scale distributions with threshold parameters. Some of the distributions above are generalized by adding a threshold parameter, denoted by γ . The addition of this threshold parameter shifts the left endpoint of the distribution away from 0. Threshold parameters are sometimes called shift, minimum, or guarantee parameters because all units survive at least until threshold time. Note that while adding a threshold parameter shifts the distribution on the time axis, the shape, and spread of the distribution are not affected. Threshold distributions are useful for fitting moderately to heavily shifted distributions. The general forms for the pdf and cdf of a log-location-scale threshold distribution are:

$$f(x; \mu, \sigma, \gamma) = \frac{1}{\sigma(x - \gamma)} \phi \left[\frac{\log(x - \gamma) - \mu}{\sigma} \right], \quad x > \gamma$$

$$F(x; \mu, \sigma, \gamma) = \Phi \left[\frac{\log(x - \gamma) - \mu}{\sigma} \right]$$

where ϕ and Φ are the pdf and cdf, respectively, for the specific distribution. Examples of specific threshold distributions are shown below for the Weibull, lognormal, Fréchet, and loglogistic distributions, where, respectively, the SEV, Normal, LEV, and logis pdfs and cdfs are appropriately substituted.

Note: If the smallest observation is a failure (not censored), JMP creates a small interval around the point and treats the observation as interval censored. This padding around the failure bounds the log-likelihood function and improves estimation. If the smallest observation is censored, then no extra padding is added to the observation.

TH Weibull

The pdf and cdf of the three-parameter Weibull distribution are:

$$f(x; \mu, \sigma, \gamma) = \frac{1}{(x - \gamma)\sigma} \phi_{\text{sev}} \left[\frac{\log(x - \gamma) - \mu}{\sigma} \right], \quad x > \gamma, \sigma > 0$$

$$F(x; \mu, \sigma, \gamma) = \Phi_{\text{sev}} \left(\frac{\log(x - \gamma) - \mu}{\sigma} \right) = 1 - \exp \left[- \left(\frac{x - \gamma}{\alpha} \right)^\beta \right], \quad x > \gamma$$

where $\mu = \log(\alpha)$, and $\sigma = 1/\beta$ and where

$$\phi_{\text{sev}}(z) = \exp[z - \exp(z)]$$

and

$$\Phi_{\text{sev}}(z) = 1 - \exp[-\exp(z)]$$

are the pdf and cdf, respectively, for the standardized smallest extreme value, SEV($\mu = 0, \sigma = 1$) distribution.

TH Lognormal

The pdf and cdf of the three-parameter lognormal distribution are:

$$f(x; \mu, \sigma, \gamma) = \frac{1}{\sigma(x - \gamma)} \phi_{\text{nor}}\left[\frac{\log(x - \gamma) - \mu}{\sigma}\right], \quad x > \gamma$$

$$F(x; \mu, \sigma, \gamma) = \Phi_{\text{nor}}\left[\frac{\log(x - \gamma) - \mu}{\sigma}\right]$$

where

$$\phi_{\text{nor}}(z) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right)$$

and

$$\Phi_{\text{nor}}(z) = \int_{-\infty}^z \phi_{\text{nor}}(w) dw$$

are the pdf and cdf, respectively, for the standardized normal, or N($\mu = 0, \sigma = 1$) distribution.

TH Fréchet

The pdf and cdf of the three-parameter Fréchet distribution are:

$$f(x; \mu, \sigma, \gamma) = \frac{1}{\sigma(x - \gamma)} \phi_{\text{lev}}\left[\frac{\log(x - \gamma) - \mu}{\sigma}\right], \quad x > \gamma$$

$$F(x; \mu, \sigma, \gamma) = \Phi_{\text{lev}}\left[\frac{\log(x - \gamma) - \mu}{\sigma}\right]$$

where

$$\phi_{\text{lev}}(z) = \exp[-z - \exp(-z)]$$

and

$$\Phi_{\text{lev}}(z) = \exp[-\exp(-z)]$$

are the pdf and cdf, respectively, for the standardized largest extreme value LEV($\mu = 0, \sigma = 1$) distribution.

TH Loglogistic

The pdf and cdf of the three-parameter loglogistic distribution are:

$$f(x; \mu, \sigma, \gamma) = \frac{1}{\sigma(x - \gamma)} \phi_{\logis} \left[\frac{\log(x - \gamma) - \mu}{\sigma} \right], \quad x > \gamma$$

$$F(x; \mu, \sigma, \gamma) = \Phi_{\logis} \left[\frac{\log(x - \gamma) - \mu}{\sigma} \right]$$

where

$$\phi_{\logis}(z) = \frac{\exp(z)}{[1 + \exp(z)]^2}$$

and

$$\Phi_{\logis}(z) = \frac{\exp(z)}{[1 + \exp(z)]} = \frac{1}{1 + \exp(-z)}$$

are the pdf and cdf, respectively, for the standardized logistic or logis distribution ($\mu = 0, \sigma = 1$).

Distributions for Defective Subpopulations

In reliability experiments, there are times when only a fraction of the population has a particular defect leading to failure. Because all units are not susceptible to failure, using the regular failure distributions is inappropriate and might produce misleading results. Use the DS distribution options to model failures that occur on only a subpopulation. The following DS distributions are available:

- DS Lognormal
- DS Weibull
- DS Loglogistic
- DS Fréchet

The pdf and cdf for defective subpopulation distributions are defined as follows:

$$f(t) = \left[p \frac{1}{t\sigma} \right] \phi \left[\frac{(\log(t) - \mu)}{\sigma} \right]$$

$$F(t) = p\Phi\left[\left(\frac{\log(t) - \mu}{\sigma}\right)\right]$$

where:

p is the defective subpopulation fraction

t is the time of measurement for the lifetime event

μ and σ are estimated by calculating the usual maximum likelihood estimations using the pdf and cdf of the corresponding defective subpopulation

$\phi(z)$ and $\Phi(z)$ are the density and cumulative distribution function, respectively, for a standard distribution. For example, for a Weibull distribution,

$$\phi(z) = \exp(z - \exp(z)) \text{ and } \Phi(z) = 1 - \exp(-\exp(z)).$$

See Tobias and Trindad (2012, p. 321) for more information about defective subpopulation models.

The defective subpopulation model is also known as a *limited failure population* model in Meeker and Escobar (1998, ch. 11).

Zero-Inflated Distributions

Zero-inflated distributions are used when some proportion (p) of the data fail at $t = 0$. When the data contain more zeros than expected by a standard model, the number of zeros is inflated. When the time-to-event data contain zero as the minimum value in the Life Distribution platform, four zero-inflated distributions are available. These distributions include:

- Zero-Inflated Lognormal (ZI Lognormal)
- Zero-Inflated Weibull (ZI Weibull)
- Zero-Inflated Loglogistic (ZI Loglogistic)
- Zero-Inflated Fréchet (ZI Fréchet)

The pdf and cdf for zero-inflated distributions are defined as follows:

$$f(t) = \left[(1-p)\frac{1}{t\sigma}\right]\phi\left[\left(\frac{\log(t) - \mu}{\sigma}\right)\right]$$

$$F(t) = p + (1-p)\Phi\left[\left(\frac{\log(t) - \mu}{\sigma}\right)\right]$$

where:

p is the proportion of zero data values

t is the time of measurement for the lifetime event

μ and σ are estimated by calculating the usual maximum likelihood estimations after removing zero values from the original data

$\phi(z)$ and $\Phi(z)$ are the density and cumulative distribution function, respectively, for a standard distribution. For example, for a Weibull distribution,

$$\phi(z) = \exp(z - \exp(z)) \text{ and } \Phi(z) = 1 - \exp(-\exp(z)).$$

See Lawless (2003, p. 34) for more information about zero-inflated distributions. Substitute $p = 1 - p$ and $S_1(t) = 1 - \Phi(t)$ to obtain the form shown above.

See Tobias and Trindade (1995, p. 232) for more information about reliability distributions. This reference gives the general form for mixture distributions. Using the parameterization in Tobias and Trindade, the form above can be found by substituting $\alpha = p$, $F_d(t) = 1$, and $F_N(t) = \Phi(t)$.

Prior Distributions for Bayesian Estimation

The following distributions are available for Location Scale Priors:

- Normal/Lognormal, with hyperparameters Location (μ) and Scale (σ). For a definition, see “Lognormal” on page 84 and “Normal” on page 87.
- Uniform, with hyperparameters Low and End, which define the support of a Uniform distribution.
- Gamma, with hyperparameters Shape and Scale. The k/θ parameterization and the probability density function is used.
- Point Mass, with hyperparameter Location. This is a degenerated prior; there is only one possible value for the parameter that we are assigning a prior distribution to. The only possible value equals the value that is entered to this Location hyperparameter.

The following distributions are available for Quantile Parameter Priors:

- Normal/Lognormal, with a 99% probability range, specifies the prior distribution using the 0.005 and 0.995 percentiles of the distribution. JMP backs out the μ and σ .
- Uniform, with hyperparameters Lower and Upper Limits, which define the support of a Uniform distribution.
- Log-Uniform, with Lower (a) and Upper (b) Limits. This distribution is uniform on the log scale between $\text{Log}(a)$ and $\text{Log}(b)$.
- Point Mass, with hyperparameter Location. This is a degenerated prior; there is only one possible value for the parameter that we are assigning a prior distribution to. The only possible value equals the value that is entered to this Location hyperparameter.

The following distributions are available for Failure Probability Priors:

- Beta, characterized by the probability density function.

- Specify the Beta prior using estimates and error percentages (mean and variance). The mean equals the number entered in to the Estimate, and the variance equals (Error Percentage / 100 * Estimate)^2.
- Specify the Beta prior using 0.005 and 0.995 percentiles of the distribution. JMP backs out the hyperparameters.

Competing Cause Details

For a competing cause model, a closed form for the aggregated distribution is defined as follows:

$$F(x) = 1 - \prod_{i=1}^k [1 - F_i(x)]$$

where the $F_i(x)$, $i = 1, \dots, k$, are individual failure distributions corresponding to causes. Confidence limits are readily available, because all involved estimates are MLEs.

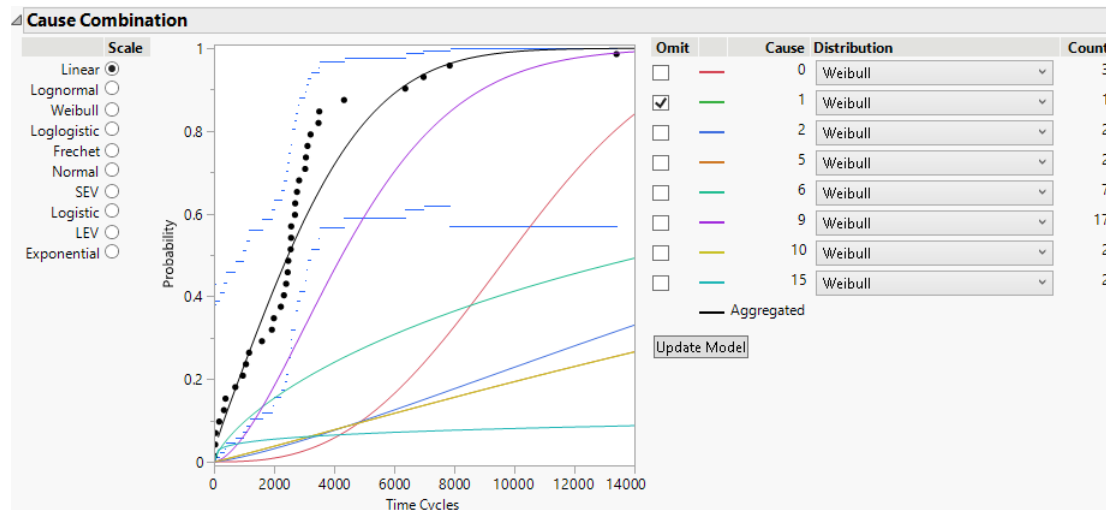
Specify a Fixed Parameter Model as a Distribution for a Cause

If a Fixed Parameter model is specified for a cause, you must fix the parameter in the Individual Causes report for that cause. Fix the parameter in the desired Parametric Estimate report in the Life Distribution - Failure Cause: <Name> report. The fixed parameter becomes part of the aggregated distribution when you click Update Model.

This example illustrates how to include the Fixed Parameter model into the aggregated distribution:

1. Select **Help > Sample Data Library** and open Reliability/Appliance.jmp.
2. Select **Analyze > Reliability and Survival > Life Distribution**.
3. Select Time Cycles and click **Y, Time to Event**.
4. Select Cause Code and click **Failure Cause**.
5. Select **Likelihood** as the Confidence Interval Method.
6. Select **Allow failure mode to use fixed parameter models**.
7. Click **OK**.

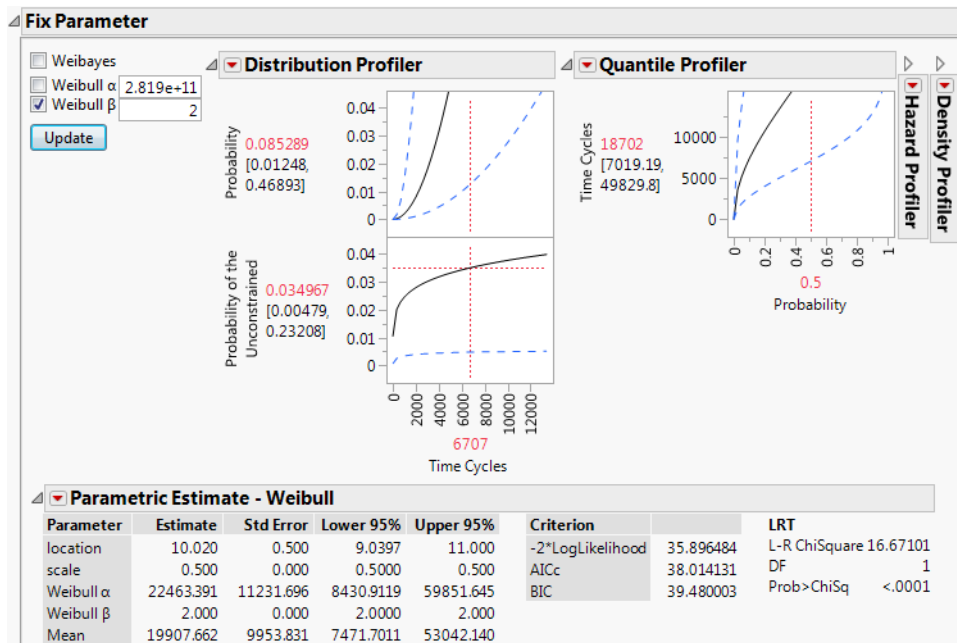
Figure 3.25 Fixed Parameter Model with Cause 1 Omitted



By default, Cause = 1 is omitted, because there are not enough data. However, you do not want this cause to be omitted.

8. Open the Individual Causes report for Cause 1. The report is called Life Distribution - Failure Cause: 1 Failure Counts: 1.
9. Click the red triangle next to Parametric Estimate - Weibull and select **Fix Parameter**.
10. Select **Weibull beta** and type 2.
11. Click **Update**.

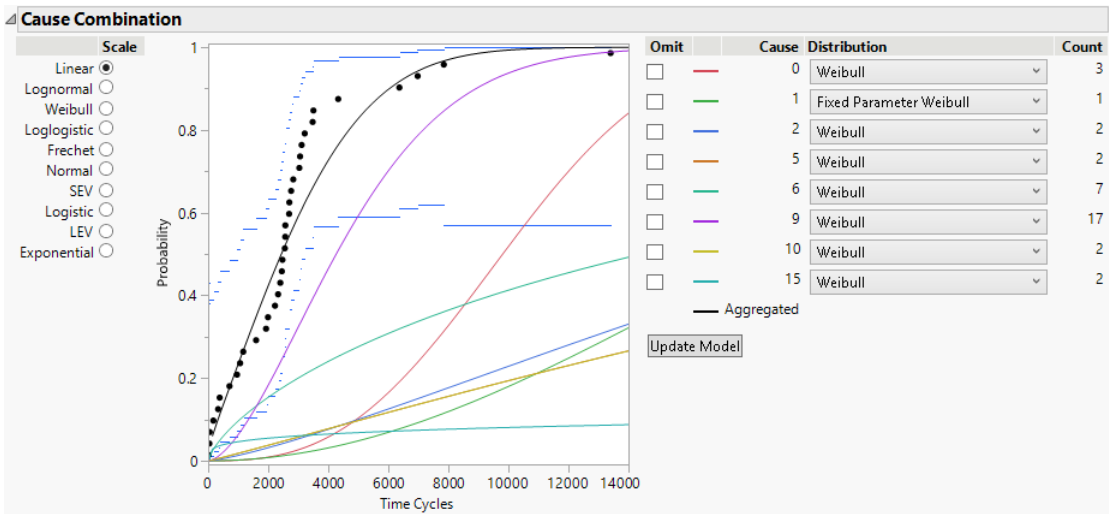
Figure 3.26 Fixed Parameter Model with Weibull Beta Specified



In the Parametric Estimate - Weibull report, assuming β equals 2, the alpha parameter is estimated to be 22463.391. Now you can use this for the failure distribution for Cause=1.

12. Scroll up to Cause Combination at the top of the report window.
13. Deselect **Omit** for Cause 1.
14. For the distribution for Cause 1, select **Fixed Parameter Weibull**.
15. Click **Update Model**.

Figure 3.27 Updated Model Showing Cause 1



Now the aggregated model uses the Fixed Parameter Weibull results for Cause 1 in the overall competing cause model.

Specify a Bayesian Model for a Cause

The steps for specifying a Bayesian model for a cause are similar to those described in “Specify a Fixed Parameter Model as a Distribution for a Cause” on page 97. Define the model in the desired Bayesian Estimation report found in the corresponding Parametric Estimate outline under Statistics in the Life Distribution report for the individual cause. See “Bayesian Estimation - <Distribution Name>” on page 48.

To incorporate a Bayesian model into the aggregated model, non-Bayesian distributions for other causes must be amenable to a simulation-based framework. For example, suppose that a model has two failure causes. One is modeled using a Weibull distribution and the other using a Bayesian approach for estimating the parameters of a second Weibull. The parameters for the first Weibull distribution, denoted by the vector θ_1 , are estimated using maximum likelihood. The parameters for the second Weibull, θ_2 , are estimated using the Bayesian approach.

The quantiles and median of the aggregated mixture distribution, denoted $F(x, \theta_1, \theta_2)$, are obtained as follows:

- A parametric bootstrap is performed for the first Weibull, yielding random samples from the asymptotic distribution of the maximum likelihood estimate θ_1 . Denote a sampled value from the asymptotic distribution of $\hat{\theta}_1$ by θ_1^* .
- A sample is drawn from the posterior distribution of θ_2 , denoted by θ_2^* .

- For each set of values θ_1^* and θ_2^* , an estimate of $F(x, \theta_1, \theta_2)$, denoted by $F^*(x, \theta_1, \theta_2)$, is obtained.
- The values $F^*(x, \theta_1, \theta_2)$ are used to obtain estimates of the quantiles and median of the aggregated distribution. These are the values displayed in the Distribution profiler at a given value of x .

Specify a Weibayes Model for a Cause

The steps for specifying a Weibayes model for a cause are similar to those described in “Specify a Fixed Parameter Model as a Distribution for a Cause” on page 97. Select the Fix Parameter option in the Parametric Estimate - Weibull outline under Statistics in the Life Distribution report for the cause. In the Fix Parameter report, check the Weibayes option. The Weibayes model is treated as a Bayesian model and a bootstrap sample is drawn from the posterior distribution of the parameter alpha. See Liu and Wang (2013).

Mean Remaining Life Calculator

Use the Configuration option in the red triangle menu to set a value for the number of simulated failure times used in computing the mean remaining life. Denote this value by m .

To obtain an estimate of the mean remaining life at time t , m samples are drawn from the aggregated distribution conditioned on survival to time t . Their average is computed.

To compute the confidence limits for the mean remaining life, you must select the box in the Configuration window. You then have the option to set the number of bootstrap samples. Denote this value by n .

To compute the confidence interval, n samples of parameter estimates are drawn from either the asymptotic distributions of the MLEs, or the posterior distributions derived using Bayesian inference. For each sample of parameter values, an aggregated distribution is formed, from which m samples are drawn to compute a mean remaining life. The samples of n mean remaining life values are used to construct the confidence interval.

Fit Mixture Save Predictions Formulas

This section gives the formulas used in calculating values in the columns saved by the Fit Mixture report option Save Predictions.

Consider the following notation:

$\hat{p}_i$ is an estimate of the mixture proportion, w_i

$\hat{F}_i$ is the estimated probability distribution function F_i

$\hat{f}_i$ is the estimated probability density function for F_i

- If the observation y is not censored, the saved value is given by the following:

$$\frac{\hat{p}_i \hat{f}_i(y)}{\sum_{i=1} \hat{p}_i \hat{f}_i(y)}$$

- If the observation is censored, the saved value is obtained by replacing the estimated density values in the formula for an uncensored observation by the following:

$$\hat{F}_i(y) \text{ for right censoring}$$

$$1 - \hat{F}_i(y) \text{ for left censoring}$$

$$\hat{F}_i(y_{high}) - \hat{F}_i(y_{low}) \text{ for interval censoring}$$

Notes Regarding Median Rank Regression

When there are no censored rows and no Weight column, median rank regression (MRR) for Weibull parameter estimates is available in the Life Distribution platform. The following conditions must be met before MRR estimates appear in the Life Distribution report:

- This Life Distribution platform preference is selected: Report Median Rank Regression Based Weibull Parameter Estimates When There Are No Censored Observations.
- There are no censored observations in the analysis. In other words, all observations represent failure times.
- There is no Weight column specified.

If all of the above conditions are met, the Parametric Estimate - Weibull report is modified to show MRR estimates as well as the usual maximum likelihood estimates (MLE). The column headings in the parameter estimates table are prefixed with MLE or MRR to denote which estimation method was used to produce each column of estimates. There are no confidence intervals available for the MRR estimates.

Caution: The use of MRR estimates is not recommended. Median rank regression uses least squares estimation to produce estimates, instead of maximum likelihood. There are no commonly accepted methods to apply least squares estimation to censored data. Further, Genschel and Meeker (2010) show that the MLE method is more precise. They provide simulation results that are based on the fact that the true estimates are known, but in reality, the accuracy of MRR estimates is unknown. Because MRR is a loosely defined and ad hoc estimation procedure, different software packages produce different MRR estimates for the same data.

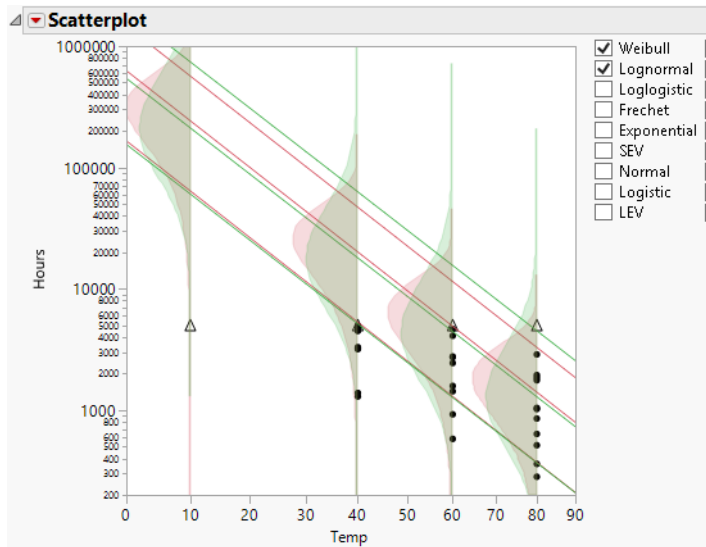
Chapter 4

Fit Life by X

Fit Single-Factor Models to Time-to Event Data

The Fit Life by X platform helps you analyze lifetime events when only one factor is present. You can choose to model the relationship between the event and the factor using various transformations, or create a custom transformation of your data. You also have the flexibility of comparing different distributions at the same factor level and comparing the same distribution across different factor levels.

Figure 4.1 Scatterplot Showing Varying Distributions and Factor Levels



Contents

Overview of the Fit Life by X Platform 107

Example of the Fit Life by X Platform 108

Launch the Fit Life by X Platform..... 110

The Fit Life by X Report..... 112

 Summary of Data 113

 Scatterplot..... 113

 Nonparametric Overlay 115

 Comparisons 116

 Results..... 120

 Custom Relationship 131

Fit Life by X Platform Options 132

Additional Examples of the Fit Life by X Platform 133

 Capacitor Example 133

 Custom Relationship Example 135

Overview of the Fit Life by X Platform

The Fit Life by X platform provides the tools needed for accelerated life-testing analysis. Accelerated tests are routinely used in industry to provide failure-time information about products or components in a relatively short time frame. Common accelerating factors include temperature, voltage, pressure, and usage rate. Results are extrapolated to obtain time-to-failure estimates at lower, normal operating levels of the accelerating factors. These results are used to assess reliability, detect and correct failure modes, compare manufacturers, and certify components.

The Fit Life by X platform includes many commonly used transformations to model physical and chemical relationships between the event and the factor of interest. Examples include transformation using Arrhenius (Celsius, Fahrenheit, and Kelvin) relationship time-acceleration factors and Voltage-acceleration mechanisms. Linear, Log, Logit, Reciprocal, Square Root, Box-Cox, Location, Location and Scale, and Custom acceleration models are also included in this platform.

You can use the DOE > Accelerated Life Test Design platform to design accelerated life test experiments. See the *Design of Experiments Guide*.

Meeker and Escobar (1998, p. 495) offer the following strategy for analyzing accelerated lifetime data:

1. Examine the data graphically. One useful way to visualize the data is by examining a scatterplot of the time-to-failure variable versus the accelerating factor.
2. Fit distributions individually to the data at different levels of the accelerating factor. Repeat for different assumed distributions.
3. Fit an overall model with a plausible relationship between the time-to-failure variable and the accelerating factor.
4. Compare the model in Step 3 with the individual analyses in Step 2, assessing the lack of fit for the overall model.
5. Perform residual and various diagnostic analyses to verify model assumptions.
6. Assess the plausibility of the data to make inferences.

Example of the Fit Life by X Platform

This example uses the Devalt.jmp sample data table, from Meeker and Escobar (1998), and can be found in the Reliability folder of the sample data. It contains time-to-failure data for a device at accelerated operating temperatures. No time-to-failure observation is recorded for the normal operating temperature of 10 degrees Celsius; all other observations are shown as time-to-failure or censored values at accelerated temperature levels of 40, 60, and 80 degrees Celsius.

1. Select **Help > Sample Data Library** and open Reliability/Devalt.jmp.
2. Select **Analyze > Reliability and Survival > Fit Life by X**.
3. Select Hours and click **Y, Time to Event**.
4. Select Temp and click **X**.
5. Select Censor and click **Censor**.
6. Leave the **Censor Code** as 1.
7. Select Weight and click **Freq**.
8. Keep **Arrhenius Celsius** as the relationship, and keep the **Nested Model Tests** option selected.
9. Select **Weibull** as the distribution.
10. Keep **Wald** as the confidence interval method.

Figure 4.2 Fit Life by X Launch Window

Models life distributions parameterized by a single regression factor.

Select Columns	Cast Selected Columns into Roles	Action
6 Columns Hours Status Weight Temp Censor x	Y, Time to Event: Hours <i>optional numeric</i> X: Temp Censor: Censor Freq: Weight By: <i>optional</i>	OK Cancel Remove Recall Help

Censor Code: 1

Relationship: Arrhenius Celsius

☒ Nested Model Tests

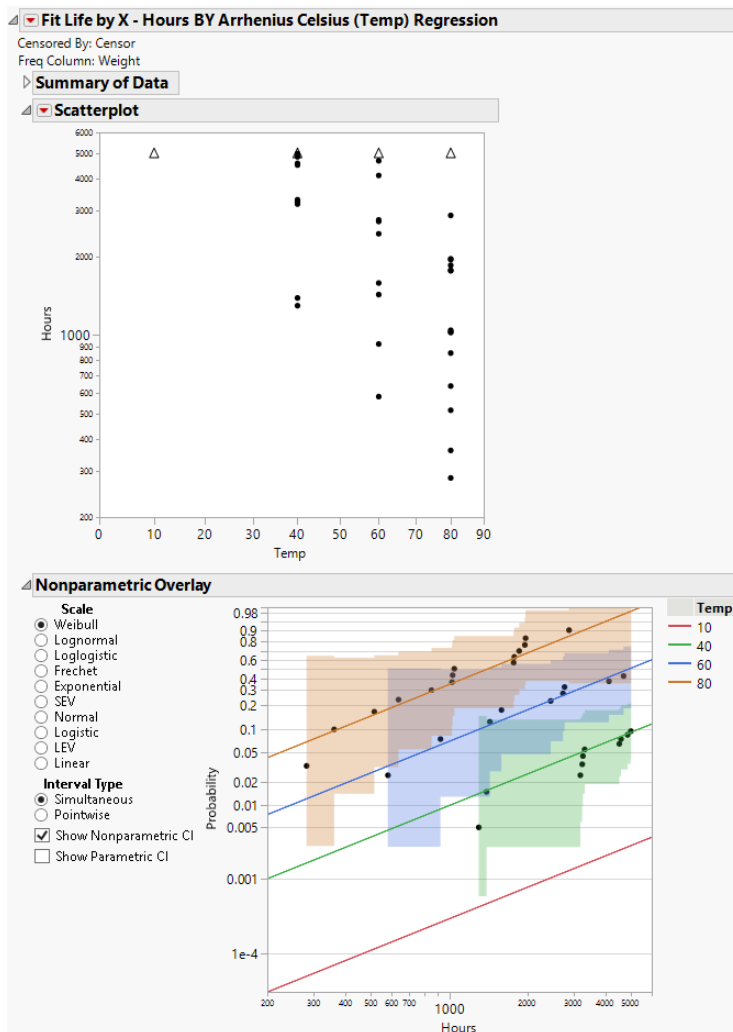
Use Condition: .

Distribution: Weibull

Select Confidence Interval Method: Wald

11. Click **OK**.

Figure 4.3 Fit Life by X Report Window for Devalt.jmp Data

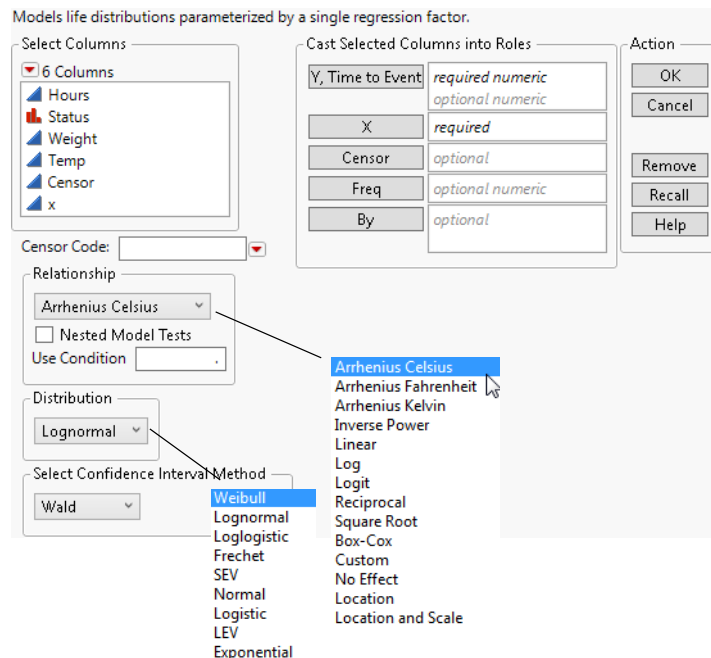


The report window shows summary data, diagnostic plots, comparison data and results, including detailed statistics and prediction profilers. Separate result sections are shown for each selected distribution. Distribution, Quantile, Hazard, Density, and Acceleration Factor Profilers are included for each of the specified distributions.

Launch the Fit Life by X Platform

Launch the Fit Life by X platform by selecting **Analyze > Reliability and Survival > Fit Life by X**. For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

Figure 4.4 The Fit Life by X Launch Window



The Fit Life by X launch window contains the following options:

Y, Time to Event Identifies the time to event (such as the time to failure) or time to censoring. With interval censoring, specify two Y variables, where one Y variable gives the lower limit and the other Y variable gives the upper limit for each unit. For more information about censoring, see [“Event Plot”](#) on page 38 in the “Life Distribution” chapter.

X Identifies the accelerating factor.

Censor Identifies censored observations. Select the value that identifies right-censored observations from the Censor Code menu beneath the Select Columns list. The Censor column is used only when one Y is entered.

Freq Identifies frequencies or observation counts when there are multiple units. If the value is 0 or a positive integer, then the value represents the frequencies or counts of observations for each row when there are multiple units recorded.

By Identifies a column that creates a report consisting of separate analyses for each level of the variable.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

Relationship Identifies the relationship between the event and the accelerating factor. Table 4.1 defines the model for each relationship.

Table 4.1 Models for Relationship Options

Relationship	Model
Arrhenius Celsius	$\mu = b_0 + b_1 * 11604.5181215503 / (X + 273.15)$
Arrhenius Fahrenheit	$\mu = b_0 + b_1 * 11604.5181215503 / ((X + 459.67) / 1.8)$
Arrhenius Kelvin	$\mu = b_0 + b_1 * 11604.5181215503 / X$
Voltage	$\mu = b_0 + b_1 * \log(X)$
Linear	$\mu = b_0 + b_1 * X$
Log	$\mu = b_0 + b_1 * \log(X)$
Logit	$\mu = b_0 + b_1 * \log(X / (1 - X))$
Reciprocal	$\mu = b_0 + b_1 / X$
Square Root	$\mu = b_0 + b_1 * \text{sqrt}(X)$
Box-Cox	$\mu = b_0 + b_1 * \text{BoxCox}(X)$
Location	means that μ is different for every level of X
Location and Scale	means that μ and σ are both different for every level of X (equivalent to a Life Distribution fit with X as a By variable)
Custom	user-defined μ and σ

If you select Box-Cox, a text edit box appears below the Use Condition option. Use this box to specify a lambda value. The BoxCox(X) transformation for a specified λ is defined as follows:

$$x_i^{(\lambda)} = \begin{cases} \frac{x_i^\lambda - 1}{\lambda} & \text{if } \lambda \neq 0 \\ \ln(x_i) & \text{if } \lambda = 0 \end{cases}$$

If you want to use a Custom relationship for your model, see [“Custom Relationship”](#) on page 131.

Nested Model Tests Appends a nonparametric overlay plot, nested model tests, and a multiple probability plot to the report window.

Use Condition Enables you to enter a value for the explanatory variable, X, of the acceleration factor. You can also set the use condition value after launching the platform using the Set Time Acceleration Use Condition option from the Fit Life by X red triangle menu.

Distribution Specifies one distribution (Weibull, Lognormal, Loglogistic, Fréchet, SEV, Normal, Logistic, LEV, or Exponential distributions) at a time. Lognormal is the default setting.

Select Confidence Interval Method Displays the method for computing confidence intervals for the parameters. The default is Wald, but you can select Likelihood instead. The Wald method is an approximation and runs faster. The Likelihood method provides more precise parameters but takes longer to compute.

Note: The Confidence Interval Method preference enables you to select the Likelihood method as the default confidence interval method. You can change this preference in **Preferences > Platforms > Fit Life by X**.

The Fit Life by X Report

The initial report window includes the following sections:

- [“Summary of Data”](#) on page 113
- [“Scatterplot”](#) on page 113
- [“Nonparametric Overlay”](#) on page 115
- [“Comparisons”](#) on page 116

Distribution, Quantile, Hazard, Density, and Acceleration Factor profilers, along with criteria values under Comparison Criterion can be viewed and compared.

- [“Results”](#) on page 120

Parametric estimates, covariance matrices, nested model tests, and diagnostics can be examined and compared for each of the selected distributions. You can also perform custom estimation and obtain Bayesian estimates for the distribution parameters.

Summary of Data

The Summary of Data section gives the total number of observations, the number of uncensored values, and the number of censored values (right, left, and interval). Figure 4.5 shows the summary data for the Devalt.jmp sample data table.

Figure 4.5 Summary of Data Example

Summary of Data	
Observation Used	165
Uncensored Values	33
Right Censored Values	132

Scatterplot

The Scatterplot of the lifetime event versus the explanatory variable is shown at the top of the report window. For the Devalt.jmp sample data, the Scatterplot shows Hours versus Temp. Table 4.2 indicates how each type of failure is represented on the Scatterplot in the report window. To increase the size of the markers on the graph, right-click the graph, select **Marker Size**, and then select one of the marker sizes listed.

Figure 4.6 Scatterplot of Hours versus Temp

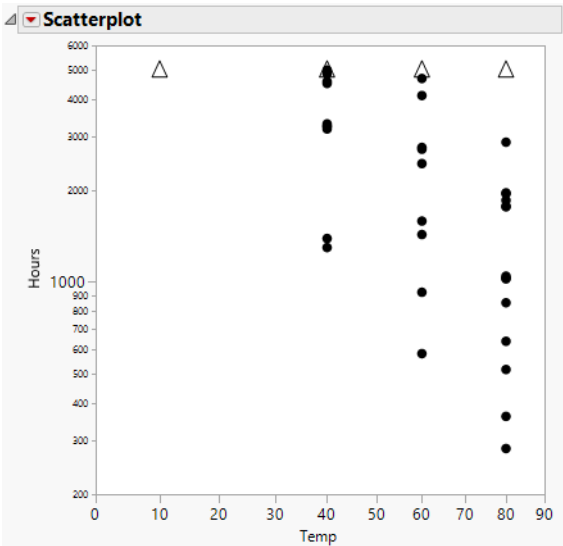


Table 4.2 Scatterplot Representation for Failure and Censored Observations

Event	Scatterplot Representation
failure	dots
right-censoring	upward triangles
left-censoring	downward triangles
interval-censoring	downward triangle on top of an upward triangle, connected by a solid line

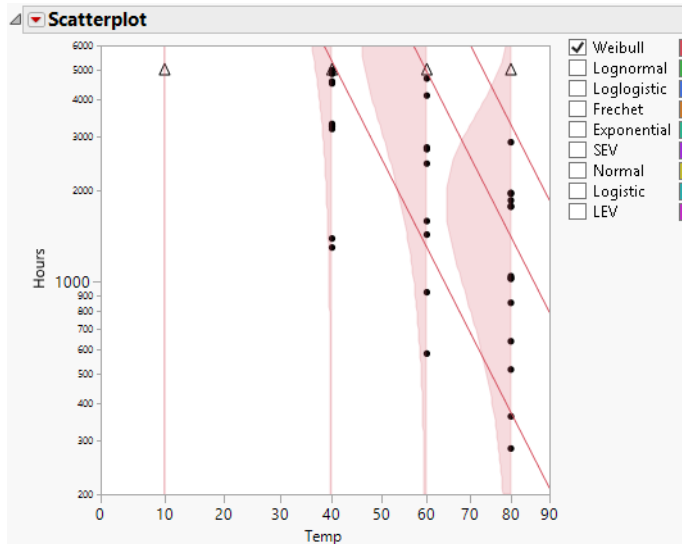
Scatterplot Options

The Scatterplot red triangle menu contains the following options:

- Add Density Curve** Specify the density curve that you want, one at a time, by entering any value within the range of the accelerating factor. You can then select different distributions by selecting the appropriate check box(es) that appear after you add a curve.
- Remove Density Curves** Displays previously entered density curve values. Remove curves by selecting the appropriate check box.
- Show Density Curves** Displays the density curves. If the **Location** or the **Location and Scale** model is fit, or if **Nested Model Tests** is selected in the launch window, then the density curves for all of the given explanatory variable levels are shown. After the curves have been created, the **Show Density Curves** option toggles the curves on and off the plot.
- Add Quantile Lines** Specify the quantile lines that you want, three at a time. You can add more quantiles by continually selecting **Add Quantile Lines**. Default quantile values are 0.1, 0.5, and 0.9. Invalid quantile values, such as missing values, are ignored. If desired, you can enter just one quantile value, leaving the other entries blank.
- Show Quantile Line CI Bands** Shows or hides confidence intervals around the quantile lines.
- Set Level of Quantile Line CI Bands** Specifies the confidence level for the confidence intervals around the quantile lines.
- Remove Quantile Lines** Displays previously entered quantile values. Remove lines by selecting the appropriate check box.
- Transposed Axes** Swaps the horizontal and vertical axes.
- Use Transformation Scale** The default view of the scatterplot incorporates the transformation scale. Select this option to switch between the linear and nonlinear scales for the horizontal axis.

Figure 4.6 shows the initial scatterplot; Figure 4.7 shows the resulting scatterplot with the **Show Density Curves** and **Add Quantile Lines** options selected displaying the curves and the lines for the various Temp levels for the Weibull distribution. You can also view density curves across all the levels of Temp for the other distributions. These distributions can be selected one at a time or can be viewed simultaneously by checking the boxes to the left of the desired distribution name(s).

Figure 4.7 Scatterplot with Density Curve and Quantile Line Options



Nonparametric Overlay

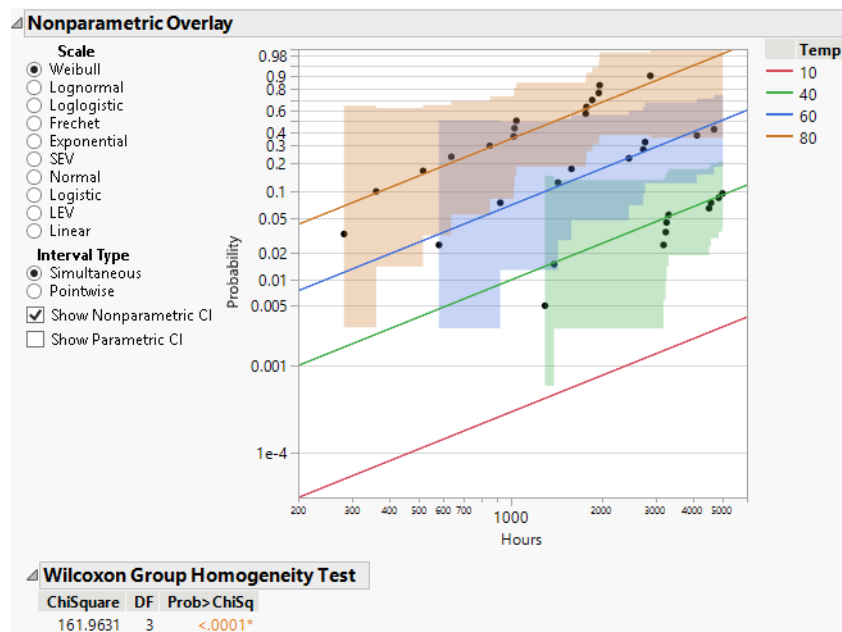
The Nonparametric Overlay plot is displayed after the scatterplot. Differences among groups can readily be detected by examining this plot. You can choose different scales for viewing these differences. You can change the interval type on a Nonparametric fit probability plot between **Simultaneous** and **Pointwise** (results displayed when **Show Nonparametric CI** is selected). You can also select whether to show parametric or nonparametric confidence intervals.

Pointwise estimates show the pointwise 95% confidence bands on the plot while simultaneous confidence intervals show the simultaneous confidence bands for all groups on the plot. Meeker and Escobar (1998, ch. 3) discuss pointwise and simultaneous confidence intervals and the motivation for simultaneous confidence intervals in a lifetime analysis.

Wilcoxon Group Homogeneity Test

For this example, the **Wilcoxon Group Homogeneity Test**, shown in Figure 4.8, indicates that there is a difference among groups. The high chi-square value and low p -value are consistent with the differences seen among the Temp groups in the Nonparametric Overlay plot.

Figure 4.8 Nonparametric Overlay Plot and Wilcoxon Test for Devalt.jmp



Comparisons

The Comparisons report section, shown in Figure 4.9, shows profilers for the selected distributions in the Nonparametric Overlay section, and includes the following tabs:

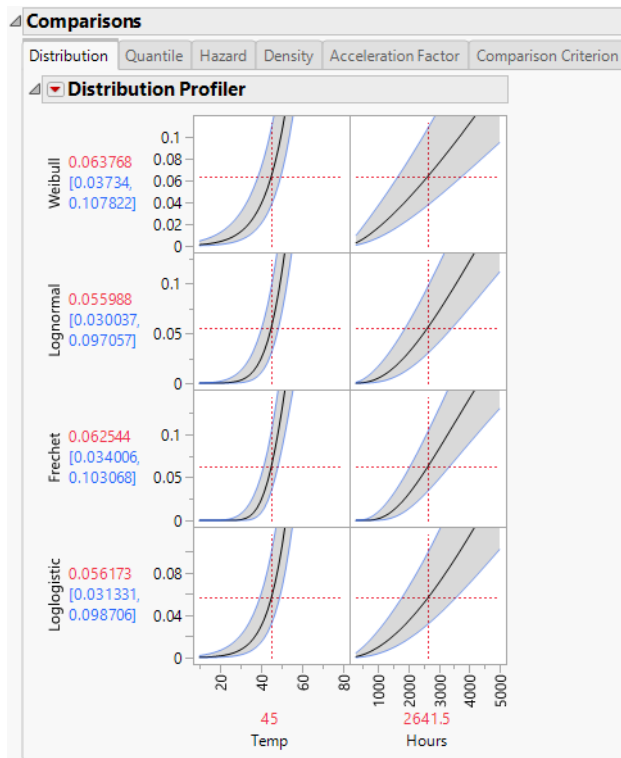
- Distribution
- Quantile
- Hazard
- Density
- Acceleration Factor
- Comparison Criterion

To show a specific profiler, select the appropriate distribution option in the Nonparametric Overlay section.

Profilers

The first five tabs show profilers for the selected distributions. Curves shown in the first four profilers correspond to both the time-to-event and explanatory variables. The Acceleration Factor profiler tab corresponds only to the acceleration factor (explanatory variable). Figure 4.9 shows the Distribution Profiler for the Weibull, Lognormal, Fréchet, and Loglogistic distributions.

Figure 4.9 Distribution Profiler

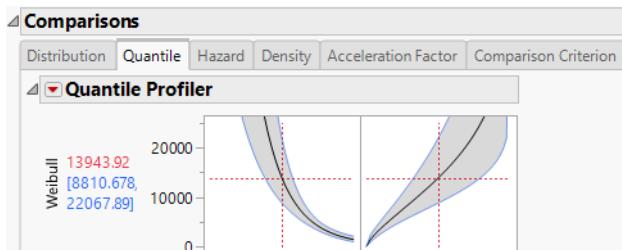


Comparable results appear on the **Quantile**, **Hazard**, and **Density** tabs. The Distribution, Quantile, Hazard, Density, and Acceleration Factor Profilers behave similarly to the Prediction Profiler in other platforms. For example, the vertical lines of Temp and Hours can be dragged to see how each of the distribution values change with temperature and time. For a detailed explanation of the Prediction Profiler, see *Profilers*.

Quantile

You can use the Quantile profiler for extrapolation. Suppose that the data are represented by a Weibull distribution. From viewing the Weibull Acceleration Factor Profiler in Figure 4.11, you see that the acceleration factor at 45 degrees Celsius is 17.42132 for a use condition temperature of 10 degrees Celsius. Select the **Quantile** tab to see the Quantile Profiler for the Weibull distribution. Select and drag the vertical line in the probability plot so that the probability reads 0.5. From viewing Figure 4.10, where the Probability is set to 0.5, you find that the quantile for the failure probability of 0.5 at 45 degrees Celsius is 13943.92 hours. So, at 10 degrees Celsius, you can expect that 50% of the units fail by $13943.92 * 17.42132 = 242921$ hours.

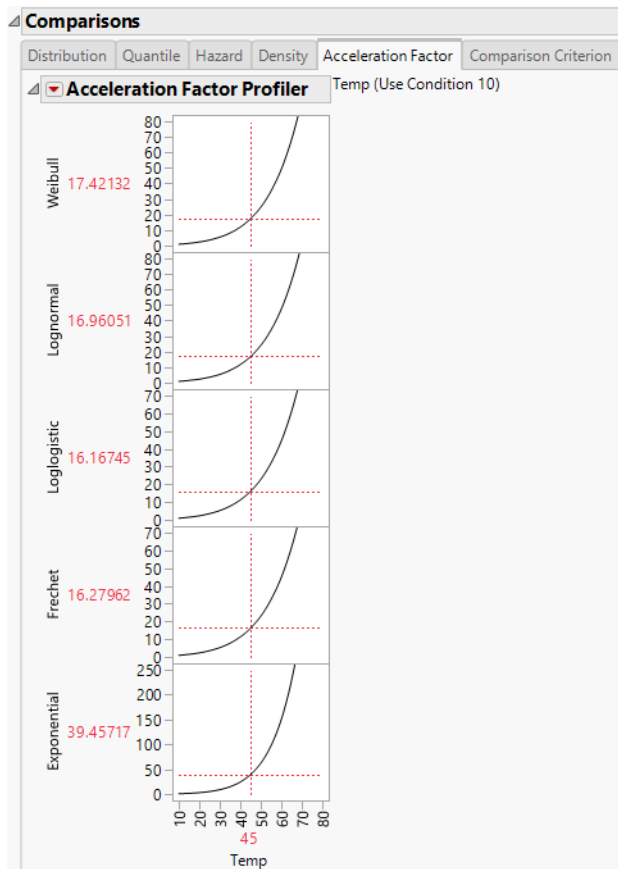
Figure 4.10 Weibull Quantile Profiler for Devalt.jmp



Acceleration Factor

Selecting the **Acceleration Factor** tab shows the Acceleration Factor Profiler for the time-to-event variable for each specified distribution. To produce Figure 4.11, select **Fit All Distributions** from the Fit Life by X red triangle menu. Modify the use condition value for the explanatory variable by selecting **Set Time Acceleration Use Condition** from the Fit Life by X red triangle menu and entering the desired value. Note that the explanatory variable and the use condition value appear beside the profiler title.

Figure 4.11 Acceleration Factor Profiler for Devalt.jmp



The Acceleration Factor Profiler lets you estimate time-to-failure for accelerated test conditions when compared with the use condition value and a parametric distribution assumption. The interpretation of a time-acceleration plot is generally the ratio of the p^{th} quantile of the use condition to the p^{th} quantile of the accelerated test condition. This relation applies only when the distribution is Lognormal, Weibull, Loglogistic, or Fréchet, and the scale parameter is constant for all levels. For more information about the parameterizations of the distributions, see “Distributions” on page 82 in the “Life Distribution” chapter.

Note: The Acceleration Factor Profiler does not appear in the following instances: when the explanatory variable is discrete; the explanatory variable is treated as discrete; a customized formula does not use a unity scale factor; or the distribution is Normal, SEV, Logistic, or LEV.

Comparison Criterion

The **Comparison Criterion** tab shows the -2Loglikelihood, AICc, and BIC criteria for the distributions of interest. Figure 4.12 shows these values for the Weibull, Lognormal, Loglogistic, and Fréchet distributions. Distributions providing better fits to the data are shown at the top of the Comparisons report, sorted by AICc.

Figure 4.12 Comparison Criterion Report Tab

Comparisons				
Distribution	Quantile	Hazard	Density	Acceleration Factor
Comparison Criterion				
Distribution	-2Loglikelihood	AICc	BIC	
Lognormal	643.40556	649.55462	658.72339	
Loglogistic	644.17962	650.32868	659.49745	
Weibull	647.23742	653.38649	662.55526	
Frechet	647.56676	653.71583	662.88460	

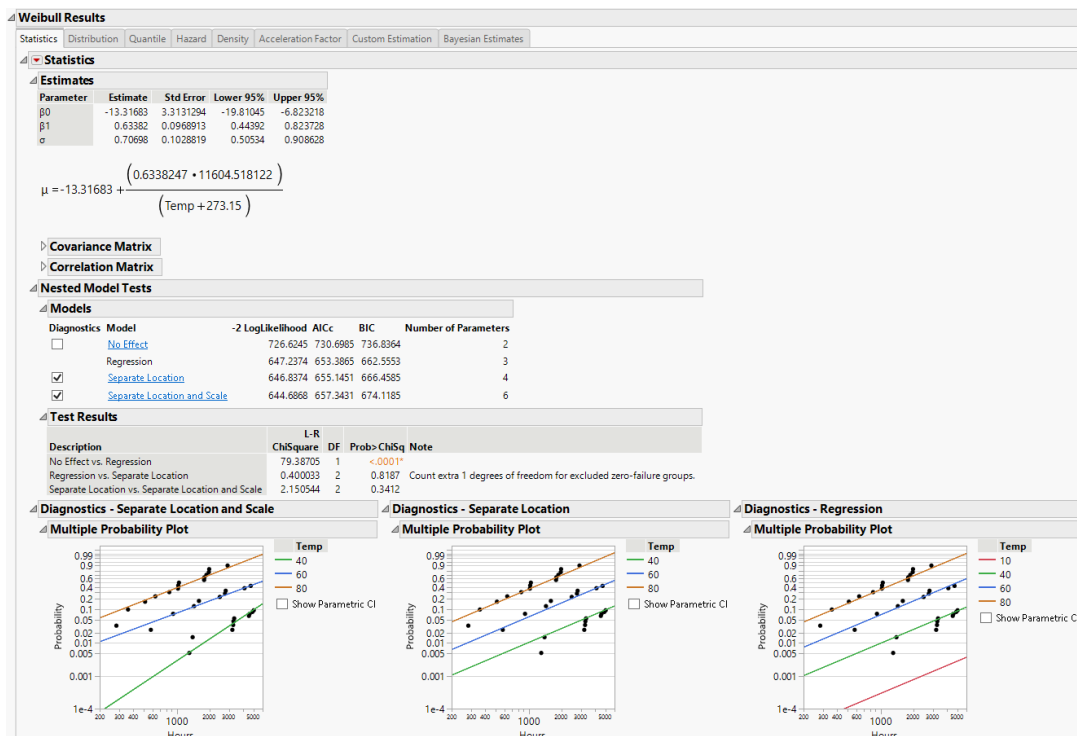
This report suggests that the Lognormal and Loglogistic distributions provide the best fits for the data, because the lowest criteria values are seen for these distributions. For more information about the criteria, see *Fitting Linear Models*.

Results

The Results section of the report window shows more detailed statistics and prediction profilers than those shown in the Comparisons report. Separate result sections are shown for each selected distribution. Figure 4.13 shows a portion of the Weibull results, Nested Model Tests, and Diagnostics plots for Devalt.jmp.

Statistical results, diagnostic plots, and Distribution, Quantile, Hazard, Density, and Acceleration Factor Profilers are included for each of your specified distributions. The Custom Estimation tab lets you estimate specific failure probabilities and quantiles, using both Wald and Profile interval methods. When the Box-Cox Relationship is selected on the platform launch window, the Sensitivity tab appears. This tab shows how the Relative Likelihood and B10 Life change as a function of Box-Cox lambda.

Figure 4.13 Weibull Distribution Nested Model Tests for Devalt.jmp Data



Statistics

For each parametric distribution, there is a Statistics section that shows parameter estimates, a covariance matrix, confidence intervals, summary statistics, and diagnostic plots. You can save probability, quantile, and hazard estimates by selecting any or all of these options from the Statistics red triangle menu for each parametric distribution. The estimates and the corresponding lower and upper confidence limits are saved as columns in your data table.

Nested Model Tests

Nested Model Tests are included, if you selected the option on the platform launch window. The Nested Model Tests include statistics and diagnostic plots for the following models:

Separate Location and Scale Assumes that the location and scale parameters are different for all levels of the explanatory variable. This option is equivalent to fitting the distribution by the levels of the explanatory variable. The Separate Location and Scale model has multiple location parameters and multiple scale parameters (Figure 4.14).

Separate Location Assumes that the location parameters are different, but the scale parameters are the same for all levels of the explanatory variable. The Separate Location model has multiple location parameters and only one scale parameter (Figure 4.15).

Regression The default model shown in the initial Fit Life by X report window (Figure 4.16).

No Effect Assumes that the explanatory variable does not affect the response. This option is equivalent to fitting all of the data values to the selected distribution. The No Effect Model has one location parameter and one scale parameter (Figure 4.17).

Separate Location and Scale, Separate Location, and Regression analyses results are shown by default. Regression parameter estimates and the location parameter formula are shown under the Estimates section, by default. The Diagnostics plots for the No Effect model can be displayed by selecting the check box to the left of No Effect under the Nested Model Tests title.

To see results for each of the models (independently of the other models), click the underlined model of interest (listed under Nested Model Tests) and then uncheck the check boxes for the other models.

If the Nested Model Tests option was not checked in the launch window, then the Separate Location and Scale and Separate Location models are not assessed. In this case, estimates are given for the regression model for each distribution that you select and the Cox-Snell Residual P-P Plot is the only diagnostic plot.

Note: When Separate Location and Scale or Separate Location models are fit for the Weibull distribution, both parameterizations of the Weibull distribution are shown in the Estimates table, as is the case in Figure 4.14 and Figure 4.15. For more information about the Weibull parameterizations, see “Weibull” on page 358 in the “Survival Analysis” chapter.

Diagnostics

The Multiple Probability Plots shown in Figure 4.13 are used to validate the distributional assumption for the different levels of the accelerating variable. If the line for each level does not run through the data points for that level, the distributional assumption might not hold. Side-by-side comparisons of the diagnostic plots provide a visual comparison for the validity of the different models. See Meeker and Escobar (1998, sec. 19.2.2) for a discussion of multiple probability plots. Each multiple probability plot has an option below the legend that enables you to show or hide shaded parametric confidence intervals for each line in the plot.

The Cox-Snell Residual P-P Plots are used to validate the distributional assumption for the data. If the data points deviate far from the diagonal, then the distributional assumption might be violated. The Cox-Snell Residual P-P Plot red triangle menu has an option called **Save Residuals** that enables you to save the residual data to the data table. See Meeker and Escobar (1998, sec. 17.6.1) for a discussion of Cox-Snell residuals.

The Residuals versus Fitted Plots are used to validate the distributional assumption for the different levels of the accelerating variable. The plots show standardized residuals on the vertical axis and the X variable on the horizontal axis. The Residuals vs Fitted red triangle menu has an option called **Save Residuals** that enables you to save the standardized residual data to the data table.

Figure 4.14 Separate Location and Scale Model with the Weibull Distribution for Devalt.jmp Data

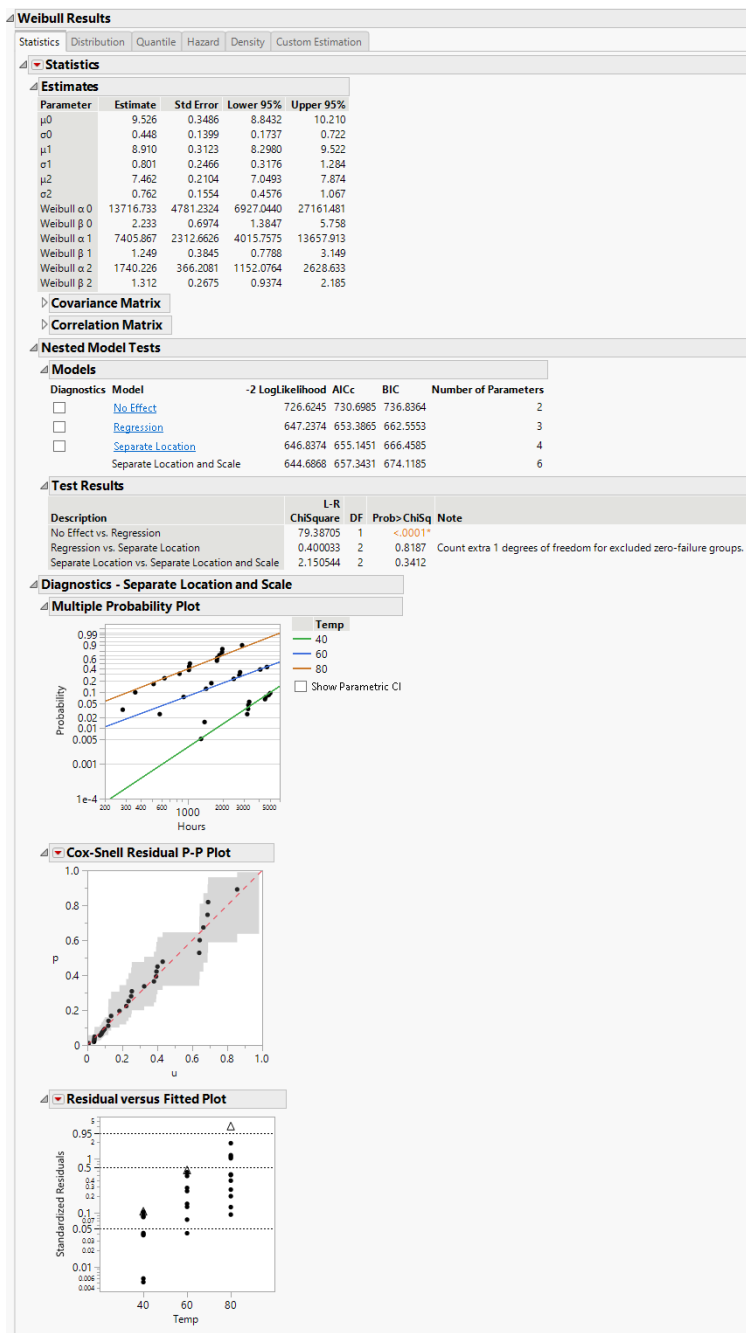


Figure 4.15 Separate Location Model with the Weibull Distribution for Devalt.jmp Data

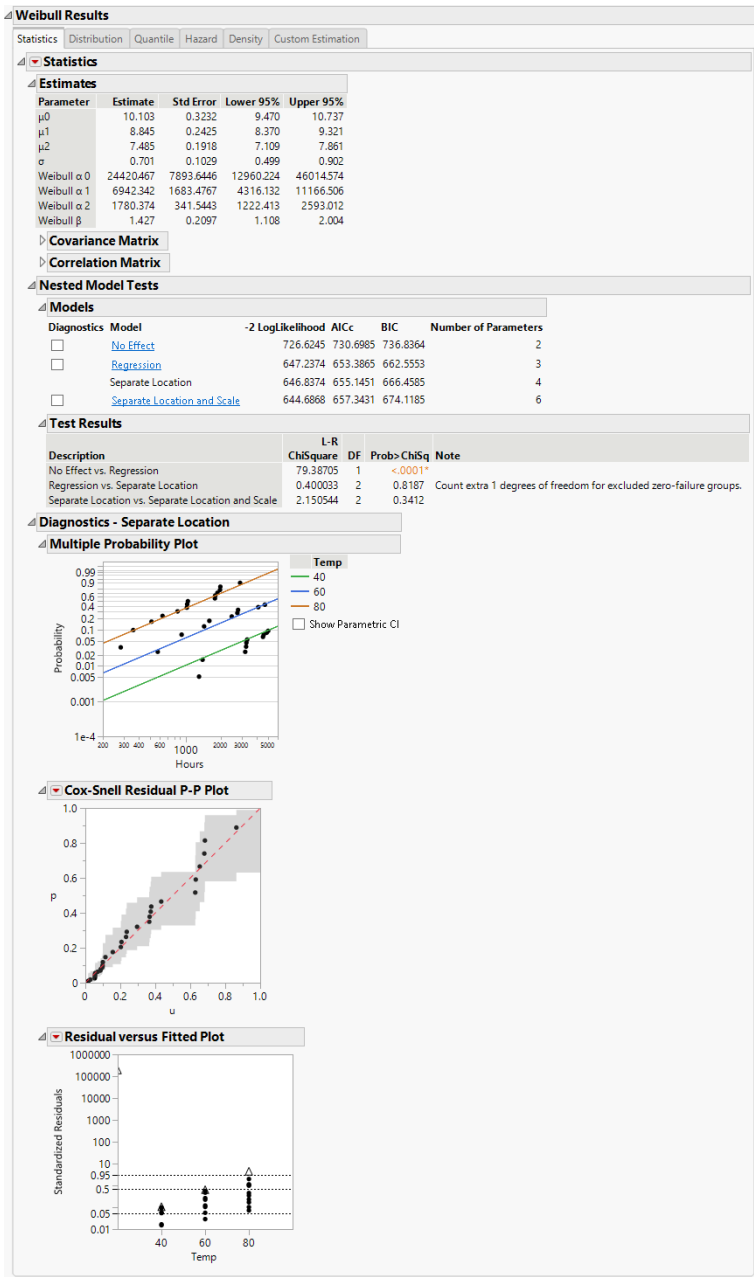


Figure 4.16 Regression Model with the Weibull Distribution for Devalt.jmp Data

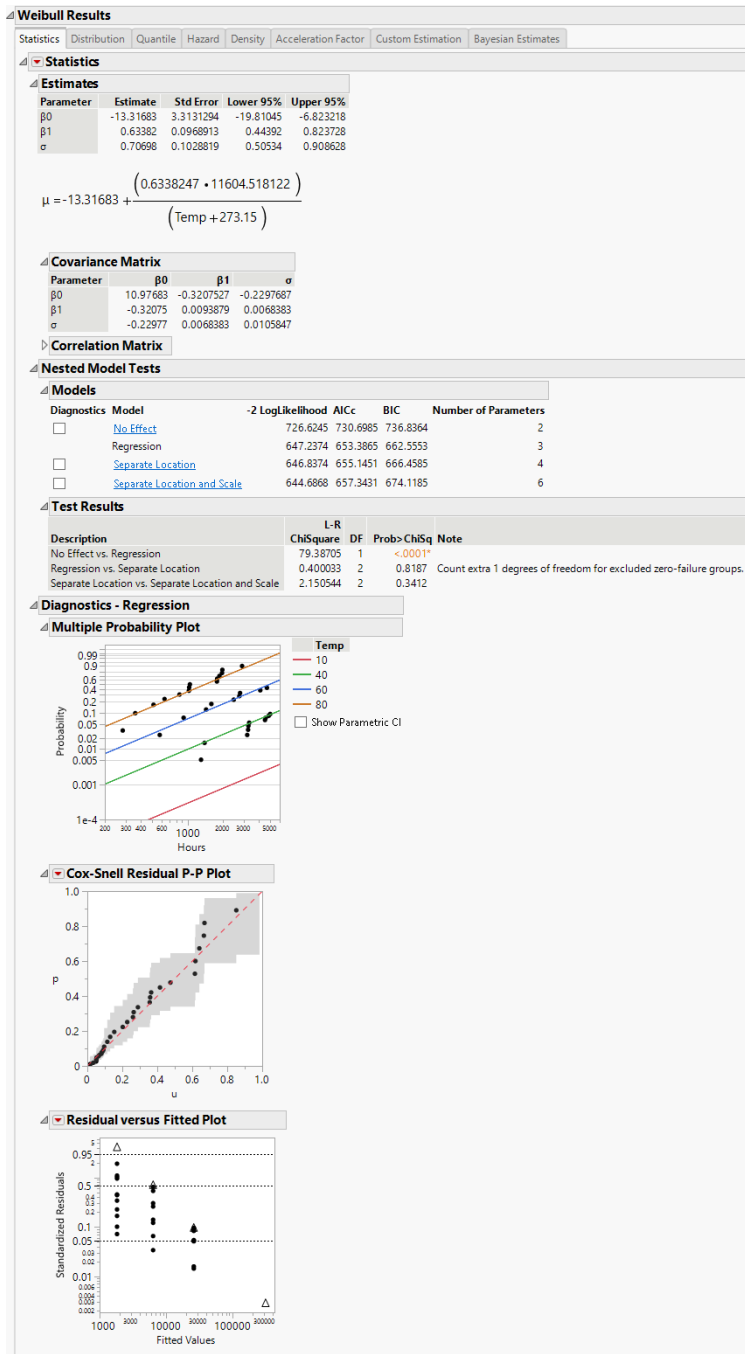
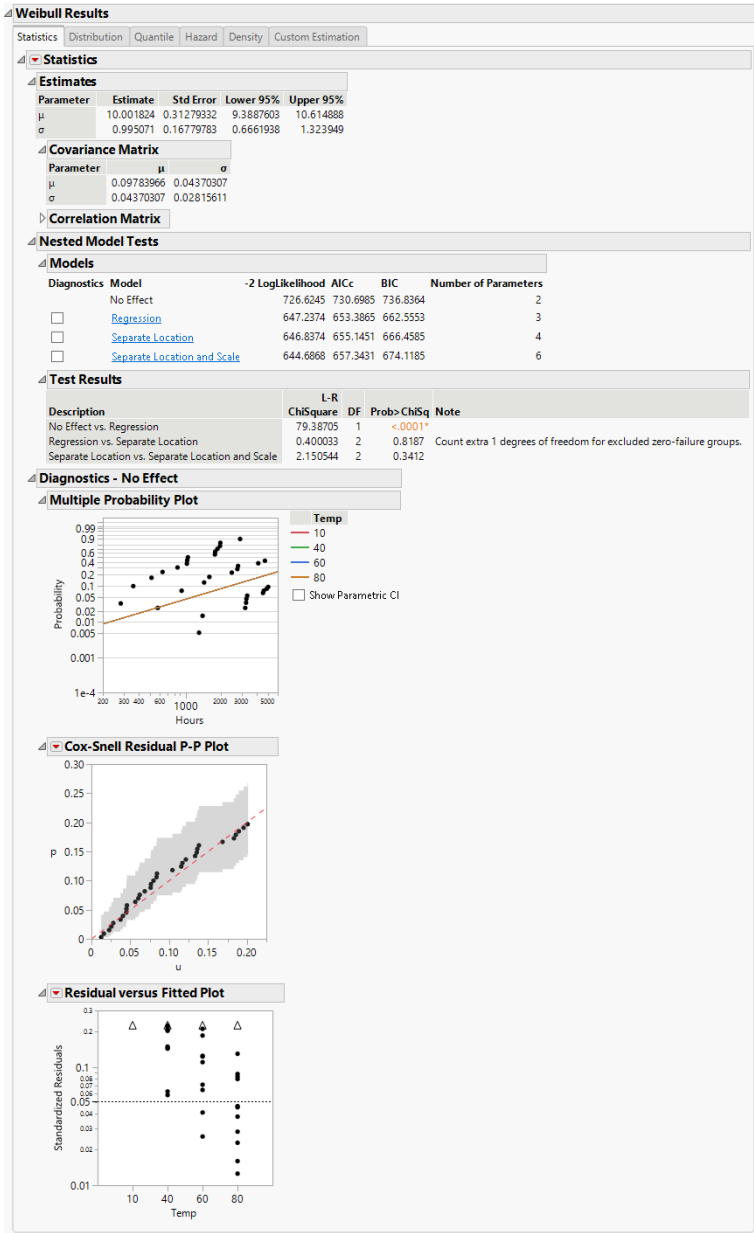


Figure 4.17 No Effect Model with the Weibull Distribution for Devalt.jmp Data

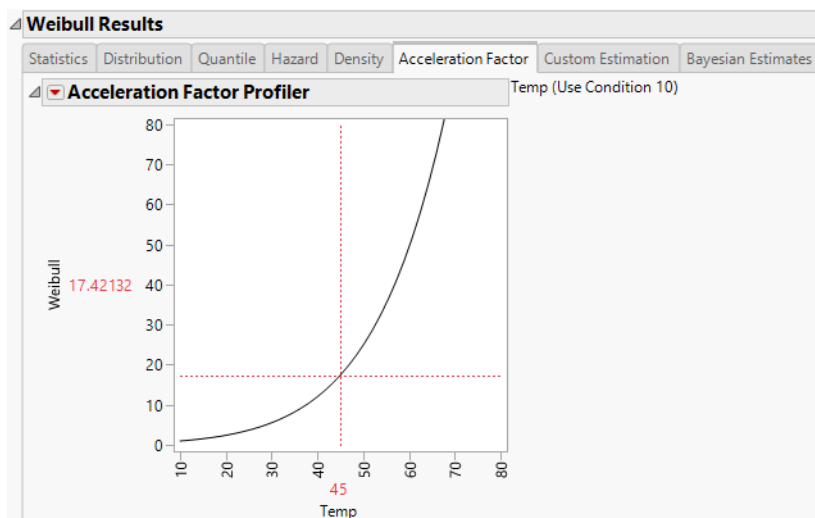


Profilers and Surface Plots

In addition to a statistical summary and diagnostic plots, the Fit Life by X report window also includes profilers and surface plots for each of your specified distributions. To view the Weibull time-accelerating factor and explanatory variable profilers, click the **Distribution** tab under Weibull Results. To see the surface plot, click the disclosure icon to the left of the Weibull title (under the profilers). The profilers and surface plot behave similarly to other platforms. See *Profilers*.

The report window also includes a tab labeled **Acceleration Factor**. Clicking the **Acceleration Factor** tab shows the Acceleration Factor Profiler. This profiler is an enlargement of the Weibull plot shown under the **Acceleration Factor** tab in the Comparisons section of the report window. Figure 4.18 shows the Acceleration Factor Profiler for the Weibull distribution of Devalt.jmp. The use condition level for the explanatory variable can be modified by selecting the **Set Time Acceleration Use Condition** option in the Fit Life by X red triangle menu.

Figure 4.18 Weibull Acceleration Factor Profiler for Devalt.jmp



Custom Estimation

For each parametric distribution, there is a Custom Estimation section that contains two reports: Estimate Quantile and Estimate Probability. For distributions with positive support, the Custom Estimation section also contains an Estimate Mean Remaining Life (MRLF) report.

Estimate Quantile

The Estimate Quantile report contains a calculator that enables you to predict quantiles for specific failure probability values. In the Estimate Quantile calculator, enter a value for Prob and the X variable. Press **Enter** to see the quantile estimates and corresponding confidence intervals. To calculate multiple quantile estimates, click the plus sign, enter another Prob value, another X value, or both, and press **Enter**. Click the minus sign to remove the last entry. If you enter more than one value in either column, the table contains all combinations of Prob and X values.

By default, Wald-based intervals are shown. Click **Likelihood CI** in the Interval Type outline to switch to likelihood-based confidence intervals. The confidence level for these intervals is determined by the **Change Confidence Level** option in the Fit Life by X red triangle menu.

Estimate Probability

The Estimate Probability report contains a calculator that enables you to predict failure and survival probabilities for specific time values. In the Estimate Probability calculator, enter a value for Time and the X variable. Press **Enter** to see the failure probability estimates and corresponding confidence intervals. To calculate multiple failure probability estimates, click the plus sign, enter another Time value, another X value, or both, and press **Enter**. Click the minus sign to remove the last entry. If you enter more than one value in either column, the table contains all combinations of Time and X values.

By default, Wald-based intervals are shown. Click **Likelihood CI** in the Interval Type outline to switch to likelihood-based confidence intervals. The confidence level for these intervals is determined by the **Change Confidence Level** option in the Fit Life by X red triangle menu.

Estimate Mean Remaining Life (MRLF)

The Estimate Mean Remaining Life (MRLF) report contains a calculator that enables you to predict mean remaining life for specific time values. In the Estimate Mean Remaining Life (MRLF) calculator, enter a value for Survival Time and the X variable. Press **Enter** to see the mean remaining life estimates and corresponding Wald-based confidence intervals. The confidence level for these intervals is determined by the **Change Confidence Level** option in the Fit Life by X red triangle menu. To calculate multiple mean remaining life estimates, click the plus sign, enter another Survival Time value, another X value, or both and press **Enter**. Click the minus sign to remove the last entry. If you enter more than one value in either column, the table contains all combinations of Survival Time and X values.

Note: When the survival time equals zero, the calculated mean remaining life (MRLF) value is equivalent to the mean time to failure (MTTF) value.

Bayesian Estimates

For each parametric distribution other than Exponential, there is a Bayesian Estimates section that enables you to obtain Bayesian parameter estimates. The Bayesian Estimates section is not available if the Relationship in the Fit Life by X launch window is Custom, No Effect, Location, or Location and Scale.

Bayesian estimation in the Fit Life by X platform is done using rejection sampling or a Markov Chain Monte Carlo (MCMC) algorithm. More specifically, the platform attempts a basic rejection sampler. If the rejection sampler produces valid results, these results are reported. If the rejection sampler cannot produce valid results, the platform uses a random walk Metropolis-Hastings algorithm and adds a note to the top of the Bayesian Estimation report. See Robert and Casella (2004).

The initial report is a control panel where you can specify prior distributions for the parameters and control aspects of the simulation. To obtain posterior estimates of the parameters, specify the prior distributions and the simulation options, and then click Fit Model.

To specify the prior distributions of the parameters, you must specify information about a quantile of the distribution and the slope β_1 and scale σ parameters. (For the Weibull distribution, you specify the Weibull β rather than σ .) The quantile is defined by two values: the probability of the quantile and the value of the X variable at the specified quantile. The default Probability value is 0.10, but you can specify a value that corresponds to the quantile of interest. Specify information about the range of the prior distribution. For Normal and Lognormal prior distributions, the range is specified in terms of Lower and Upper 99% limits. For Uniform and Log-Uniform prior distributions, the range is specified in terms of the Lower and Upper limits. See Meeker and Escobar (1998). The initial values that are provided are estimates consistent with the maximum likelihood estimates in the Statistics section of the report.

The following options for the simulation appear below the prior distribution specification table:

Number of Monte Carlo Iterations Sets the sample size that will be drawn from the posterior distribution after a burn-in procedure is completed.

Random Seed Sets the initial state of the simulation. By default, it is the clock time. The number should be a positive integer greater than 1. If you specify 1, the current clock time is used.

Bayesian Estimates - Result <N> Report

After you specify prior distributions and the simulation options, click **Fit Model** to perform the simulation. A Bayesian Estimates - Result <N> report is provided for each simulation. This report contains the following headings:

Priors Shows the specifications that you entered in the Bayesian Estimates report to run the simulation. The Prior report also contains the random seed.

Posterior Estimates Shows five marginal statistics that describe the posterior distribution of β_0 , β_1 , σ , and the quantile. The marginal statistics are the median, 0.025 quantile (Lower Bound), 0.975 quantile (Upper Bound), mean, and standard deviation computed from the Monte Carlo samples. If the Weibull distribution is specified, this table contains the posterior estimate of the Weibull β instead of σ .

To compute statistics for other derived variables based on the posterior estimates of the generic parameters, click the Export Monte Carlo Samples link.

Posterior Scatter Plot Shows two scatter plots of values from the Monte Carlo simulation. The scatter plot on the left shows the values of the posterior parameters as they are specified in the Priors report. The scatter plot on the right shows the values of the posterior parameters as they are specified in the Posterior Estimates report.

Profilers Shows two profilers based on samples from the posterior distribution. The values shown in the profilers, at the specified values of the X and Time variables, are calculated using the following steps:

- For each set of sampled parameter values from the posterior distribution, the values of the cumulative distribution function and the quantile function are calculated at the specified values of the X and Time variables.
- The predicted values of the cumulative distribution function and the quantile function are the medians of the calculated values.
- The upper and lower confidence limits are the 0.025 and 0.975 quantiles of the calculated values. The confidence level for these limits is determined by the **Change Confidence Level** option in the Fit Life by X red triangle menu.

Distribution Profiler Shows the parametric cumulative distribution function as a function of the X variable and Time.

Quantile Profiler Shows the parametric quantile function as a function of the X variable and a specified probability.

Bayesian Estimates - Result <N> Options

The Bayesian Estimates - Result <N> red triangle menu contains the following options:

Remove Removes the current Bayesian Estimates report from the Fit Life by X report.

Export Monte Carlo Samples Saves the results of the Monte Carlo simulation to a new data table. You can use this table to compute statistics of the posterior estimates.

Custom Relationship

If you want to use a custom transformation to model the relationship between the lifetime event and the accelerating factor, use the **Custom** option. This option is found in the list under Relationship in the launch window. Enter comma delimited values into the entry fields for the location (μ) and scale (σ) parameters. For the Devalt.jmp sample data, an example entry for μ could be “1, log(:Temp), log(:Temp)^2”, and an entry for σ could be “1, log(:Temp)”, where 1 indicates that an intercept is included in the model. Select the **Use Exponential Link** check box to ensure that the sigma parameter is positive.

Figure 4.19 Custom Relationship Specification in Fit Life by X Launch Window

Models life distributions parameterized by a single regression factor.

Select Columns

6 Columns

- Hours
- Status
- Weight
- Temp
- Censor
- x

Censor Code: 1

Relationship

Custom

☐ Nested Model Tests

Use Condition 10

$\mu =$ 1, log(:Temp), log(:Temp)^2

$\sigma =$ 1, log(:Temp)

☒ Use Exponential Link

Distribution

Weibull

Select Confidence Interval Method

Wald

Cast Selected Columns into Roles

Y, Time to Event	Hours <i>optional numeric</i>
X	Temp
Censor	Censor
Freq	Weight
By	<i>optional</i>

Action

OK

Cancel

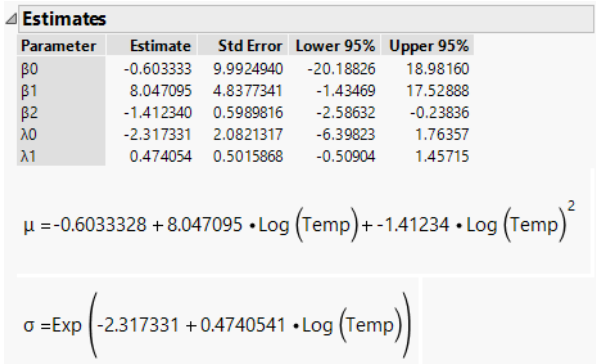
Remove

Recall

Help

After selecting **OK**, location and scale transformations are created and included at the bottom of the Estimates report section.

Figure 4.20 Weibull Estimates and Formulas for Custom Relationship



For an example of how to use a custom transformation, see [“Custom Relationship Example”](#) on page 135. Analysis proceeds similarly to the [“Example of the Fit Life by X Platform”](#) on page 108, where the **Arrhenius Celsius** Relationship was specified.

Fit Life by X Platform Options

The Fit Life by X red triangle menu contains the following options:

- Fit Lognormal** Fits a lognormal distribution to the data.
- Fit Weibull** Fits a Weibull distribution to the data.
- Fit Loglogistic** Fits a loglogistic distribution to the data.
- Fit Fréchet** Fits a Fréchet distribution to the data.
- Fit Exponential** Fits an exponential distribution to the data.
- Fit SEV** Fits a smallest extreme value (SEV) distribution to the data.
- Fit Normal** Fits a normal distribution to the data.
- Fit Logistic** Fits a logistic distribution to the data.
- Fit LEV** Fits a largest extreme value (LEV) distribution to the data.
- Fit All Distributions** Fits all distributions to the data.
- Set Time Acceleration Use Condition** Enables you to enter a use condition value for the explanatory variable of the acceleration factor in a pop-up window.
- Change Confidence Level** Enables you to enter a desired confidence level, for the plots and statistics, in a pop-up window. The default confidence level is 0.95.

Tabbed Report Lets you specify how you want the report window displayed. Two options are available: Tabbed Overall Report and Tabbed Individual Report. Tabbed Individual Report is checked by default. You can select one, both or none.

Show Surface Plot Shows or hides the surface plot for the distribution on and off in the individual distribution results section of the report. The surface plot is shown in the Distribution, Quantile, Hazard, and Density sections for the individual distributions, and it is on by default.

Show Points Shows or hides the data points on and off in the Nonparametric Overlay plot and in the Multiple Probability Plots. The points are shown in the plots by default. If this option is unchecked, step functions are shown instead.

Scatterplot Shows a scatterplot of a lifetime event versus an explanatory variable.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Additional Examples of the Fit Life by X Platform

- [“Capacitor Example”](#)
- [“Custom Relationship Example”](#)

Capacitor Example

This example uses Capacitor ALT.jmp and can be found in the Reliability folder of the sample data. It contains simulated data that gives censored observations for three levels of temperature, for a reliability study. Observations are shown as censored values at temperature levels of 85, 105, and 125 degrees Celsius.

1. Select **Help > Sample Data Library** and open Reliability/Capacitor ALT.jmp.

2. Select **Analyze > Reliability and Survival > Fit Life by X**.
3. Select Hours and click **Y, Time to Event**.
4. Select Temperature and click **X**.
5. Select Censor and click **Censor**.
6. Leave the **Censor Code** as Censored.
7. Select Freq as **Freq**.
8. Keep **Arrhenius Celsius** as the relationship, and keep the **Nested Model Tests** option selected.
9. Select **Weibull** as the distribution.
10. Keep **Wald** as the confidence interval method.

Figure 4.21 Fit Life by X Launch Window

Models life distributions parameterized by a single regression factor.

Select Columns	Cast Selected Columns into Roles	Action										
▼ 4 Columns ▲ Hours ▲ Temperature ▲ Freq ▲ Censor	<table><tr><td>Y, Time to Event</td><td>▲ Hours <i>optional numeric</i></td></tr><tr><td>X</td><td>▲ Temperature</td></tr><tr><td>Censor</td><td>▲ Censor</td></tr><tr><td>Freq</td><td>▲ Freq</td></tr><tr><td>By</td><td><i>optional</i></td></tr></table>	Y, Time to Event	▲ Hours <i>optional numeric</i>	X	▲ Temperature	Censor	▲ Censor	Freq	▲ Freq	By	<i>optional</i>	OK Cancel Remove Recall Help
Y, Time to Event	▲ Hours <i>optional numeric</i>											
X	▲ Temperature											
Censor	▲ Censor											
Freq	▲ Freq											
By	<i>optional</i>											

Censor Code:

Censored ▼

Relationship

Arrhenius Celsius ▼

☒ Nested Model Tests

Use Condition

85

Distribution

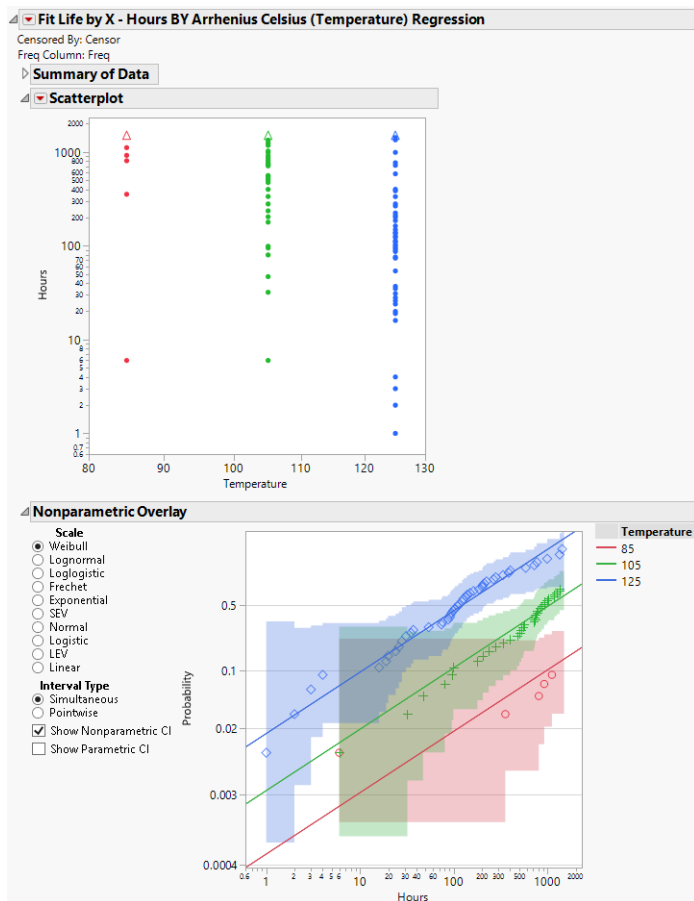
Weibull ▼

Select Confidence Interval Method

Wald ▼

11. Click **OK**.

Figure 4.22 Fit Life by X Report Window for Capacitor ALT.jmp Data



The report window shows summary data, diagnostic plots, comparison data and results, including detailed statistics and prediction profilers. Separate result sections are shown for each selected distribution. Distribution, Quantile, Hazard, Density, and Acceleration Factor Profilers are included for each of the specified distributions.

Custom Relationship Example

To create a quadratic model with $\text{Log}(\text{Temp})$ for the Weibull location parameter and a log-linear model with $\text{Log}(\text{Temp})$ for the Weibull scale parameter, follow these steps:

1. Select **Help > Sample Data Library** and open Reliability/Devalt.jmp.
2. Select **Analyze > Reliability and Survival > Fit Life by X**.
3. Select Hours and click **Y, Time to Event**.
4. Select Temp and click **X**.

- 5. Select **Censor** and click **Censor**.
- 6. Select **Weight** and click **Freq**.
- 7. Select **Custom** as the Relationship from the list.
- 8. In the entry field for μ , enter 1, log(:Temp), log(:Temp)^2.
(The 1 indicates that an intercept is included in the model.)
- 9. In the entry field for σ , enter 1, log(:Temp).
- 10. Select the check box for **Use Exponential Link**.
- 11. Deselect the check box for **Nested Model Tests**.
- 12. In the entry field for Use Condition, enter 10.
- 13. Select **Weibull** as the Distribution.

Figure 4.23 shows the completed launch window using the **Custom** option.

Note: The Nested Model Tests check box is not checked for non-constant scale models. Nested Model test results are not supported for this option.

- 14. Click **OK**.

Figure 4.23 Custom Relationship Specification in Fit Life by X Launch Window

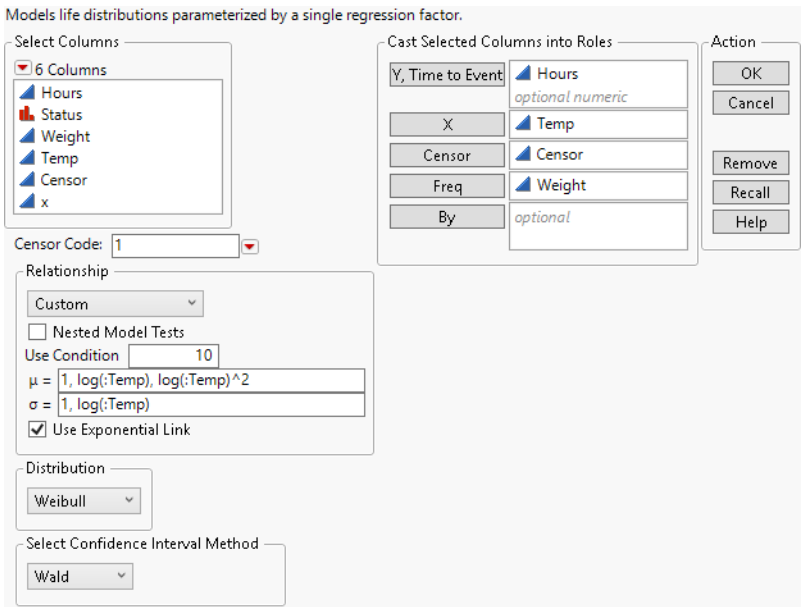


Figure 4.24 shows the location and scale transformations, which are subsequently created and included at the bottom of the Estimates report section.

Analysis proceeds similarly to the “[Example of the Fit Life by X Platform](#)” on page 108, where the **Arrhenius Celsius** Relationship was specified.

Figure 4.24 Weibull Estimates and Formulas for Custom Relationship

Estimates				
Parameter	Estimate	Std Error	Lower 95%	Upper 95%
β_0	-0.603333	9.9924940	-20.18826	18.98160
β_1	8.047095	4.8377341	-1.43469	17.52888
β_2	-1.412340	0.5989816	-2.58632	-0.23836
λ_0	-2.317331	2.0821317	-6.39823	1.76357
λ_1	0.474054	0.5015868	-0.50904	1.45715

$$\mu = -0.6033328 + 8.047095 \cdot \text{Log}(\text{Temp}) + -1.41234 \cdot \text{Log}(\text{Temp})^2$$

$$\sigma = \text{Exp} \left(-2.317331 + 0.4740541 \cdot \text{Log}(\text{Temp}) \right)$$

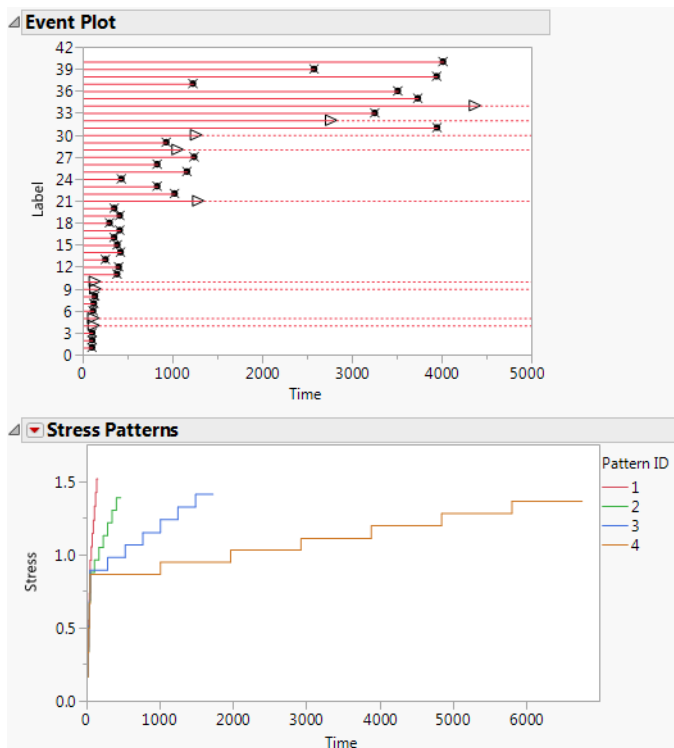
Chapter 5

Cumulative Damage

Model Product Deterioration with Variable Stress over Time

Cumulative damage models, which include step-stress models, enable you to analyze an accelerated life test where the stress levels might be changed over time. The stress can be applied by many different forces: load, temperature, pressure. A typical cumulative damage experiment consists of multiple test units. Each unit has an initial stress level, and the stress level can be changed throughout the experiment. The response is the failure time or time-to-event. The platform plots the failure events and enables you to fit multiple distributions to your data.

Figure 5.1 Example of Cumulative Damage Report



Contents

Overview of the Cumulative Damage Platform..... 141

Example of the Cumulative Damage Platform..... 141

Launch the Cumulative Damage Platform 145

 Time to Event..... 146

 Stress Pattern 147

The Cumulative Damage Report 150

 Event Plot 150

 Stress Patterns Report 150

 Model List..... 151

 Model Results 151

Cumulative Damage Platform Options 152

 Stress Patterns Options..... 152

Additional Example of the Cumulative Damage Platform..... 154

 Example of Simulating New Data 154

Overview of the Cumulative Damage Platform

A cumulative damage experiment, also called a *varying-stress experiment*, is an accelerated life test where the stress levels can change over time. The stress can be applied by many different forces: load, temperature, pressure. A typical cumulative damage experiment consists of multiple test units. Each unit has an initial stress level, and the stress level can be changed throughout the experiment.

The most common cumulative damage experiment is a *step-stress experiment*. A step-stress experiment uses multiple units with varying levels of stress applied. Stress can be applied using factors such as temperature, pressure, or voltage. For each unit, there is an initial stress level. At specified time points, the stress levels are adjusted based on different patterns of stress levels. Between stress level changes, the stress level remains constant.

The Cumulative Damage platform also includes three other varying-stress pattern models:

- In a ramp-stress experiment, the stress levels start at an initial value and then increase linearly over time at a specified slope.
- In a sinusoid-stress experiment, the stress levels fluctuate in a periodic fashion that is defined by a sine wave.
- In a piecewise ramp-stress experiment, the stress levels are defined at specified time points similar to the step-stress case. However, the stress level is not required to stay constant between time points. Rather, it changes linearly from a starting stress level to an ending stress level between time points. If a pair of starting and ending stress levels are equal, the interval is equivalent to a step-stress interval.

For more information about varying-stress and step-stress models, see Nelson (2004, ch. 10).

Example of the Cumulative Damage Platform

A step-stress experiment is conducted on 40 test units at varying stress conditions. Your goal is to estimate the probability of failure at 10,000 time units given a stress level of 0.75. These data tables are based on data from Nelson (2004, ch. 10).

The Reliability/CD Step Stress Pattern.jmp data table contains a column called Pattern ID that identifies four different stress patterns. The stress level at a particular step is the ratio of Voltage to Thickness. (Note that these two columns are hidden.) Thickness is held constant for each stress pattern. However, Voltage is set to different levels and increases within each pattern.

The Reliability/CD Step Stress.jmp data table contains the time to failure data.

1. Select **Help > Sample Data Library** and open Reliability/CD Step Stress.jmp and Reliability/CD Step Stress Pattern.jmp.

The CD Step Stress table contains failure time data:

- The Time column gives the failure times.
- The Pattern ID column identifies the stress pattern.
- The Censor column indicates whether the failure time is exact or censored.

Each row of the table corresponds to one test unit.

The CD Step Stress Pattern table contains the four stress patterns (identified as 1 through 4). The levels of the stress factor, **Stress**, are varied within each value of the **Pattern ID** column. The **Duration** column represents how many time units a particular level of the stress factor lasted.

2. Select **Analyze > Reliability and Survival > Cumulative Damage**.

The launch window has two sections: one for the failure time data (Time-to-Event) and one for the stress pattern data (Stress Pattern).

3. Click **Select Table** in the Time-to-Event panel.

A Time-to-Event Data Table window appears, which prompts you to specify the data table for the failure time data.

4. Select CD Step Stress and click **OK**.

The columns from this table now populate the **Select Columns** list in the Time-to-Event panel.

5. Select Time and click **Time to Event**.

6. Select Censor and click **Censor**.

7. Select Pattern ID for **Pattern ID**.

8. Click **Select Table** in the Stress Pattern panel.

9. Select CD Step Stress Pattern and click **OK**.

The columns from this table now populate the **Select Columns** list in the Stress Pattern panel.

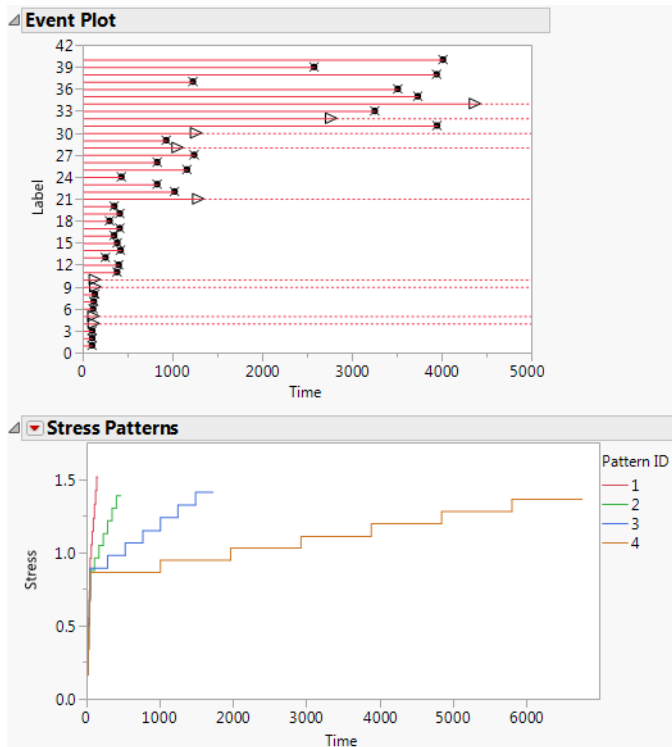
10. Select Duration and click **Stress Duration**.

11. Select Stress and click **Stress**.

12. Select Pattern ID and click **Pattern ID**.

13. Click **OK**.

Figure 5.2 Event Plot and Stress Patterns Plot



The initial report contains the Event Plot and a plot of the defined stress patterns. All four stress patterns increase the stress level quickly over the first 40 time units, after which they increase at much different rates.

14. Click the Cumulative Damage red triangle and select **Fit All**.

Figure 5.3 Model List Report

Model List				
Distribution	-2 Log Likelihood	Nparm	AICc	BIC
Exponential	398.48312	2	402.80745	405.86088
Weibull	398.4822	3	405.14887	409.54884
Loglogistic	400.93264	3	407.59931	411.99928
Lognormal	402.3488	3	409.01547	413.41544
Frechet	410.92577	3	417.59243	421.99241

From the Model List report, you determine that the best fitting distribution is the Exponential distribution.

15. In the Results report, scroll to the Exponential report.
16. In the Distribution Profiler report, set the current value of Stress to 0.75.
17. Set the current value of Time to 10000.

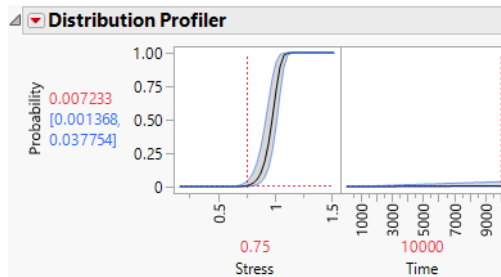
Figure 5.4 Distribution Profiler for Exponential Distribution at Specified Settings

Figure 5.4 shows that the predicted probability of failure for a test unit under constant stress of 0.75 at 10000 time units is 0.007233, with a 95% confidence interval of 0.001368 to 0.037754.

Launch the Cumulative Damage Platform

Launch the Cumulative Damage platform by selecting **Analyze > Reliability and Survival > Cumulative Damage**. You must supply the Cumulative Damage platform with two data tables as input. The first data table is time-to-event data for each unit under test. The second data table defines the stress patterns used for each unit.

Figure 5.5 The Cumulative Damage Launch Window

The screenshot shows the 'Cumulative Damage Launch Window' with the 'Step Stress' tab selected. The window is organized into several sections:

- Time-to-Event:** Includes a 'Select Table' button, a 'Select Columns' list with a red triangle menu (showing 'optional'), and a 'Cast Selected Columns into Roles' table.

Column	Role
Time to Event	required numeric
Censor	optional numeric
Freq	optional numeric
Pattern ID	required
- Censor Code:** A dropdown menu.
- Relationship:** A dropdown menu set to 'Arrhenius Celsius'.
- Distribution:** A dropdown menu set to 'Weibull'.
- Stress Pattern:** Includes a 'Select Table' button, a 'Select Columns' list with a red triangle menu (showing 'optional'), and a 'Cast Selected Columns into Roles' table.

Column	Role
Stress Duration	required numeric
Stress	required numeric
Pattern ID	required
- Pattern Continuation:** Radio buttons for 'Terminate' (selected), 'Extend', and 'Repeat'.
- Action:** Buttons for 'OK', 'Cancel', 'Remove', 'Recall', and 'Help'.

The launch window includes a separate tab for each step-stress data format. For more information about the various stress patterns, see [“Stress Pattern”](#) on page 147. For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

Each of the step-stress format tabs contains two panels for specifying variables for the model:

- The Time-to-Event panel is common to all of the step-stress formats. This panel is similar to the Fit Life by X platform launch window.
- The Stress Pattern panel is used to describe the type of stress. Consequently, its format depends on the selected step-stress tab.

Time to Event

The Time to Event panel in the Cumulative Damage launch window contains the following options:

Time to Event Identifies the time to event (such as the time to failure) or time to censoring. For interval censoring, your data table should contain two columns, where one gives the lower bound and the other gives the upper bound of the failure time for each unit. Enter the two censoring columns as Time to Event. For more information about censoring, see [“Event Plot”](#) on page 38 in the “Life Distribution” chapter.

Censor Identifies right-censored observations. Select the value that identifies right-censored observations from the Censor Code menu beneath the Select Columns list. The Censor column is used only when one Y is entered.

Freq Identifies frequencies or observation counts when there are multiple units. If the value is 0 or a positive integer, then the value represents the frequencies or counts of observations with the given row’s settings.

Pattern ID Contains values that specify the stress pattern in the Stress Pattern data table that was used for the given row.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

Relationship Identifies the relationship between the event and the stress factor. Table 5.1 defines the model for each relationship.

Table 5.1 Models for Relationship Options

Relationship	Model
Arrhenius Celsius	$\mu = b_0 + b_1 * 11604.5181215503 / (X + 273.15)$

Table 5.1 Models for Relationship Options (*Continued*)

Relationship	Model
Arrhenius Fahrenheit	$\mu = b_0 + b_1 * 11604.5181215503 / ((X + 459.67) / 1.8)$
Arrhenius Kelvin	$\mu = b_0 + b_1 * 11604.5181215503 / X$
Inverse Power (default)	$\mu = b_0 + b_1 * \log(X)$
Linear	$\mu = b_0 + b_1 * X$
Log	$\mu = b_0 + b_1 * \log(X)$
Logit	$\mu = b_0 + b_1 * \log(X / (1 - X))$
Reciprocal	$\mu = b_0 + b_1 / X$
Square Root	$\mu = b_0 + b_1 * \text{sqrt}(X)$
Box-Cox	$\mu = b_0 + b_1 * \text{BoxCox}(X)$
Custom	user-defined (Available only in the Step Stress panel.)

The BoxCox(X) transformation is defined as follows:

$$x_i^{(\lambda)} = \begin{cases} \frac{x_i^\lambda - 1}{\lambda} & \text{if } \lambda \neq 0 \\ \ln(x_i) & \text{if } \lambda = 0 \end{cases}$$

If you select **Custom**, additional controls appear that require you to define the Custom transformation that models the relationship between the lifetime event and the stress factor.

If you want to use a Custom relationship for your model, see [“Custom Relationship”](#) on page 131 in the “Fit Life by X” chapter.

Distribution Specifies an initial time-to-failure distribution. Select from Weibull, Lognormal, Loglogistic, Fréchet, or Exponential. Weibull is the default setting. For more information about the distributions, see [“Distributions”](#) on page 82 in the “Life Distribution” chapter.

Stress Pattern

Specify the stress patterns used in the experiment using the second panel.

Step Stress Pattern

The Step Stress pattern has stress levels that are changed at arbitrary time points. The duration of each stress step and associated stress level must be specified in ascending time order.

The Stress Pattern panel in the Step Stress tab contains the following options:

Stress Duration The column that contains the length in time units of each stress step.

Stress The column that contains the level of the stress setting.

Pattern ID The column that contains a unique identifier for the stress pattern. This column is used to match stress patterns in the Stress Pattern data table with the observations in the Time-to-Event data table.

Pattern Continuation Specifies how to handle failures that occur after the final time period in the defined stress pattern. This panel contains the following options:

Terminate A failure that occurs at a time beyond the final time period in the defined stress pattern produces an error, and the model is not fit.

Extend A failure that occurs at a time beyond the final time period in the defined stress pattern assumes the same stress level as the level in the final time period.

Repeat A failure that occurs at a time beyond the final time period in the defined stress pattern assumes the same stress level as if the stress pattern were being repeated. For example, if a failure occurs 10 time units after the final time period in the defined stress pattern, then the stress level at that failure time is set to the stress level at 10 time units after the beginning of the defined stress pattern.

Note: The default Pattern Continuation setting for the Step Stress Pattern is Terminate.

Ramp Stress Pattern

The Ramp Stress pattern defines stress as a linear function of time. Each pattern is defined by an intercept (the stress level at time zero) and a slope (the increase in the stress level for every one time unit). Each pattern is described in a single row in the stress pattern data table.

Intercept The column that contains the intercept for each pattern.

Slope The column that contains the slope for each pattern.

Pattern ID The column that contains a unique identifier for the stress pattern. This column is used to match stress patterns in the Stress Pattern data table with the observations in the Time-to-Event data table.

Sinusoid Stress Pattern

The Sinusoid Stress pattern defines stress as a periodic function. The pattern is defined by a level, an amplitude, a period, and a phase. Each pattern is described in a single row in the stress pattern data table. The pattern is defined as follows:

$$S(t) = \text{level} + \text{amplitude} * \sin(\text{phase} + (2 * \pi * t) / \text{period})$$

Level The column that contains the level for each pattern.

Amplitude The column that contains the amplitude for each pattern.

Period The column that contains the period for each pattern.

Phase The column that contains the phase for each pattern.

Pattern ID The column that contains a unique identifier for the stress pattern. This column is used to match stress patterns in the Stress Pattern data table with the observations in the Time-to-Event data table.

Piecewise Ramp Stress Pattern

The Piecewise Ramp Stress pattern defines stress as a piecewise linear function of time. The line segments for the stress level over time can be disjoint or continuous. Line segments can also be flat, so that step stress and ramp stress can be combined. The line segments are defined in the stress pattern data table by the time duration of the segment and the start and end levels of the stress setting.

Stress Duration The column that contains the length in time units of each stress step.

Stress Ramp The two columns that contain the stress levels at the start and end of the step.

Pattern ID The column that contains a unique identifier for the stress pattern. This column is used to match stress patterns in the Stress Pattern data table with the observations in the Time-to-Event data table.

Pattern Continuation The Pattern Continuation panel enables you to specify the stress levels that occur after the final time period in the defined stress pattern. This panel contains the following options:

Terminate A failure that occurs at a time beyond the final time period in the defined stress pattern produces an error, and the model is not fit.

Extend A failure that occurs at a time beyond the final time period in the defined stress pattern assumes the same stress level as the level in the final time period.

Repeat A failure that occurs at a time beyond the final time period in the defined stress pattern assumes the same stress level as if the stress pattern were being repeated. For example, if a failure occurs 10 time units after the final time period in the defined stress pattern, then the stress level at that failure time is set to the stress level at 10 time units after the beginning of the defined stress pattern.

Note: The default Pattern Continuation setting for the Piecewise Ramp Stress Pattern is Terminate.

The Cumulative Damage Report

After you click OK, the Cumulative Damage report window appears. By default, the Cumulative Damage report contains an event plot, a stress patterns report, a model list, and model results.

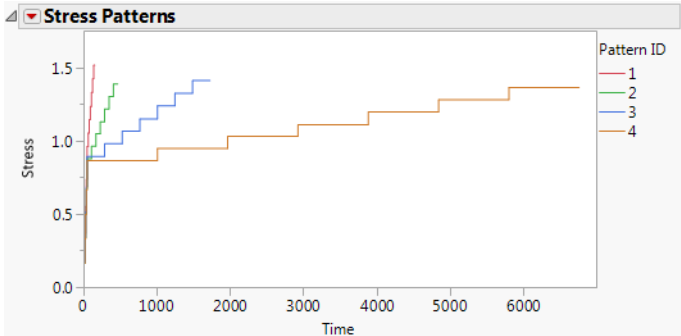
Event Plot

The Event Plot in Cumulative Damage displays time to failure or censoring. See [“Event Plot”](#) on page 38 in the “Life Distribution” chapter.

Stress Patterns Report

The Stress Patterns report shows a plot of the stress level over time for each of the stress pattern IDs. The Simulate option in this report enables you to simulate new data. Figure 5.6 shows an example of the Stress Patterns report plot. See [“Stress Patterns Options”](#) on page 152.

Figure 5.6 Stress Patterns Report



Model List

The Model List report provides the -2Loglikelihood, number of parameters, AICc, and BIC statistics for each fitted distribution. Smaller values of each of these statistics (other than number of parameters) indicate a better fit. For more information about these statistics, see *Fitting Linear Models*.

Model Results

The Results report contains a separate report for each fitted distribution. Each report contains the following:

Parameter Estimates Shows the estimates, standard errors, and Wald-based 95% confidence intervals.

The fitted equation appears below the Parameter Estimates table. This equation takes into account the fitted parameter estimates and the relationship specified in the launch window.

Distribution Profiler Shows cumulative failure probability as a function of time.

Quantile Profiler Shows failure time as a function of cumulative probability.

Hazard Profiler Shows the hazard rate as a function of time.

Density Profiler Shows the density function for the distribution.

Probability Plot of Standardized Residuals Shows a plot of standardized residuals for the model on a probability scale axis.

The Probability Plot of Standardized Residuals red triangle menu has a Save Residuals option that saves the standardized residuals to three new columns in the failure time data table. The three columns are Left Residuals, Right Residuals, and Residual Weight.

Cox-Snell Residual P-P Plot Shows a residual plot that is used to validate the distributional assumption for the data.

The Cox-Snell Residual P-P Plot red triangle menu has a Save Residuals option that saves the Cox-Snell residuals to three new columns in the failure time data table. The three columns are Left Residuals, Right Residuals, and Residual Weight. See Meeker and Escobar (1998, sec. 17.6.1) for a discussion of Cox-Snell residuals.

Cumulative Damage Platform Options

The Cumulative Damage red triangle menu contains the following options:

Fit All Fits the distributions that were not selected in the launch window. The available distributions are listed under “[Distribution](#)” on page 147.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

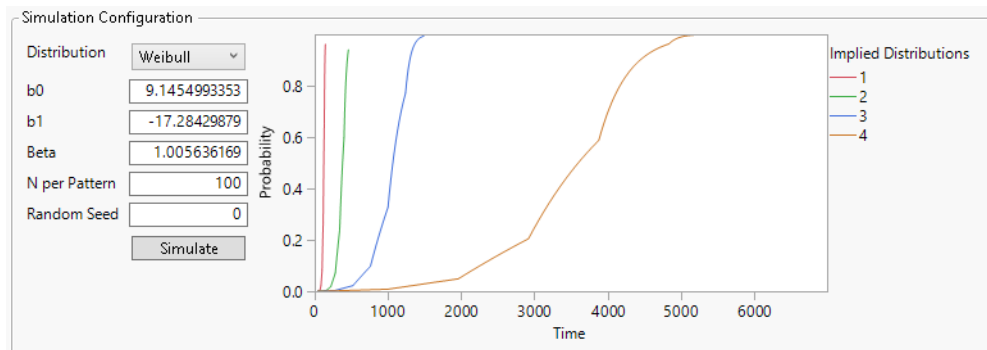
Stress Patterns Options

The Stress Pattern red triangle menu contains the Simulate option, which shows or hides the Simulation Configuration panel.

The Simulation Control Panel

Figure 5.7 shows the Simulation Configuration Panel. The initial values for Distribution and the parameter settings are determined by the parameter estimates for the distribution specified in the Cumulative Damage launch window. The graph shows the estimated failure distribution functions over time for the stress patterns defined in the Stress Pattern data table.

Figure 5.7 Simulation Configuration Panel



The Simulation Configuration panel enables you to simulate new failure time data based on a distribution and stress pattern. The stress pattern defined in the Stress Pattern data table used in the launch of the platform is also used for the simulation. This panel contains the following options:

Distribution The distribution to be used for the simulation. The available distributions are the same as in the Cumulative Damage launch window. For more information about the distributions, see “Distributions” on page 82 in the “Life Distribution” chapter.

b0 The intercept for the location parameter of the distribution.

b1 The slope for the location parameter of the distribution.

lambda (Available only for the Box-Cox relationship.) The lambda value for the Box-Cox relationship.

b2, s0, s1, and so on (Available only for the Custom relationship.) Other parameters that are defined in the Custom relationship in the Cumulative Damage launch window.

Beta (Available only for the Weibull distribution.) The Beta parameter of the Weibull distribution.

sigma (Available only for the Lognormal, Loglogistic, and Fréchet distributions.) The sigma parameter of the distribution.

N per Pattern The number of points generated in the simulation for each stress pattern.

Random Seed (Optional) A nonzero random seed that ensures the reproducibility of simulation results.

Termination (Not available when the specified Pattern Continuation in the Cumulative Damage launch window is Terminate.) A time beyond which surviving test units are censored.

Simulate

The plot in the Simulation Configuration panel shows the implied distributions for each of the stress patterns over time. Click the Simulate button to generate a new JMP data table that contains the results of the simulation.

Additional Example of the Cumulative Damage Platform

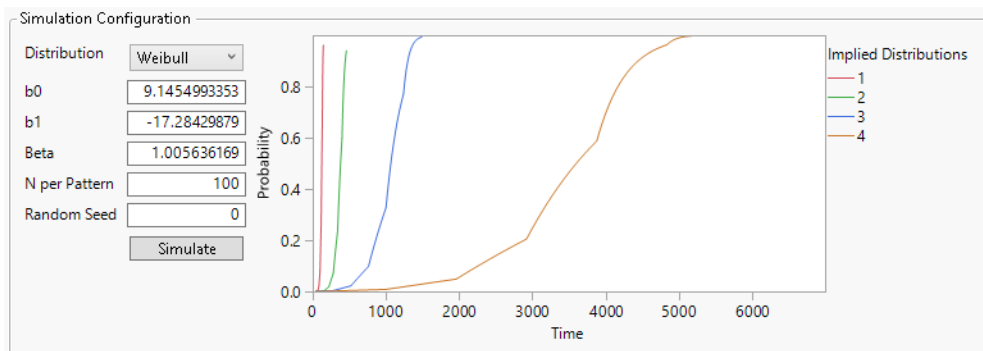
This section contains additional examples using the Cumulative Damage platform.

Example of Simulating New Data

This example illustrates using the Simulation Configuration panel in the Cumulative Damage report window to generate new step stress data. This example uses the same data as the example in [“Example of the Cumulative Damage Platform”](#) on page 141.

1. Select **Help > Sample Data Library** and open Reliability/CD Step Stress.jmp and Reliability/CD Step Stress Pattern.jmp.
2. In the CD Step Stress data table, run the script Cumulative Damage.
3. Click the Stress Patterns red triangle and select **Simulate**.

Figure 5.8 Simulation Configuration Panel



The Simulation Configuration panel appears in the Stress Patterns report. The selection for Distribution is Weibull. The fitted values for b0 and b1 are used as initial values for the simulation.

4. Select Exponential for **Distribution**.
5. Enter 10 for **b0**.

6. Enter -18 for **b1**.
7. (Optional) Enter 14678 for **Random Seed**.
8. Click **Simulate**.

Figure 5.9 Partial Results of Simulation

	Time Left	Time Right	Pattern ID
1	140.67948106	140.67948106	1
2	139.34770497	139.34770497	1
3	136.95430716	136.95430716	1
4	136.30050091	136.30050091	1
5	145	•	1
6	118.5123739	118.5123739	1
7	119.03562567	119.03562567	1
8	145	•	1
9	127.19420127	127.19420127	1
10	118.48391055	118.48391055	1
11	103.54004804	103.54004804	1
12	145	•	1
13	121.04292162	121.04292162	1
14	136.15486483	136.15486483	1
15	143.65666248	143.65666248	1
16	123.04501594	123.04501594	1
17	145	•	1
18	124.26692664	124.26692664	1
19	120.96407573	120.96407573	1

Figure 5.9 shows a partial listing of the simulated data table. The stress pattern with Pattern ID equal to 1 is defined only up to 145 time units. Since the Pattern Configuration setting in the launch window was set to Terminate, the simulation censors any simulated values at 145 for stress pattern 1.

Chapter 6

Recurrence Analysis

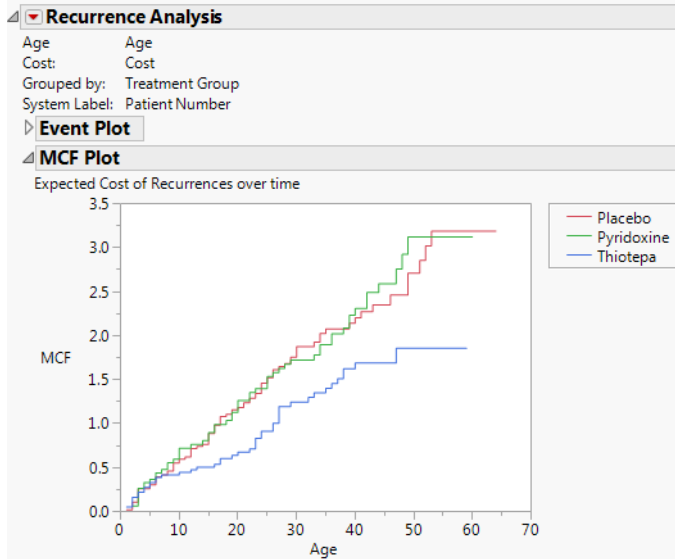
Model the Frequency or Cost of Recurrent Events over Time

The Recurrence Analysis platform analyzes event times, where the events can recur several times for each unit, item, or person. In an industrial setting, these events can occur when a unit breaks down, is repaired, and then put back into service after the repair. The units are followed until they are ultimately taken out of service.

In a medical setting, recurrence analysis can be used to analyze data from continuing treatments of a long-term disease, such as the recurrence of tumors in patients receiving treatment for bladder cancer.

The goal of the analysis is to obtain the mean cumulative function (MCF), which shows the total cost per unit as a function of time. Cost can be a count of the number of repairs, or it can be the actual cost of repair.

Figure 6.1 Recurrence Analysis Example



Contents

Overview of the Recurrence Analysis Platform 159

Example of the Recurrence Analysis Platform 159

Launch the Recurrence Analysis Platform 162

Recurrence Analysis Platform Options 163

 Fit Model 164

Additional Examples of the Recurrence Analysis Platform 168

 Bladder Cancer Recurrences Example 169

 Diesel Ship Engines Example 172

Overview of the Recurrence Analysis Platform

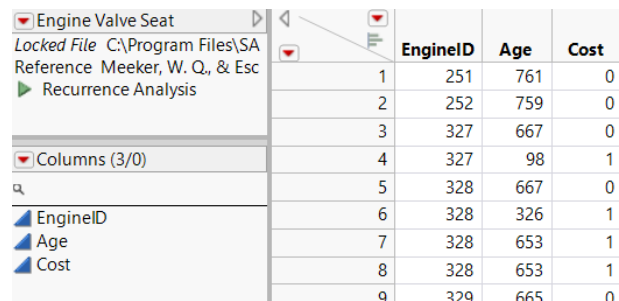
Recurrent event data involves the cumulative frequency or cost of repairs as units age. In JMP, the Recurrence Analysis platform analyzes recurrent events data.

The data for recurrence analysis have one row for each observed event and a closing row with the last observed age of a unit. Any number of units or systems can be included. In addition, these units or systems can include any number of recurrences.

Example of the Recurrence Analysis Platform

A typical unit might be a system, such as a component of an engine or appliance. For example, consider the sample data table *Engine Valve Seat.jmp*, which records valve seat replacements in locomotive engines. See Meeker and Escobar (1998, p. 395) and Nelson (2003). A partial listing of this data is shown in Figure 6.2. The *EngineID* column identifies a specific locomotive unit. *Age* is time in days from beginning of service to replacement of the engine valve seat. Note that an engine can have multiple rows with its age at each replacement and its cost, corresponding to multiple repairs. Here, *Cost*=0 indicates the last observed age of a locomotive.

Figure 6.2 Partial Engine Valve Seat Data Table

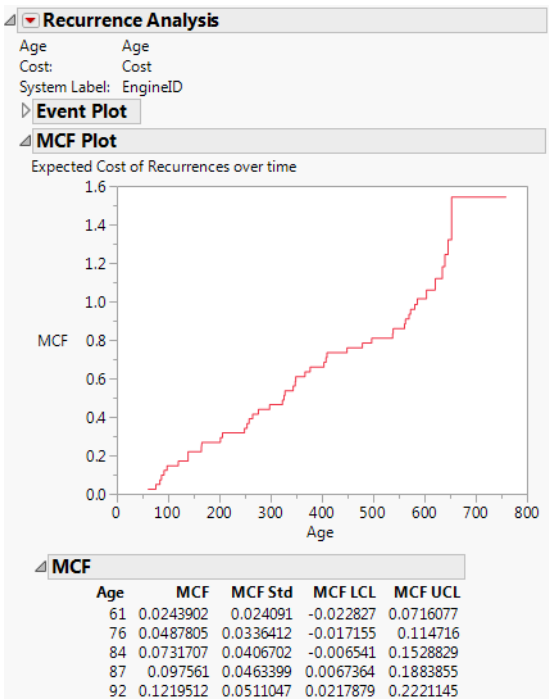


EngineID	Age	Cost
1	251	761
2	252	759
3	327	667
4	327	98
5	328	667
6	328	326
7	328	653
8	328	653
9	329	665

Complete the launch window as shown in Figure 6.5.

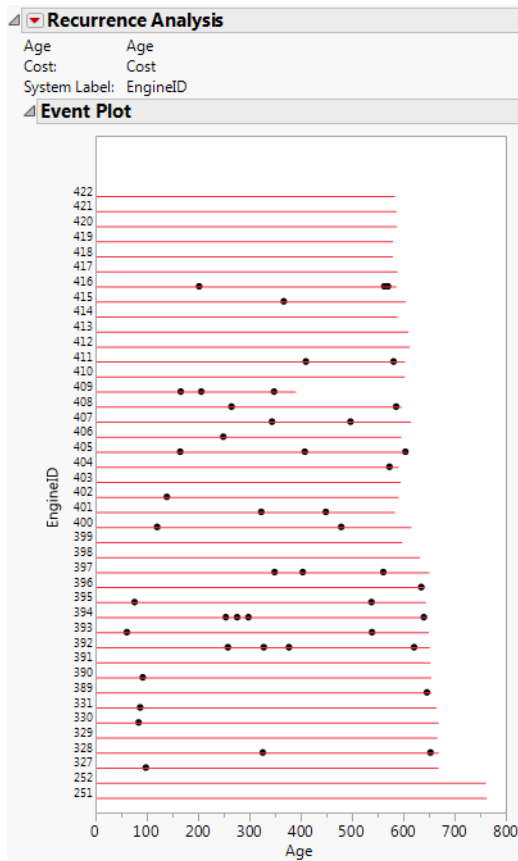
When you click **OK**, the Recurrence platform shows the reports in Figure 6.3 and Figure 6.4. The MCF plot shows the sample mean cumulative function. For each age, this is the nonparametric estimate of the mean cumulative cost or number of events per unit. This function goes up as the units get older and total costs grow. The plot in Figure 6.3 shows that about 580 days is the age that averages one repair event.

Figure 6.3 MCF Plot and Partial Table for Recurrence Analysis



The event plot in Figure 6.4 shows a time line for each unit. There are markers at each time of repair, and each line extends to that unit’s last observed age. For example, unit 409 was last observed at 389 days and had three valve replacements.

Figure 6.4 Event Plot for Valve Seat Replacements



Launch the Recurrence Analysis Platform

Launch the Recurrence Analysis platform by selecting **Analyze > Reliability and Survival > Recurrence Analysis**.

Figure 6.5 Recurrence Analysis Launch Window

Analyzes recurring event history.

Select Columns

▼ 3 Columns

- EngineID
- Age
- Cost

☐ First Event is Start Timestamp

Age Scaling: No scaling ▼

Default End Timestamp:

Cast Selected Columns into Roles

Y, Age, Event Timestamp	Age
Label, System ID	EngineID
Cost	Cost
Grouping	optional
Cause	optional
Timestamp at Start	optional numeric
Timestamp at End	optional numeric
By	optional

•If the Y column is an event timestamp rather than an age, then you need to specify information that allows JMP to calculate age.
 •JMP calculates age by subtracting the values in the Timestamp at Start column. If you have starting times as event records, select the First Event is Start Timestamp option.
 •Since timestamps are usually coded as seconds, and modeling is usually done in other time units, specify the Age Scaling option.
 •Recurrence also needs an end time (out-of-service or end-of-study), which is usually a record in the data where cost=0. If end times are given for all units, specify the Timestamp at End column. If end times are not given for all units and you are using timestamps, specify the Default End Timestamp.

Action

OK

Cancel

Remove

Recall

Help

For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

Y, Age, Event Timestamp Specifies either the unit's age at the time of an event or the timestamp of the event. If the Y column is an event timestamp, then you must specify the start and the end timestamp so that JMP can calculate age.

Label, System ID Identifies the unit for each event and censoring age.

Cost Identifies a column that must contain one of the following values:

- A 1, indicating that an event has occurred (a unit failed or was repaired, replaced, or adjusted). When indicators (1s) are specified, the MCF is the mean cumulative count of events per unit as a function of age.
- A cost for the event (the cost of the repair, replacement, or adjustment). When costs are specified, the MCF is a mean cumulative cost per unit as a function of age and the markers in the Event Plot are sized by the cost values.

- A zero, indicating that the unit went out-of-service, or is no longer being studied. All units (each System ID) must have one row with a zero for this column, with the **Y, Age, Event Timestamp** column containing the final observed age. If each unit does not have exactly one last observed age in the table (where the Cost column cell is zero), then an error message appears.

Note: Cost indicators for **Recurrence Analysis** are the reverse of censor indicators seen in **Life Distribution** or **Survival Analysis**. For the cost variable, the value of 1 indicates an event, such as repair; the value of 0 indicates that the unit is no longer in service. For the censor variable, the value of 1 indicates censored values, and the value of 0 indicates the event or failure of the unit (non-censored value).

Grouping Produces separate MCF estimates for the different groups that are identified by this column.

Cause Specifies multiple failure modes.

Timestamp at Start Specifies the column with the origin timestamp. If you have starting times as event records, select the **First Event is Start Timestamp** option instead. JMP calculates age by subtracting the values in this column.

Timestamp at End Specifies the column with the end-of-service timestamp. If end times are given for all units, specify that column here. If end times are not given for all units, specify the **Default End Timestamp** option instead. But if you have a record in which Cost is equal to zero, JMP uses that record as the end timestamp and you do not need to specify this role.

Age Scaling Specifies the time units for modeling. For example, if your timestamps are coded in seconds, you can change them to hours.

Recurrence Analysis Platform Options

The Recurrence Analysis red triangle menu contains the following options:

MCF Plot Shows or hides the mean cumulative function (MCF) plot.

MCF Confid Limits Shows or hides lines corresponding to the approximate 95% confidence limits of the mean cumulative function (MCF).

Event Plot Shows or hides the Event Plot. If a Cost column is specified in the launch window, the markers in the Event Plot are sized by the values in the Cost column.

Calendar Event Plot (Available only when the events are designated by a timestamp rather than an age.) Shows or hides the Calendar Event Plot, which shows events with calendar

date on the horizontal axis. This plot is next to the Event Plot and the units in each plot are aligned vertically.

Plot Interarrival by Age Shows or hides the Interarrival by Age plot, which plots the time between successive events on the vertical axis and the event times on the horizontal axis. You can use this plot to determine whether there are changes in the time between events in your data. For a recurrence analysis, the interarrival times should be independent and identically distributed over time. For more information about interarrival plots, see Tobias and Trindade (2012, p. 420).

Plot MCF Differences (Available only when you specify a grouping variable.) Shows or hides a plot of the difference of MCFs, including a 95% confidence interval for that difference. The MCFs are significantly different where the confidence interval lines do not cross the zero line. This option is available only when you specify a grouping variable.

MCF Plot Each Group (Available only when you specify a grouping variable.) Shows or hides a report that contains a mean cumulative function (MCF) plot for each level of the grouping variable.

This option can be used to get an MCF Plot for each unit if the **Label, System ID** variable is also specified as the **Grouping** variable.

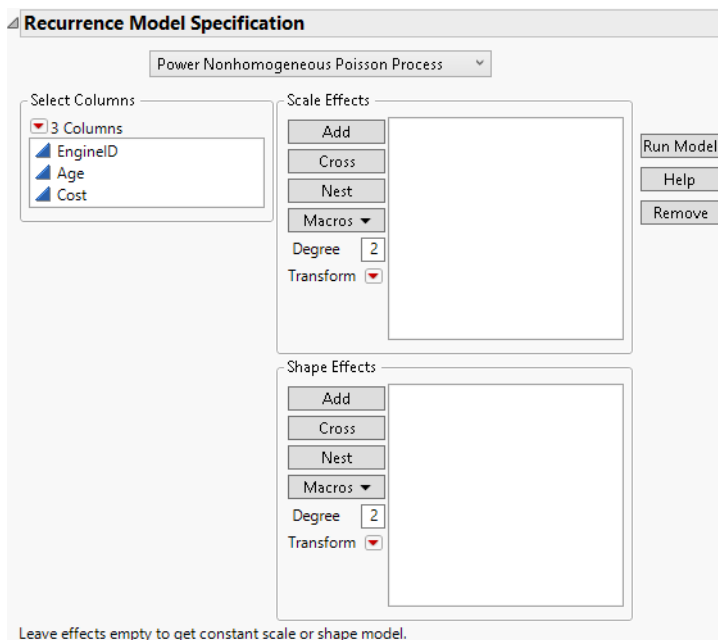
Fit Model Enables you to fit models for the Recurrence Intensity and Cumulative functions. See “[Fit Model](#)” on page 164.

Fit Model

The Fit Model option is used to fit models for the Recurrence Intensity and Cumulative functions. There are four models available for describing the intensity and cumulative functions. You can fit the models with constant parameters, or with parameters that are functions of effects.

Select Fit Model from the Recurrence Analysis red triangle menu to produce the Recurrence Model Specification window shown in Figure 6.6.

Figure 6.6 Recurrence Model Specification



You can select one of four models, with the following Intensity and Cumulative functions:

Power Nonhomogeneous Poisson Process

$$I(t) = \left(\frac{\beta}{\theta}\right) \left(\frac{t}{\theta}\right)^{\beta-1}$$

$$C(t) = \left(\frac{t}{\theta}\right)^{\beta}$$

Proportional Intensity Poisson Process

$$I(t) = \delta t^{\delta-1} e^{\gamma}$$

$$C(t) = t^{\delta} e^{\gamma}$$

Loglinear Nonhomogeneous Poisson Process

$$I(t) = e^{\gamma + \delta t}$$

$$C(t) = \frac{I(t) - I(0)}{\delta} = \frac{e^{\gamma + \delta t} - e^{\gamma}}{\delta}$$

Homogeneous Poisson Process

$$I(t) = e^{\gamma}$$

$$C(t) = te^{\gamma}$$

where t is the age of the product.

Table 6.1 defines each model parameter as a scale parameter or a shape parameter.

Table 6.1 Scale and Shape Parameters

Model	Scale Parameter	Shape Parameter
Power NHPP	θ	β
Proportional Intensity PP	γ	δ
Loglinear NHPP	γ	δ
Homogeneous PP	γ	none

Note the following:

- For the Recurrence Model Specification window (Figure 6.6), if you include Scale Effects or Shape Effects, the scale and shape parameters in Table 6.1 are modeled as functions of the effects. To fit the models with constant scale and shape parameters, do not include any Scale Effects or Shape Effects.
- The Homogeneous Poisson Process is a special case compared to the other models. The Power NHPP and the Proportional Intensity Poisson Process are equivalent for one-term models, but the Proportional Intensity model seems to fit more reliably for complex models.

Click **Run Model** to fit the model and see the model report.

Figure 6.7 Model Report

Fitted Recurrence Model		
Power Nonhomogeneous Poisson Process		
Parameter Estimates		
Parameter	Estimate	Std Error
θ Intercept	553.64302	57.863577
β Constant	1.3995793	0.2005022
-2Loglikelihood 692.9806		
Converged in Gradient, 8 iterations		
Covariance of Estimates		
Covariance		
	θ Intercept	β Constant
θ Intercept	3348.19	1.88248
β Constant	1.88248	0.0402
Correlation		
	θ Intercept	β Constant
θ Intercept	1.0000	0.1623
β Constant	0.1623	1.0000

The Fitted Recurrence Model red triangle menu contains the following options:

Profiler Shows or hides the Profiler showing the Intensity and Cumulative functions.

Effect Marginals Evaluates the parameter functions for each level of the categorical effect, holding other effects at neutral values. This helps you see how different the parameter functions are between groups. This is available only when you specify categorical effects.

Test Homogeneity Tests if the process is homogeneous. This option is not available for the Homogeneous Poisson Process model.

Effect Likelihood Ratio Test Produces a test for each effect in the model. This option is available only if there are effects in the model.

Specific Intensity and Cumulative Computes the intensity and cumulative values associated with particular time and effect values. The confidence intervals are profile likelihood intervals.

Specific Time for Cumulative Computes the time associated with a particular number of recurrences and effect values.

Save Intensity Formula Saves the Intensity formula to the data table.

Save Cumulative Formula Saves the Cumulative formula to the data table.

Publish Intensity Formula Creates the Intensity formula and saves it as a formula column script in the Formula Depot platform. If a Formula Depot report is not open, this option creates a Formula Depot report. See *Predictive and Specialized Modeling*.

Publish Cumulative Formula Creates the Cumulative formula and saves it as a formula column script in the Formula Depot platform. If a Formula Depot report is not open, this option creates a Formula Depot report. See *Predictive and Specialized Modeling*.

Simulate from Model Enables you to simulate new data from the estimated recurrence model. This option creates a new data table of simulated observations that are based on options that are specified in the Simulate from Model window. See [“Simulate from Model”](#) on page 168.

Remove Fit Removes the model report.

Simulate from Model

When you select the Simulate from Model option from the Fitted Recurrence Model red triangle menu, the Simulate from Model window appears. This window contains the following specifications for the simulation:

Maximum Number of Events Specifies the maximum number of events to be simulated in each level of the System ID.

Maximum Age Specifies the maximum time for a simulated event.

Number of Units Specifies the number of units in each level of the System ID.

Note: If the model contains terms other than the Intercept and Constant terms, the simulated data table contains this number of units for all level combinations of the regression terms. If the regression term is continuous, it is divided into 5 levels.

After you click OK, a new data table that contains the results appears. This table contains a script that enables you to analyze the simulated observations in the Recurrence platform.

Additional Examples of the Recurrence Analysis Platform

- [“Bladder Cancer Recurrences Example”](#)
- [“Diesel Ship Engines Example”](#)

Bladder Cancer Recurrences Example

The sample data file *Bladder Cancer.jmp* contains data on cancerous bladder tumor recurrences from the Veteran's Administration Co-operative Urological Research Group. See Andrews and Herzberg (1985, table 45). All patients presented with superficial bladder tumors, which were removed upon entering the trial. Each patient was then assigned to one of three treatment groups: placebo pills, pyridoxine (vitamin B6) pills, or periodic chemotherapy with Thiotepea. The following analysis of tumor recurrence explores the progression of the disease, and whether there is a difference among the three treatments.

Launch the platform with the options shown in Figure 6.8.

Figure 6.8 Bladder Cancer Launch Window

Analyzes recurring event history.

Select Columns

7 Columns

Patient Number

Treatment Group

Cause of Death

Initial Number of Tumors

Initial Size of Tumors

Age

Cost

☐ First Event is Start Timestamp

Age Scaling: No scaling

Default End Timestamp

Cast Selected Columns into Roles

Y, Age, Event Timestamp	Age
Label, System ID	Patient Number
Cost	Cost
Grouping	Treatment Group
Cause	optional
Timestamp at Start	optional numeric
Timestamp at End	optional numeric
By	optional

- If the Y column is an event timestamp rather than an age, then you need to specify information that allows JMP to calculate age.
- JMP calculates age by subtracting the values in the Timestamp at Start column. If you have starting times as event records, select the First Event is Start Timestamp option.
- Since timestamps are usually coded as seconds, and modeling is usually done in other time units, specify the Age Scaling option.
- Recurrence also needs an end time (out-of-service or end-of-study), which is usually a record in the data where cost=0. If end times are given for all units, specify the Timestamp at End column. If end times are not given for all units and you are using timestamps, specify the Default End Timestamp.

Action

OK

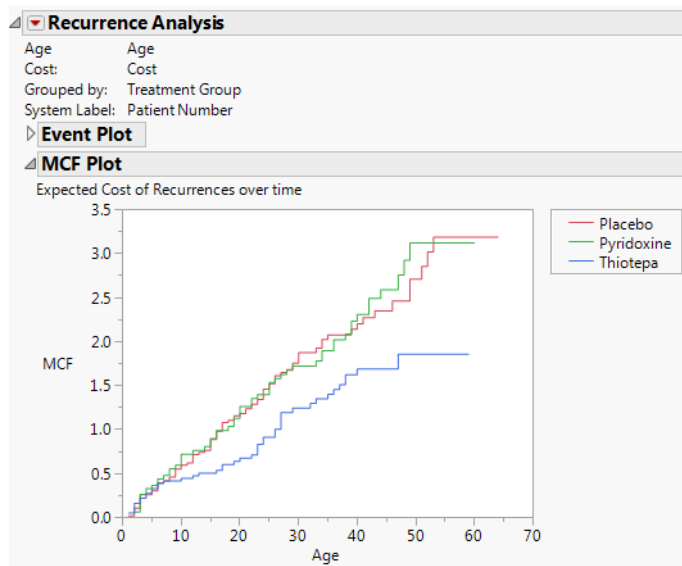
Cancel

Remove

Recall

Help

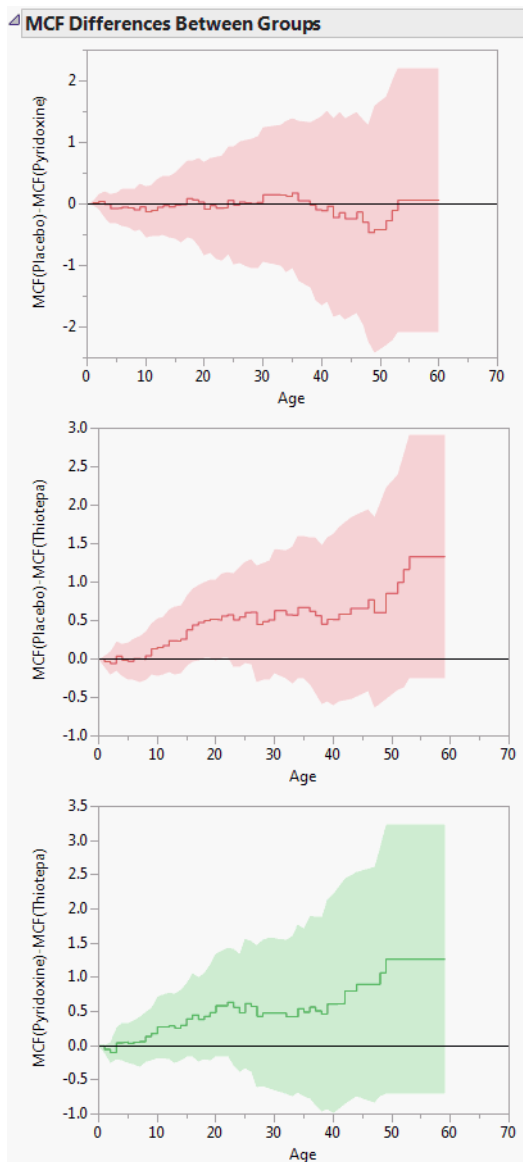
Figure 6.9 shows the MCF plots for the three treatments.

Figure 6.9 Bladder Cancer MCF Plot

Note that all three of the MCF curves are essentially straight lines. The slopes (rates of recurrence) are therefore constant over time, implying that patients do not seem to get better or worse as the disease progresses.

To examine if there are differences among the treatments, click the Recurrence Analysis red triangle and select **Plot MCF Differences**.

Figure 6.10 MCF Differences



To determine whether there is a statistically significant difference between treatments, examine the confidence limits on the differences plot. If the limits do not include zero, the treatments are convincingly different. The graphs in Figure 6.10 show there is no significant difference among the treatments.

Diesel Ship Engines Example

The sample data table Diesel Ship Engines.jmp contains data on engine repair times for two ships (Grampus4 and Halfbeak4) that have been in service for an extended period of time. See Meeker and Escobar (1998). You want to examine the progression of repairs and gain a sense of how often repairs might need to be done in the future. These observations can help you decide when an engine should be taken out of service.

1. Select **Help > Sample Data Library** and open Reliability/Diesel Ship Engines.jmp.
2. Ensure that rows 57 and 129 are set as Excluded.

Note: If they are not set to Excluded, select rows 57 and 129 and select **Rows > Exclude/Unexclude**.

3. Select **Analyze > Reliability and Survival > Recurrence Analysis**.
4. Complete the launch window as shown in Figure 6.11.

Figure 6.11 Diesel Ship Engines Launch Window

Analyzes recurring event history.

Select Columns

▼ 7 Columns

- Unit
- kHours
- Cost
- System ID
- orig time
- event time
- end time

☐ First Event is Start Timestamp

Age Scaling DateTime to Hour ▼

Default End Timestamp

Cast Selected Columns into Roles

Y, Age, Event Timestamp	event time
Label, System ID	System ID
Cost	optional numeric
Grouping	System ID
Cause	optional
Timestamp at Start	orig time
Timestamp at End	end time
By	optional

•If the Y column is an event timestamp rather than an age, then you need to specify information that allows JMP to calculate age.
 •JMP calculates age by subtracting the values in the Timestamp at Start column. If you have starting times as event records, select the First Event is Start Timestamp option.
 •Since timestamps are usually coded as seconds, and modeling is usually done in other time units, specify the Age Scaling option.
 •Recurrence also needs an end time (out-of-service or end-of-study), which is usually a record in the data where cost=0. If end times are given for all units, specify the Timestamp at End column. If end times are not given for all units and you are using timestamps, specify the Default End Timestamp.

Action

OK

Cancel

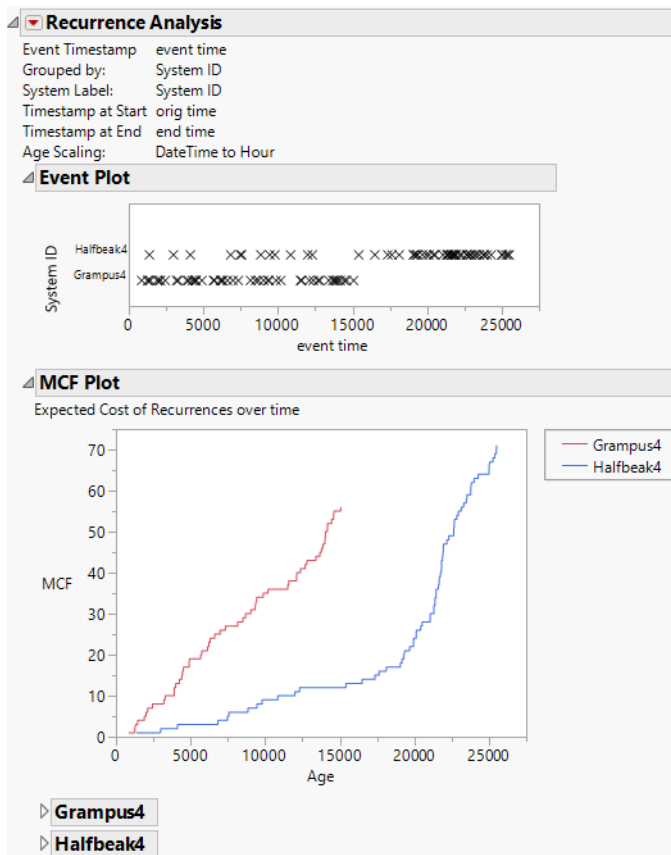
Remove

Recall

Help

5. Click **OK**.

Figure 6.12 Diesel Ship Engines Report

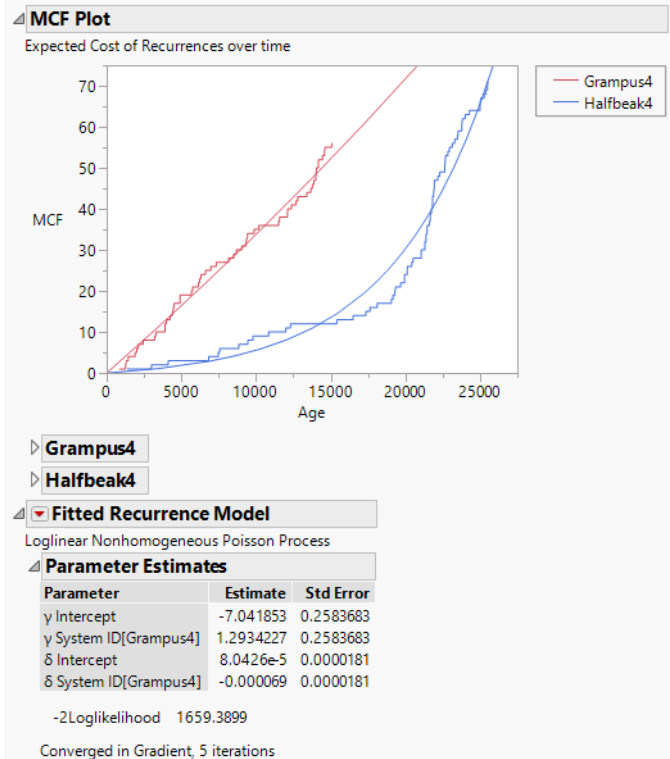


Looking at the Event Plot, you can see that repairs for the Grampus4 engine have been relatively consistent. Repairs for the Halfbeak4 engine have been more sporadic, and there appears to be an abrupt increase in repairs somewhere around the 19,000 hour mark. This increase is even more obvious in the MCF Plot.

Continue your analysis by fitting a parametric model to help predict future performance.

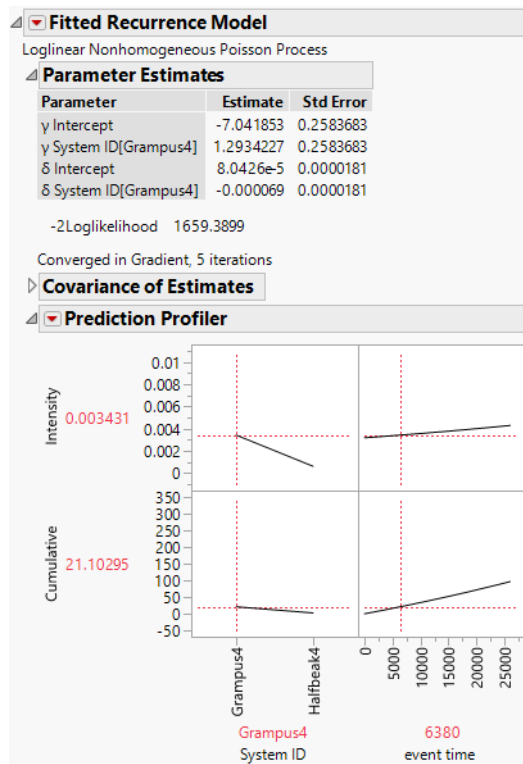
6. Click the Recurrence Analysis red triangle and select **Fit Model**.
7. In the Recurrence Model Specification, select the **Loglinear Nonhomogeneous Poisson Process**.
8. Add the System ID column as both a Scale Effect and a Shape Effect.
9. Click **Run Model**.

Figure 6.13 Diesel Ship Engines Fitted Model



10. Click the Fitted Recurrence Model red triangle and select **Profiler**.

Figure 6.14 Diesel Ship Profiler



Compare the number of future repairs for the Grampus4 engine to the Halfbeak4 engine. Change the event time value to see the effect on the cumulative number of future repairs.

- To see how many repairs will be needed after 30,000 hours of service, type 30,000 for the event time. The Grampus4 engine will require about 114 repairs. To see the values for Halfbeak4, click and drag the dotted line from Grampus4 to Halfbeak4. The Halfbeak4 engine will require about 140 repairs.
- To see how many repairs will be needed after 80,000 hours of service, type 80,000 for the event time. The Halfbeak4 engine will require about 248,169 repairs. Click and drag the dotted line from Halfbeak4 to Grampus4. The Grampus4 engine will require about 418 repairs.

You can conclude that in the future, the Halfbeak4 engine will require many more repairs than the Grampus4 engine.

Chapter 7

Degradation

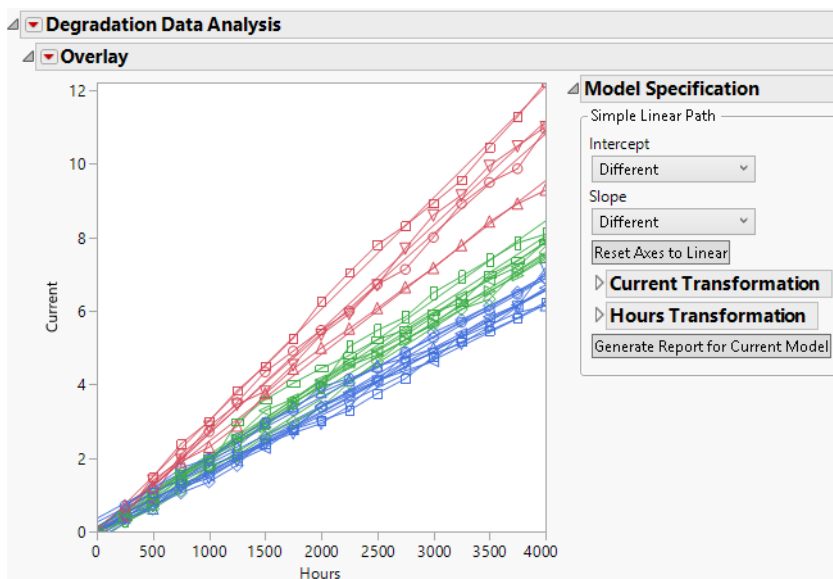
Model Product Deterioration over Time

Using the Degradation platform, you can analyze degradation data to predict pseudo failure times. These pseudo failure times can then be analyzed by other reliability platforms to estimate failure distributions.

Both linear and nonlinear degradation paths can be modeled. You can specify an accelerating factor to analyze accelerated degradation data.

You can also perform stability analysis, which is useful when setting pharmaceutical product expiration dates.

Figure 7.1 Degradation Analysis Example



Contents

Overview of the Degradation Platform 179

Example of the Degradation Platform 179

Launch the Degradation Platform 181

The Degradation Reports 183

Model Specification..... 186

 Simple Linear Path 186

 Nonlinear Path..... 188

Inverse Prediction 196

Prediction Graph..... 198

Degradation Platform Options 199

Model Reports 201

 Model Lists..... 201

 Reports 202

Destructive Degradation 203

Stability Analysis..... 208

Overview of the Degradation Platform

In reliability analyses, the primary objective is to model the failure times of the product under study. In many situations, these failures occur because the product degrades (weakens) over time. But, sometimes failures do not occur. In these situations, modeling the product degradation over time is helpful in making predictions about failure times. When an accelerating factor is included in the data, you can fit an accelerated degradation model.

The Degradation platform can model data that follows linear or nonlinear degradation paths. If a path is nonlinear, transformations are available to linearize the path. If linearization is not possible, you can specify a nonlinear model.

You can also use the Degradation platform to perform stability analysis. Three types of linear models are fit, and an expiration date is estimated. Stability analysis is used in setting pharmaceutical product expiration dates.

Example of the Degradation Platform

This example uses the GaAs Laser.jmp data table from Meeker and Escobar (1998), which contains measurements of the percent increase in operating current taken on several gallium arsenide lasers. When the percent increase reaches 10%, the laser is considered to have failed.

1. Select **Help > Sample Data Library** and open Reliability/GaAs Laser.jmp.
2. Select **Analyze > Reliability and Survival > Degradation**.
3. Select Current and click **Y, Response**.
4. Select Hours and click **Time**.
5. Select Unit and click **Label, System ID**.
6. Type 10 in the text box for **Upper Spec Limit**.
7. Click **OK**.

Figure 7.2 Initial Degradation Report

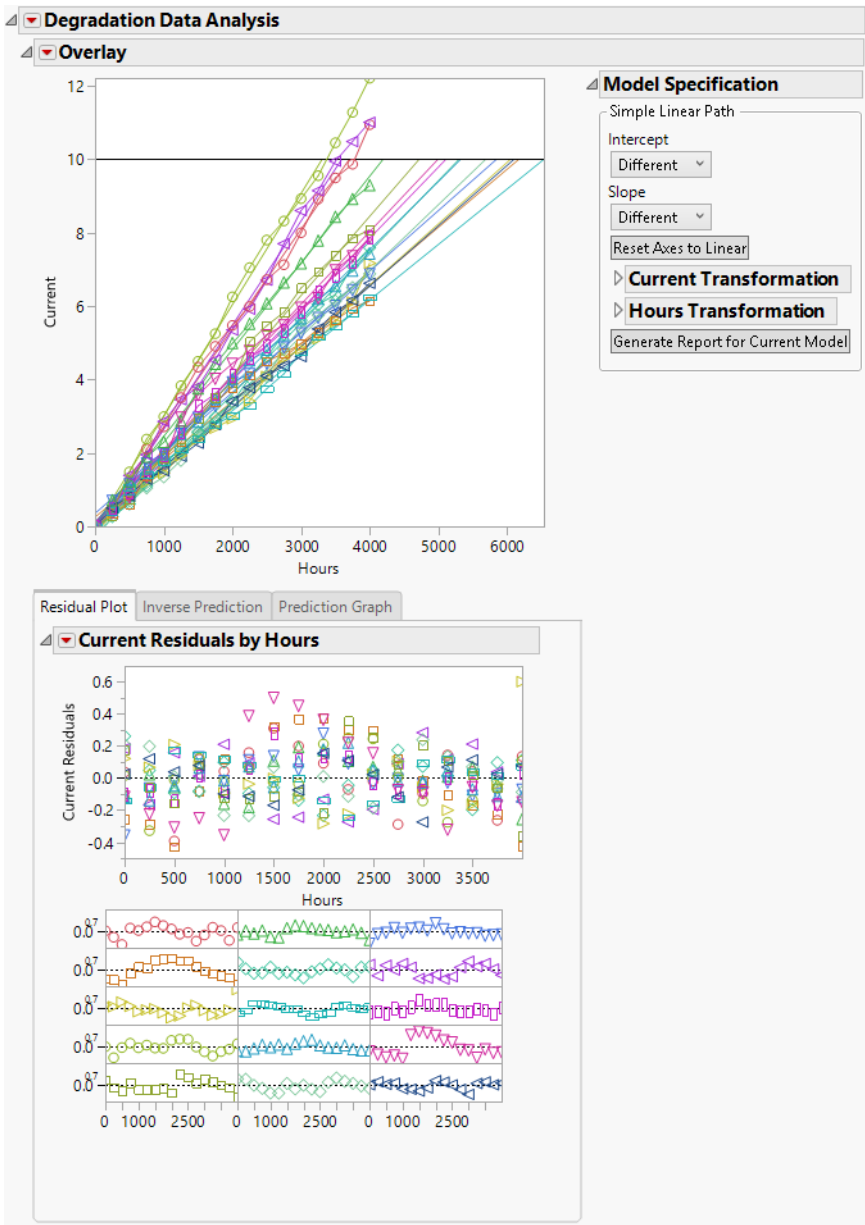


Figure 7.2 shows the initial Degradation report. The Overlay plot shows the measurements of Current versus Time for each unit in the data. The horizontal line at Current = 10 corresponds to the upper specification limit at 10%. Units with values above this limit are considered to have failed. Three of the fifteen units have reached that point by the end of the study period. The Inverse Prediction outline shows the predicted Hours value for which each unit fails, based on your specified model.

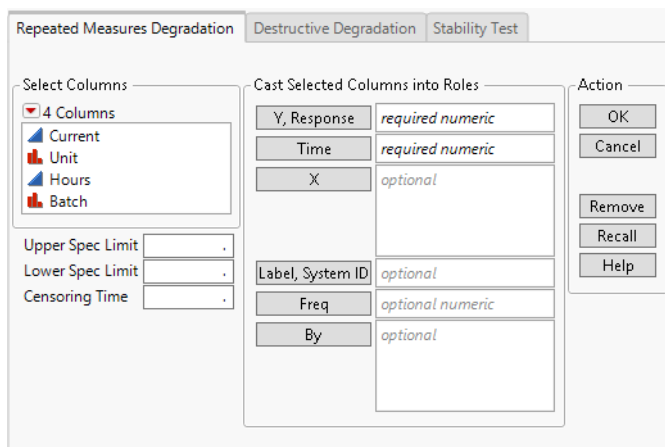
The default model fits a separate slope and intercept for each unit, using linear transformations. You can fit other models using the Model Specification outline.

The Residual Plot tab in Figure 7.2 shows residuals based on your specified model. The top plot shows residuals for all units plotted against Hours and overlaid on one plot. The bottom plot shows individual plots of the residuals for each unit in a rectangular array.

Launch the Degradation Platform

Launch the Degradation platform by selecting **Analyze > Reliability and Survival > Degradation**. Figure 7.3 shows the Degradation launch window using the GaAs Laser.jmp data table (located in the Reliability folder). For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

Figure 7.3 The Degradation Launch Window



Analysis Types

The launch window is split into three tabs, representing three different types of analyses:

Repeated Measures Degradation Performs linear or nonlinear degradation analysis. This option allows only one Y, Response variable. It does not allow censoring.

Destructive Degradation Choose this type of analysis if units are destroyed during the measurement process. This option allows censoring. See [“Destructive Degradation”](#) on page 203.

Note: The Destructive Degradation platform provides a flexible collection of predefined models for destructive testing. See the [“Destructive Degradation”](#) chapter on page 213.

Stability Test Performs a stability analysis for setting pharmaceutical product expiration dates. This option allows only one Y, Response variable. See [“Stability Analysis”](#) on page 208.

Launch Window Options

The launch window contains the following options:

Y, Response Assign the column with degradation measurements.

Time Assign the column containing the time values.

X (Available only in the Repeated Measures Degradation and Destructive Degradation tabs.) Assign an explanatory variable. Use this role to specify the accelerating factor in an accelerated degradation model.

Label, System ID (Available only in the Repeated Measures Degradation and Stability Test tabs.) Assign the column that designates the unit IDs.

Freq Assign a column giving a frequency for each row.

Censor (Available only in the Destructive Degradation tab.) Assign a column that designates if a unit is censored.

By Assign a variable to produce an analysis for each level of the variable.

Censor Code (Available only in the Destructive Degradation tab.) Specify the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

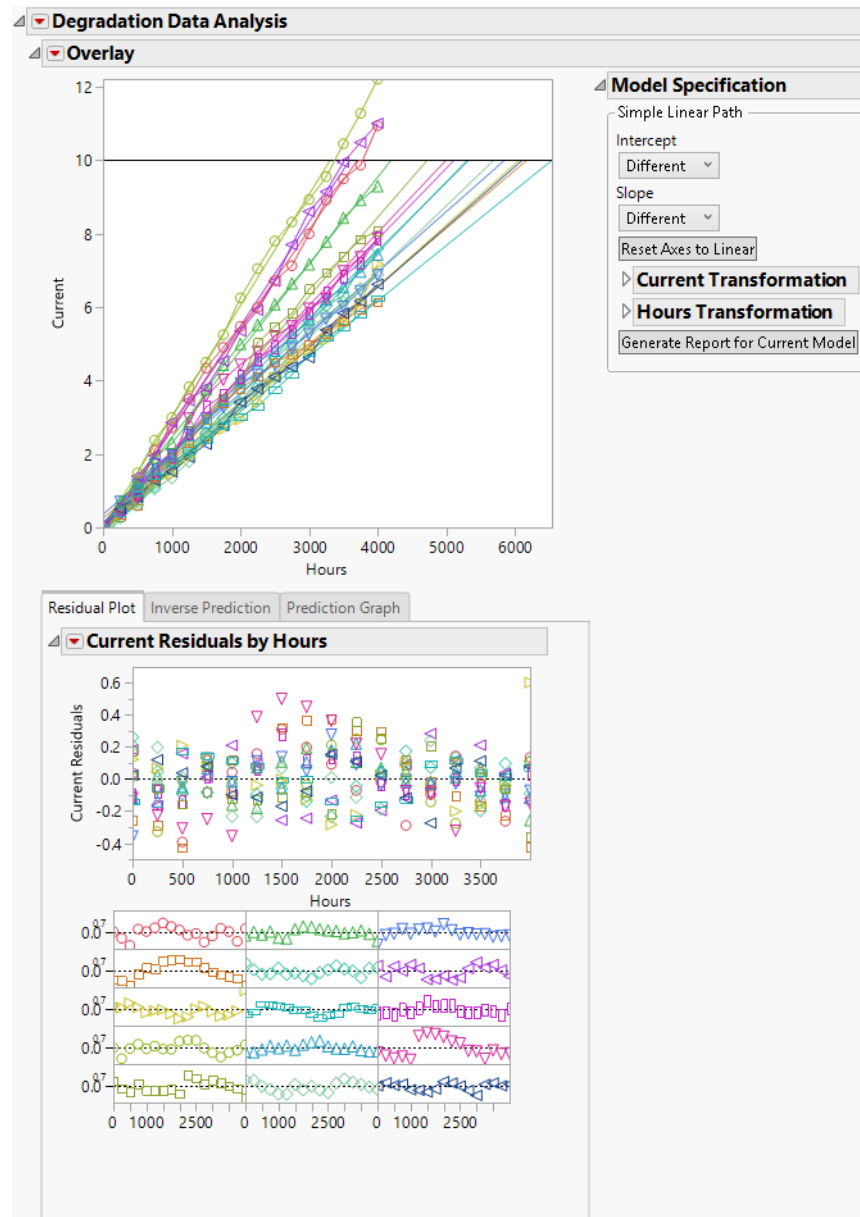
Upper Spec Limit Assign an upper specification limit. (Optional except for Stability Test tab.)

Lower Spec Limit Assign a lower specification limit. (Optional except for Stability Test tab.)

Censoring Time (Available only in the Repeated Measures Degradation and Stability Test tabs.) Assign a Time value that represents censoring of pseudo failures when you use Inverse Prediction. See [“Inverse Prediction”](#) on page 196.

The Degradation Reports

The Repeated Measures Degradation and Destructive Degradation methods show a report that fits a default model. As shown in the Model Specification outline in Figure 7.4, this model fits each unit with its own intercept and slope, using a linear transformation of the response and time columns. Separate intercepts and slopes are fit for each value of the Label, System ID variable, or, if only an X variable is specified, separate intercepts and slopes are fit for each level of the X variable. The Stability Test method fits three models.

Figure 7.4 Initial Repeated Measures Degradation Report with Transformation Outlines Open


To reproduce this example, see [“Example of the Degradation Platform”](#) on page 179.

The reports for Repeated Measures Degradation, Destructive Degradation, and Stability Test include the following:

Overlay

An Overlay plot of the Y, Response variable versus the Time variable. In this example, the plot is of Current versus Hours. The Overlay plot red triangle menu has the Save Estimates option, which creates a new data table containing the estimated slopes and intercepts for all units.

Model Specification

Specify your model and generate a report for that model. See [“Model Specification”](#) on page 186. (Available only for the Repeated Measures Degradation and Destructive Degradation methods.)

Stability Tests Outline

Compare models and estimate expiration dates. See [“Stability Analysis”](#) on page 208. (Available only for the Stability Test method.)

Reports

Shows analysis results for three different models and the best model. See [“Stability Analysis”](#) on page 208. (Available for the Stability Test method by default. Also available for the Repeated Measures Degradation and Destructive Degradation methods after you click Generate Report for Current Model.)

Tabbed Reports

Residual Plot Shows a single residual plot with all the units overlaid and separate residual plots for each unit arranged in a rectangular grid. The red triangle menu has the following options:

Save Residuals Saves the residuals of the current model to a new data table.

Jittering Adds random noise to the points in the time direction. This is useful for visualizing the data if there are a lot of points clustered together.

Separate Groups Adds space between the groups to visually separate the groups. This option appears only when an X variable is specified on the platform launch window.

Jittering Scale Changes the magnitude of the jittering and group separation. This option appears only if Jittering is selected.

Inverse Prediction Enables you to predict the time at which the Y variable reaches a specified value. See [“Inverse Prediction”](#) on page 196.

Prediction Graph Enables you to predict the Y variable for a specified Time value. See “[Prediction Graph](#)” on page 198.

Model Specification

You can use the Model Specification outline to specify the model that you want to fit to the degradation data. There are two types of Model Specifications:

Simple Linear Path Used to model linear degradation paths, or nonlinear paths that can be transformed to linear. See “[Simple Linear Path](#)” on page 186.

Nonlinear Path Used to model nonlinear degradation paths, especially those that cannot be transformed to linear. See “[Nonlinear Path](#)” on page 188.

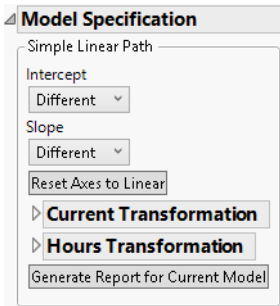
To change between the two specifications, use the Degradation Path Style submenu from the Degradation Data Analysis red triangle menu.

Simple Linear Path

To model linear degradation paths, select **Degradation Path Style > Simple Linear Path** from the Degradation Data Analysis red triangle menu.

Use the Simple Linear Path Model specification to specify the form of the linear model that you want to fit to the degradation path. You can model linear paths, or nonlinear paths that can be transformed to linear.

Figure 7.5 Simple Linear Path Model Specification



The Simple Linear Path model specification contains the following options:

Intercept Specifies the form of the intercept in the model:

Different Fits a different intercept for each level of the ID variable.

Common in Group Fits the same intercept for each level of the ID variable in the same level of the X variable, and fits different intercepts between levels.

Common Fits the same intercept for all levels of the ID variable.

Zero Restricts the intercept to be zero for all levels of the ID variable.

Slope Specifies the form of the slope in the model:

Different Fits a different slope for each level of the ID variable.

Common in Group Fits the same slope for each level of the ID variable in the same level of the X variable, and fits different slopes between levels.

Common Fits the same slope for all levels of the ID variable.

Reset Axes to Linear Resets the Overlay plot axes to their initial settings.

<Y, Response> Transformation If a transformation on the Y variable can linearize the degradation path, select the transformation (Linear, $\ln(x)$, $\exp(x)$, x^2 , $\sqrt{x}$ or Custom) here. For more information about the Custom option, see [“Custom Transformations”](#) on page 187.

<Time> Transformation If a transformation for the Time variable can linearize the degradation path, select the transformation (Linear, $\ln(x)$, x^2 , $\sqrt{x}$ or Custom) here. For more information about the Custom option, see [“Custom Transformations”](#) on page 187.

Generate Report for Current Model Creates a report for the current model settings. This includes a Model Summary report, and Estimates report giving the parameter estimates. See [“Model Reports”](#) on page 201.

Custom Transformations

If you need to perform a transformation that is not given, use the Custom option. For example, to transform the response variable using $\exp(-x^2)$, enter the transformation as shown in the Scale box in Figure 7.6. Also, enter the inverse transformation in the Inverse Scale box.

Note: JMP automatically attempts to solve for the inverse transformation. If it can solve for the inverse, it automatically enters it in the Inverse Scale box. If it cannot solve for the inverse, you must enter it manually.

Figure 7.6 Custom Transformation Options

Current Transformation

☐ Linear
☐ $\ln(x)$
☐ $\exp(x)$
☐ x^2
☐ $\sqrt{x}$
☒ Custom

Scale
 Function({x}, $\exp(-x^2)$)

Inverse Scale
 Function({x}, $(-\text{Log}(x))^{(1/2)}$)

Linear

Hours Transformation

Name the transformation using the text box. When finished, click the **Use & Save** button to apply the transformation. Select a transformation from the menu if you have created multiple custom transformations. Click the **Delete** button to delete a custom transformation.

Nonlinear Path

To model nonlinear degradation paths, select **Degradation Path Style > Nonlinear Path** from the Degradation Data Analysis red triangle menu. This is useful if a degradation path cannot be linearized using transformations, or if you have a custom nonlinear model that you want to fit to the data.

To facilitate explaining the Nonlinear Path Model Specification, open the Device B.jmp data table. The data consists of power decrease measurements taken on 34 units, across four levels of temperature. Follow these steps:

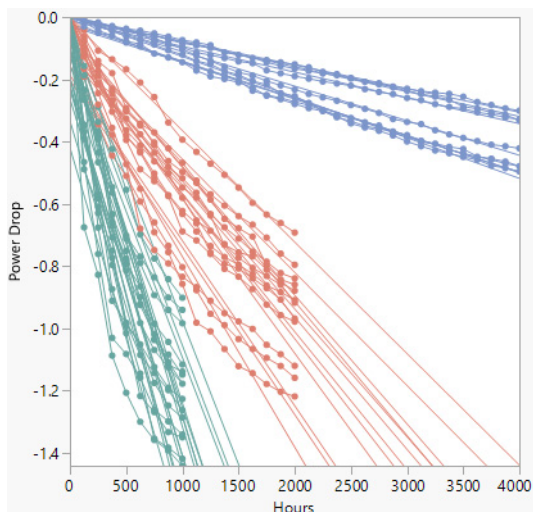
1. Select **Help > Sample Data Library** and open Reliability/Device B.jmp.
2. Select **Analyze > Reliability and Survival > Degradation**.
3. Select Power Drop and click **Y, Response**.
4. Select Hours and click **Time**.
5. Select Degrees C and click **X**.

The temperature setting is the accelerating factor in the experiment.

6. Select Device and click **Label, System ID**.
7. Click **OK**.

The initial overlay plot of the data appears.

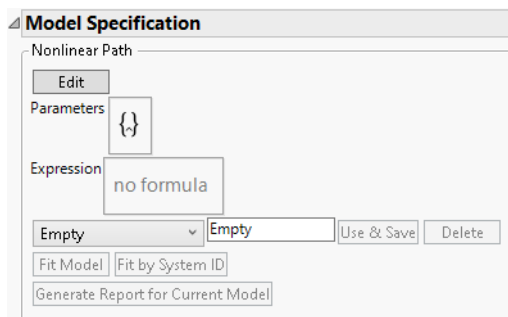
Figure 7.7 Device B Overlay Plot



The degradation paths appear linear for the first several hundred hours, but then start to curve. To fit a nonlinear model, select **Degradation Path Style > Nonlinear Path** from the Degradation Data Analysis red triangle menu to show the Nonlinear Path Model Specification outline (Figure 7.8).

Note: To view the Edit button displayed in Figure 7.8, you must select the interactive formula editor preference (**File > Preferences > Platforms > Degradation > Use Interactive Formula Editor**) before launching the Degradation platform.

Figure 7.8 Initial Nonlinear Model Specification Outline



The first step to create a model is to select one of the options on the menu initially labeled **Empty**:

- For more information about Reaction Rate models, see [“Reaction Rate Models”](#) on page 190.

- For more information about Constant Rate models, see [“Constant Rate Models”](#) on page 190.
- For more information about using a Prediction Column, see [“Prediction Columns”](#) on page 191.

Reaction Rate Models

The Reaction Rate options are applicable when the degradation occurs from a single chemical reaction, and the reaction rate is a function of temperature only. Select **Reaction Rate** or **Reaction Rate Type 1** from the menu shown in Figure 7.8. Although similar to the Reaction Rate model, the Reaction Rate Type 1 model contains an offset term that changes the basic assumption concerning the response value’s sign.

The Setup window prompts you to select the temperature scale, and the baseline temperature. The baseline temperature is used to generate initial estimates of parameter values. The baseline temperature should be representative of the temperatures used in the study.

For this example, select **Reaction Rate** and then select Celsius as the Temperature Unit. Click **OK** to return to the report. For more information about all the features for Model Specification, see [“Model Specification Details”](#) on page 192.

Constant Rate Models

The Constant Rate option is for modeling degradation paths that are linear with respect to time (or linear with respect to time after transforming the response or time). The reaction rate is a function of temperature only.

Select **Constant Rate** from the menu shown in Figure 7.8. The Constant Rate Model Settings window prompts you to enter transformations for the Path, Rate, and Time.

Figure 7.9 Constant Rate Transformation

The dialog box is titled "Constant Rate Transformation". It contains three panels for selecting transformations:

- Path Transformation:**
 - No Transformation (selected)
 - Exp
 - Log
 - Custom
- Rate Transformation:**
 - No Transformation
 - Arrhenius Kelvin
 - Arrhenius Celsius (selected)
 - Arrhenius Fahrenheit
 - Power
 - Exponential
- Time Transformation:**
 - No Transformation (selected)
 - Sqrt
 - Custom

Below these panels is the **Path Definition** area, which contains the following formula:

$$:Power Drop = Beta0 + Exp \left(Beta1 + Beta2 \cdot \left(\frac{-11604.51812}{(Celsius + 273.15)} \right) \right) \cdot :Hours + e$$

At the bottom right of the dialog are "OK" and "Cancel" buttons.

Once a selection is made for each transformation, the associated formula appears in the lower left corner as shown in Figure 7.9.

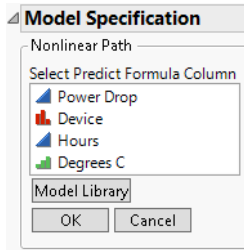
After all selections are made, click **OK** to return to the report. For more information about all the features for Model Specification, see ["Model Specification Details"](#) on page 192.

Prediction Columns

The Prediction Column option enables you to use a custom model that is stored as a formula in a data table column. The easiest approach is to create the formula column before launching the Degradation platform. You can also create the formula column from within the Degradation platform if you want to use one of the built-in models in the Nonlinear Model Library.

For more information about how to create a custom model and store it as a column formula, see ["Fit a Custom Model"](#) on page 194 or *Predictive and Specialized Modeling*.

Select **Prediction Column** from the list that appears beneath the Expression area in Figure 7.8. The Model Specification outline changes to prompt you to select the column that contains the model.

Figure 7.10 Column Selection


At this point, do one of three things:

- If the model that you want to use already exists as a formula in a column of the data table, select the corresponding column here, and then click **OK**. You are returned to the Nonlinear Path Model Specification. For more information about all the features for that specification, see [“Model Specification Details”](#) on page 192.
- If the model that you want to use does not already exist in the data table, you can click the **Model Library** button to use one of the built-in models. For more information about using the Model Library button, see [“Model Library”](#) on page 195 or *Predictive and Specialized Modeling*. After the model is created, select **Redo > Redo Analysis** from the Degradation Data Analysis red triangle menu. Then, return to the column selection shown in Figure 7.10. Select the column that contains the model, and then click **OK**. You are returned to the Nonlinear Path Model Specification. For more information about all the features for that specification, see [“Model Specification Details”](#) on page 192.
- If the model that you want to use is not in the data table, and you do not want to use one of the built-in models, then you are not ready to use this model specification. First, create the model, relaunch the Degradation platform, and then return to the column selection (Figure 7.10). Select the column containing the model, and then click **OK**. You are returned to the Nonlinear Path Model Specification. For more information about all the features for that specification, see [“Model Specification Details”](#) on page 192.

Model Specification Details

After you select one of the model types and supply the required information, you are returned to the Nonlinear Path Model Specification window.

Note: To view the Edit button displayed in Figure 7.11, you must select the interactive formula editor preference (**File > Preferences > Platforms > Degradation > Use Interactive Formula Editor**) before launching the Degradation platform.

A model is now shown in the script box that uses the **Parameter** statement. Initial values for the parameters are estimated from the data. For more information about creating models that use parameters, see [“Fit a Custom Model”](#) on page 194 or *Predictive and Specialized Modeling*.

If desired, enter a name in the text box to name the model. For this example, use the name “Device RR”. After that, click the **Use & Save** button to enter the model and activate the other buttons and features. Figure 7.11 shows the Model Specification window after clicking the Use & Save button.

Figure 7.11 Model Specification

Model Specification

Nonlinear Path

Edit

Parameters {Dlnf == -1.4423, Ru = 0.0004946233, Ea = 0.6822868283}

Expression

$$Dlnf \cdot \left(1 - \exp \left(-Ru \cdot \exp \left(Ea \cdot \left(\frac{11604.518122}{(193.5 + 273.15)} - \frac{11604.518122}{(\text{Degrees C} + 273.15)} \right) \right) \cdot \text{Hours} \right) \right)$$

Device RR Use & Save Delete

Fit Model Fit by System ID

Generate Report for Current Model

Parameter	Estimate	Low	High	Fixed
Dlnf	-1.4423	-1.8189	-1.0657	☐
Ru	0.00049	0.00028	0.00071	☐
Ea	0.68229	0.65228	0.7123	☐

- The **Fit Model** button is used to fit the model to the data.
- The **Fit by System ID** is used to fit the model to every level of **Label, System ID**.
- The **Delete** button is used to delete a model from the model menu.
- The **Generate Report for Current Model** button creates a report for the current model settings. See “[Model Reports](#)” on page 201.

The initial parameter values are shown at the bottom, along with sliders for visualizing how changes in the parameters affect the model. The fitted lines are shown on the Overlay plot. Move the parameter sliders to see how changes affect the fitted lines.

Here are the parameters for the Reaction Rate model (Meeker and Escobar 1998):

- **Dlnf** (D_{∞}) - asymptotic degradation level
- **Ru** (R_U) - reaction rate at use temperature (temp_U)
- **Ea** (Ea) - reaction-specific activation energy

The above parameters are calculated as follows:

$$D(t; \text{temp}) = D_{\infty} \times \{1 - \exp[-R_U \times AF(\text{temp}) \times t]\}$$

where R_U is the reaction rate at use temperature $temp_U$, $R_U \times AF(temp)$ is the reaction rate at a general temperature $temp$, and for $temp > temp_U$, $AF(temp) > 1$

and

$$AF(temp) = \frac{R(temp)}{R(temp_U)} = \exp\left[Ea\left(\frac{11604.5181215503}{(temp_UK)} - \frac{11604.5181215503}{(tempK)}\right)\right]$$

where $temp_UK$ and $tempK$ are temperatures expressed on the Kelvin scale.

To compute the optimal values for the parameters, click the **Fit Model** or **Fit by System ID** button.

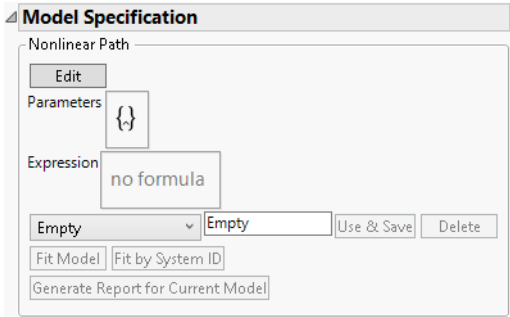
To fix a value for a parameter, check the box under **Fixed** for the parameter. When fixed, that parameter is held constant in the model fitting process.

Entering a Model with the Formula Editor

You can use the Formula Editor to enter a model. Click the **Edit** button to open the Formula Editor to enter parameters and the model. For more information about entering parameters and formulas in the Formula Editor, see *Using JMP*.

Note: To view the Edit button displayed in Figure 7.12, you must select the interactive formula editor preference (**File > Preferences > Platforms > Degradation > Use Interactive Formula Editor**) before launching the Degradation platform.

Figure 7.12 Alternate Model Specification Report



Fit a Custom Model

If you want to fit a custom model, you must first create a formula column with initial parameter estimates. This method requires a few more steps than fitting a built-in model, but it allows any nonlinear model to be fit. Also, you can provide a custom loss function, and specify several other options for the fitting process.

1. Open your data table.

2. Create a new column in the data table.
3. Open the Formula Editor for the new column.
4. Select **Parameters** from the list in the lower left corner.
5. Click **New Parameter**.
6. Enter the name of the parameter.
7. Enter the initial value of the parameter.

Repeat steps 4 to 6 to create all the parameters in the model.

8. Build the model formula using the data table columns, parameters, and formula editor functions.
9. Click **OK**.

Parameters for Models with a Grouping Variable

In the formula editor, when you add a parameter, note the check box for **Expand Into Categories, selecting column**. This option is used to add several parameters (one for each level of a categorical variable for example) at once. When you select this option, a window appears that enables you to select a column. After selection, a new parameter appears in the Parameters list with the name *D_column*, where D is the name that you gave the parameter. When you use this parameter in the formula, a Match expression is inserted, containing a separate parameter for each level of the grouping variable.

Model Library

The Model Library can assist you in creating a formula column with parameters and initial values. Click **Model Library** under Model Specification to open the library. Select a model in the list to see its formula in the **Formula** box.

Click **Show Graph** to show a 2-D theoretical curve for one-parameter models and a 3-D surface plot for two-parameter models. No graph is available for models with more than two explanatory (X) variables. Use the slider bars to change the default starting values for the parameters. You can also click the values and enter new values directly.

The **Reset** button sets the starting values for the parameters back to their default values.

Click **Show Points** to overlay the actual data points to the plot. A window opens, asking you to assign columns into X and Y roles, and an optional Group role. The Group role allows for fitting the model to every level of a categorical variable. If you specify a Group role here, also specify the same column in the Label, System ID role in the platform launch window.

For most models, the starting values are constants. Showing points enables you to adjust the parameter values to see how well the model fits for different values of the parameters.

Click **Make Formula** to create a new column in the data table. This column has the formula as a function of the specified X variable and uses the parameter values specified in the graph window.

Note: If you click **Make Formula** before you click the **Show Graph** or **Show Points** buttons, you are asked to provide the X and Y roles, and an optional Group role. After that, you are brought back to the plot so that you have the option to adjust the starting values for the parameters. Once the starting values for the parameters are satisfactory, click **Make Formula** again to create the new column.

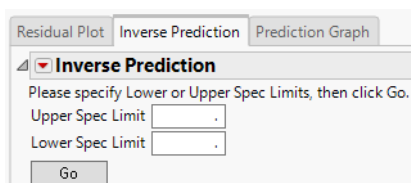
Once the formula is created in the data table, click the Degradation Data Analysis red triangle and select **Redo > Redo Analysis**. Then, return to the column selection shown in [Figure 7.10](#) on page 192. Select the column that contains the model, and then click **OK**. You are returned to the Nonlinear Path Model Specification. For more information about all the features for that specification, see [“Model Specification Details”](#) on page 192.

Note: You can customize the models included in the Nonlinear Model Library by modifying the built-in script named NonlinLib.jsl. This script is located in the Resources/Builtins folder in the folder that contains JMP (Windows) or in the Application Package (macOS).

Inverse Prediction

Use the Inverse Prediction tab to predict the time at which the Y variable reaches a specified value. These times are sometime called pseudo failure times.

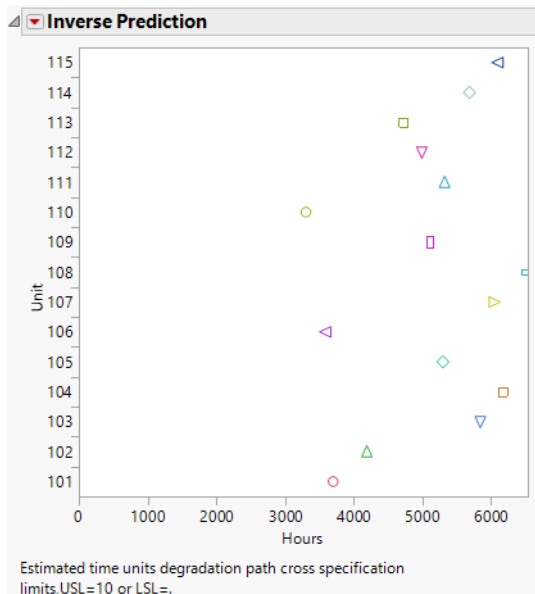
Figure 7.13 Inverse Prediction Tab



Enter either the Lower or Upper Spec Limit. Generally, if your Y variable decreases over time, then enter a Lower Spec Limit. If the Y variable increases over time, then enter an Upper Spec Limit.

For the GaAs Laser example, enter 10 for the Upper Spec Limit and click **Go**. A plot is produced showing the estimated times until the units reach a 10% increase in operating current.

Figure 7.14 Inverse Prediction Plot



The Inverse Prediction red triangle menu contains the following options:

Save Crossing Time Saves the pseudo failure times to a new data table. The table contains a Life Distribution or Fit Life by X script that can be used to fit a distribution to the pseudo failure times. When one of the Inverse Prediction Interval options is enabled, the table also includes the intervals.

Set Upper Spec Limit Sets the upper specification limit.

Set Lower Spec Limit Sets the lower specification limit.

Set Censoring Time Sets the censoring time. The plot updates to show the Censoring Time as a dotted vertical line. If Inverse Prediction Interval > No Interval is selected, observations that exceed the Censoring Time are displayed on horizontal lines starting at the Censoring Time. If Confidence Interval or Prediction Interval is selected, horizontal lines extend indefinitely to the right of observations whose upper limits exceed the Censoring Time. The Censoring Time is reflected in data tables constructed using Save Crossing Time and Generate Pseudo Failure Data.

Use Interpolation through Data Specifies the use of linear interpolation between points (instead of the fitted model) to predict when a unit crosses the specification limit. The behavior depends on whether a unit has observations that exceed the specification limit.

- If a unit has observations exceeding the specification limit, the inverse prediction is the linear interpolation between the observations that surround the specification limit.

- If a unit does not have observations exceeding the specification limit, the inverse prediction is censored and has a value equal to the maximum observed time for that unit.

Inverse Prediction Interval Specifies that confidence or prediction intervals for the pseudo failure times are shown on the Inverse Prediction plot. When intervals are enabled, the intervals are also included in the data table that is created when using the Save Crossing Time option.

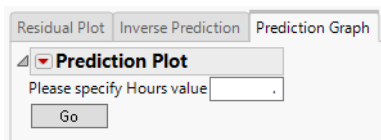
Inverse Prediction Alpha Specifies the alpha level used for the intervals.

Inverse Prediction Side Specifies one or two sided intervals.

Prediction Graph

Use the Prediction Graph tab to predict the Y variable for a specified Time value.

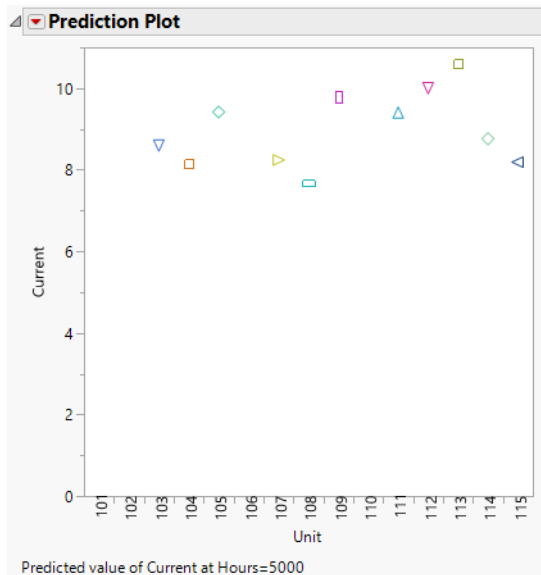
Figure 7.15 Prediction Plot Tab



The screenshot shows a software interface with three tabs: 'Residual Plot', 'Inverse Prediction', and 'Prediction Graph'. The 'Prediction Graph' tab is active. Below the tabs, there is a section titled 'Prediction Plot' with a small icon of a plot. Below this title, there is a text prompt 'Please specify Hours value' followed by a text input field. Below the input field is a 'Go' button.

For the GaAs Laser example, no data was collected after 4000 hours. If you want to predict the percent increase in operating current after 5000 hours, enter 5000 and click **Go**. A plot is produced showing the estimated percent decrease after 5000 hours for all the units.

Figure 7.16 Prediction Plot



The Prediction Plot red triangle menu contains the following options:

Save Predictions Saves the predicted Y values to a data table. When one of the Longitudinal Prediction Interval options is enabled, the table also includes the intervals.

Longitudinal Prediction Interval Shows or hides confidence or prediction intervals for the estimated Y on the Prediction Plot. When intervals are enabled, the intervals are also included in the data table that is created when using the Save Predictions option.

Longitudinal Prediction Time Specifies the time value for which you want to predict the Y.

Longitudinal Prediction Alpha Specifies the alpha level used for the intervals.

Degradation Platform Options

The Degradation Data Analysis red triangle menu contains the following options:

Path Definition The Y variable at a given time is assumed to have a distribution. You can model the mean, location parameter, or median of that distribution.

Mean Path Specifies that the mean is the parameter to be modeled.

Location Parameter Path Specifies that the location parameter of the distribution is the parameter to be modeled.

Median Path Specifies that the median of the distribution is the parameter to be modeled.

When the Location Parameter or Median Path option is selected, a menu appears in the Model Specification that enables you to select the distribution of the response.

Degradation Path Style Contains options for selecting the style of degradation path to fit.

Simple Linear Path Enables you to fit linear degradation paths and nonlinear paths that can be transformed to linear paths. See [“Simple Linear Path”](#) on page 186.

Nonlinear Path Enables you to fit nonlinear degradation paths. See [“Nonlinear Path”](#) on page 188.

Graph Options Contains options for modifying the appearance of graphs in the report.

Connect Data Markers Shows or hides lines connecting the points on the Overlay plot.

Show Fitted Lines Shows or hides the fitted lines on the Overlay plot.

Show Spec Limits Shows or hides the specification limits on the Overlay plot.

Show Residual Plot Shows or hides the residual plot.

Show Inverse Prediction Plot Shows or hides the inverse prediction plot

Show Curve Interval Shows or hides the confidence intervals on the fitted lines on the Overlay plot.

Curve Interval Alpha Enables you to change the alpha used for the confidence interval curves.

Show Median Curves Shows or hides median lines on the plot when the Path Definition is set to Location Parameter Path.

Show Legend Shows or hides a legend for the markers used on the Overlay plot.

No Tab List Shows or hides the Residual Plot, Inverse Prediction, and Prediction Graph in tabs or in stacked reports.

Prediction Settings Contains options for modifying the settings used in the model predictions.

Upper Spec Limit Specifies the upper specification limit.

Lower Spec Limit Specifies the lower specification limit.

Censoring Time Specifies the censoring time. See [“Inverse Prediction”](#) on page 196.

Baseline Specifies the normal use conditions for an X variable in nonlinear degradation paths. A path with this value usually results in an Overlay plot.

Inverse Prediction Specifies the interval type, alpha level, and one- or two-sided intervals for inverse prediction. To do inverse prediction, you must also specify the lower or upper specification limit. See [“Inverse Prediction”](#) on page 196.

Longitudinal Prediction Specifies the Time value, interval type, and alpha level for longitudinal prediction. See [“Prediction Graph”](#) on page 198.

Applications Contains options for further analysis of the degradation data.

Generate Pseudo Failure Data Creates a data table giving the predicted time each unit crosses the specification limit. The table contains a Life Distribution or Fit Life by X script that can be used to fit a distribution to the pseudo failure times.

Test Stability Performs stability analysis. See [“Stability Analysis”](#) on page 208.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Model Reports

When the **Generate Report for Current Model** button is clicked, summary reports are added in two places:

- An entry is added to the Model List report. See [“Model Lists”](#) on page 201.
- An entry is added to the Reports report. See [“Reports”](#) on page 202.

Model Lists

The Model List report gives summary statistics and other options for every fitted model. Figure 7.17 shows an example of the Model List with summaries for three models.

Figure 7.17 Model List

Model List									
Display	Model Type	Report	Nparm	-2LogLikelihood	AICc	BIC	SSE	DF	Description
☐	1 Simple Linear Path	☒	30	-179.299	-110.995	-13.0606	7.39094	225	Intercept:Different; Slope:Different; Y
☐	2 Simple Linear Path	☒	16	-115.57	-81.2841	-26.9096	9.489362	239	Intercept:Common; Slope:Different; Y
☒	3 Simple Linear Path	☒	2	756.146	760.1936	767.2286	289.6476	253	Intercept:Common; Slope:Common; Y

- Display** Select the model that you want represented in the Overlay plot, Residual Plot, Inverse Prediction plot, and Prediction Graph.
- Model Type** Gives the type of path, either linear or nonlinear.
- Report** Select the check boxes to display the report for a model. For more information about the reports, see [“Reports”](#) on page 202.
- Nparm** Gives the number of parameters estimated for the model.
- 2LogLikelihood** Gives twice the negative of the log-likelihood. See *Fitting Linear Models*.
- AICc** Gives the corrected Akaike Criterion. See *Fitting Linear Models*.
- BIC** Gives the Bayesian Information Criterion. See *Fitting Linear Models*.
- SSE** Gives the error sums-of-squares for the model.
- DF** Gives the error degrees of freedom.
- Description** Gives a description of the model.

Reports

The Reports outline node gives details about each model fit. The report includes a Model Summary report, and an Estimate report.

The Model Summary report contains the following information:

- <Y, Response> Scale** The transformation on the response variable.
- <Time> Scale** The transformation on the time variable.
- SSE** The error sums-of-squares.
- Nparm** The number of parameters estimated for the model.
- DF** The error degrees of freedom.
- RSquare** The R-square.
- MSE** The mean square error.

The Estimate report contains the following information:

Parameter The name of the parameter.

Estimate The estimate of the parameter.

Std Error The standard error of the parameter estimate.

t Ratio The t statistic for the parameter, computed as the Estimate divided by the Std Error.

Prob>|t| The p -value for a two-sided test for the parameter.

Destructive Degradation

To measure a product characteristic, sometimes the product must be destroyed. For example, when measuring breaking strength, the product is stressed until it breaks. Regular degradation analysis no longer applies in these situations. You can handle these situations in one of two ways:

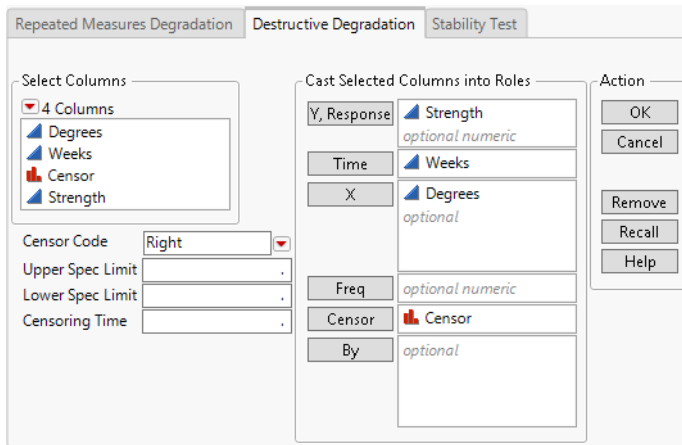
- If your failure time model is a standard one, it might be covered by the Destructive Degradation platform. See the [“Destructive Degradation”](#) chapter on page 213.
- If you are using custom transformations of the time or response variables and custom nonlinear models, use the Degradation platform. In the launch window, select the Destructive Degradation tab.

Example of Accelerated Destructive Degradation

This example fits a custom nonlinear model. The data consist of measurements on the strength of an adhesive bond. The product is stressed until the bond breaks, and the required breaking strength is recorded. Because units at normal use conditions are unlikely to break, the units were tested at several levels of an acceleration factor (temperature). You want to estimate the strength (in newtons) of units after 52 weeks (one year) at use conditions of 25°C.

Complete the Launch Window

1. Select **Help > Sample Data Library** and open Reliability/Adhesive Bond.jmp.
2. Select **Analyze > Reliability and Survival > Degradation**.
3. Select the **Destructive Degradation** tab.
4. Select Strength and click **Y, Response**.
5. Select Weeks and click **Time**.
6. Select Degrees and click **X**.
7. Select Censor and click **Censor**.

Figure 7.18 Completed Launch Window


8. Click **OK**.

Define and Fit the Model

1. Select **Lognormal** from the menu in the Location Parameter Path Specification panel in the Model Specification outline.

This specifies a lognormal transformation for the response, Strength.

2. Click the Degradation Data Analysis red triangle and select **Degradation Path Style > Nonlinear Path**.

This adds a script window to the report. You insert a script that specifies your model in this window.

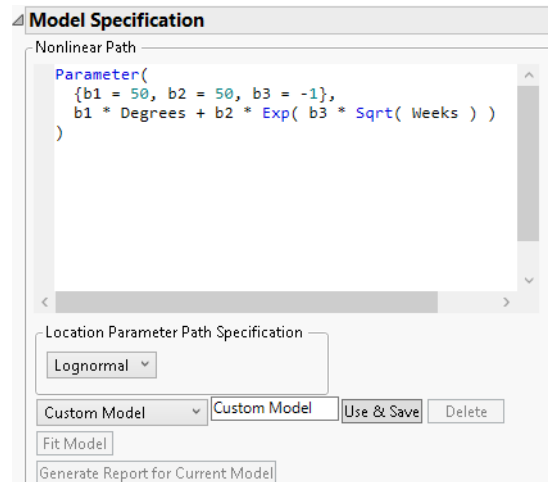
3. Copy the JSL formula below and paste it into the script window under Nonlinear Path:

```
Parameter(
  {b1 = 50, b2 = 50, b3 = -1},
  b1 * :Degrees + b2 * Exp( b3 * Sqrt( :Weeks ) )
)
```

The script defines a model for Strength in terms of parameters b1, b2, and b3. The script specifies initial values for the parameters.

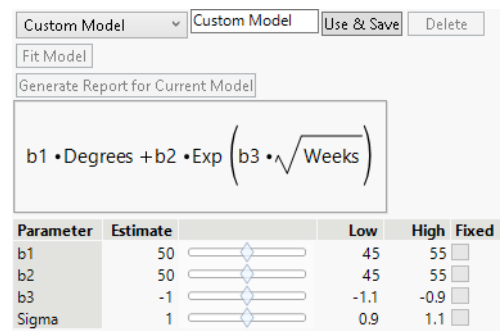
4. In the text box at the bottom of the report, change Empty to Custom Model.

Figure 7.19 Nonlinear Path Script



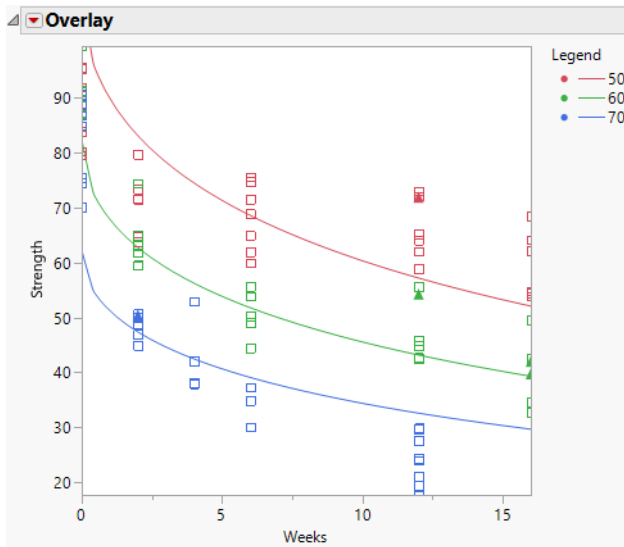
5. Click **Use & Save**.

Figure 7.20 Updated Model Specification Outline



The report is updated to include controls for changing initial values for the parameters. Initial values for the parameters are set to the values specified in the formula in the script editor.

6. Click **Fit Model**.

Figure 7.21 Plot of Fitted Model


The Parameter panel in the Model Specification outline is updated to show the parameter estimates for the fitted model. The model fit is shown in the Overlay plot. You can drag the axes to show the points, as is done in Figure 7.21. The legend identifies the curves.

Obtain Predicted Values and Prediction Intervals

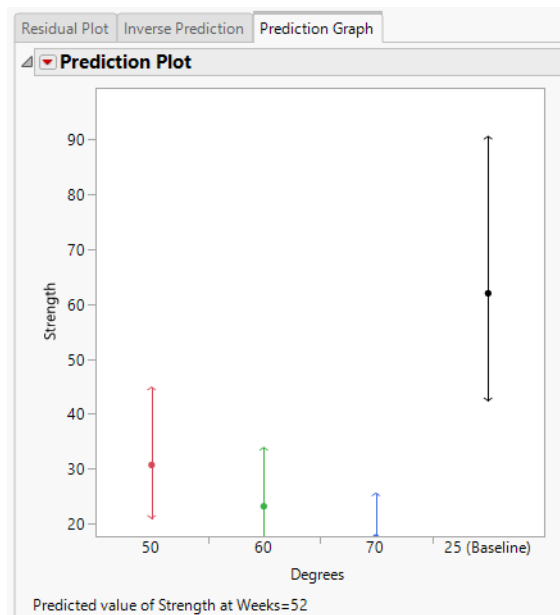
Next, you obtain a predicted value and prediction interval for Strength after 52 weeks at the baseline use condition of 25°C.

1. Click the Degradation Data Analysis red triangle and select **Prediction Settings**.
2. In the **Prediction Settings** window:
 - Enter 25 for Baseline.
 - Enter 52 for Time in the Longitudinal Prediction panel.
 - Select Prediction Interval from the menu in the Longitudinal Prediction panel.

Figure 7.22 Prediction Settings

3. Click **OK**.
4. Select the **Prediction Graph** tab in the report.

Figure 7.23 Prediction Plot at Weeks = 52



The Prediction Plot shows the predicted value and its 95% level prediction interval above the axis label of 25 (Baseline).

5. Click the Prediction Plot red triangle and select **Save Predictions**.

Predicted values for the three values of **Degrees** and for the desired baseline of 25 degrees are saved to a data table. The predicted strength of the adhesive bond after 1 year at 25° C is 61.96, with a prediction interval of 42.42 to 90.50.

Stability Analysis

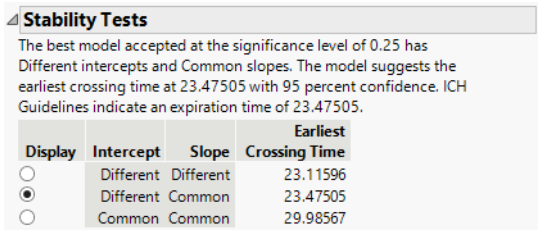
Stability analysis is used in setting pharmaceutical product expiration dates. Three linear degradation models are fit, and an expiration date is estimated following International Conference on Harmonisation (ICH) guidelines. The ICH guidelines are used for the general framework of determining if batches can be pooled for expiration dating (ICH Q1E 2003). For specific implementation details, see the STAB macro and FDA guidelines in Chow (2007, Appendix B).

The Stability Tests report summarizes three degradation models and their corresponding earliest crossing times. The best model is selected and displayed in the Overlay plot. The models are listed in the order of complexity.

The first model fits different intercepts and different slopes to each batch. (In the procedure steps and model comparisons described below, this is the *full model*.) When this model is used for estimating the expiration date, the mean square error (MSE) is not pooled across batches. Prediction intervals are computed for each batch using individual mean squared errors, and the interval that crosses the specification limit first is used to estimate the expiration date. The earliest crossing times are based on 95% two-sided prediction intervals when there are two specification limits provided; the earliest crossing times are based on 95% one-sided prediction intervals when there is only one specification limit provided.

The second model fits different intercepts to each batch, but fits a common slope across all batches. The third model fits a common intercept and a common slope across all batches. When this model is appropriate, it provides the expiration date furthest into the future.

Figure 7.24 Stability Tests Summary



The Model Comparisons section of the Stability Tests report summarizes the tests of significance for each of the stability models. The Legend describes each source. The procedure for determining the best model for expiration dating considers the *p*-values for Sources C and B. The sources are listed below in reverse order to accommodate the order of steps in the procedure.

Tip: The ICH guidelines specify using a significance level of 0.25 to determine the appropriate model. You can follow the steps below and compare the p -values at each source to a different significance level. If this results in a different model being selected, you can use the radio buttons in the Display column of the table in the Stability Tests summary report.

Source E Specifies the sum of squared responses minus the Source D sum of squares. The value in the Mean Square column for Source E is the Source E SS value divided by the Source E degrees of freedom, which are equal to the number of parameters in the full model (different intercepts and different slopes).

Source D Specifies the error sum of squares and corresponding MSE for the full model (different intercepts and different slopes).

Source C Specifies the test of equal slopes. This is a test of the second model (different intercepts and common slope) versus the full model (different intercepts and different slopes).

- If the p -value is less than 0.25, the slopes are assumed to be different across batches. The procedure stops and the full model (different intercepts and different slopes) is used to estimate the expiration date.
- If the p -value is greater than or equal to 0.25, the slopes are assumed to be common across batches and you then evaluate the intercepts using Source B.

Source B Specifies the test for equal intercepts. This is a test of the third model (common intercepts and common slopes) versus the second model (different intercepts and common slope).

- If the p -value is less than 0.25, the intercepts are assumed to be different across batches, and the second model (different intercepts and common slope) is used to estimate the expiration date.
- If the p -value is greater than or equal to 0.25, the intercepts are assumed to be common across batches, and the third model (common intercepts and common slopes) is used to estimate the expiration date.

Source A Specifies the test of the third model (common intercepts and common slopes) versus the full model (different intercepts and different slopes). This test is not used in the procedure to determine the best model for expiration dating.

Figure 7.25 Stability Model Comparisons

Model Comparisons					
Source	DF	SS	Mean Square	F Statistic	Prob>F
A	6	61.48956	10.24826	10.10836	<.0001*
B	3	60.48866	20.16289	19.88764	<.0001*
C	3	1.000906	0.333635	0.329081	0.8043
D	28	28.38752	1.01384		
E	8	360214.5	45026.81		

Legend					
Source		Intercept	Slope	Intercept	Slope
A		Different	Different	Common	Common
B		Different	Common	Common	Common
C		Different	Different	Different	Common
D	Residual				
E	Whole Model				

The Reports section contains models fit by the Test Stability option. The following four models are fit by the Test Stability option:

- A different intercept, different slope model where the MSE (mean squared error) is pooled across batches. (This is an alternative to the first model in the Summary report.)
- A different intercept, common slope model. (This is the second model in Summary report.)
- A common intercept, common slope model. (This is the third model in Summary report.)
- A different intercept, different slope model where the MSE (mean squared error) is not pooled across batches. (This is the first model in the Summary report.)

Caution: In addition to the four models fit by the Test Stability option, the Reports section also contains any models fit in the Degradation platform prior to running the Test Stability option.

Example

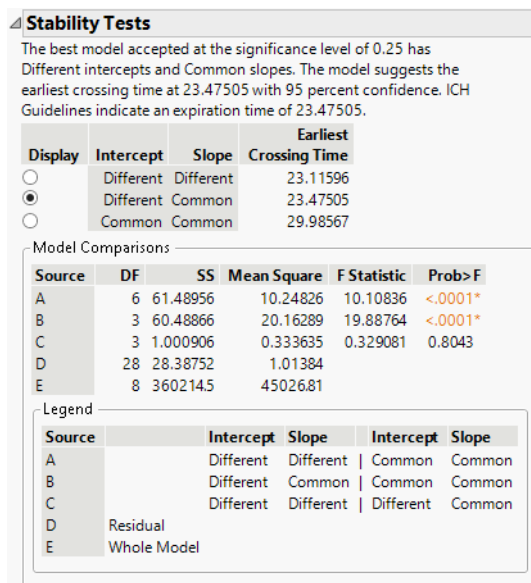
Use the data in the *Stability.jmp* sample data table to establish an expiration data for a new product. The data consists of product concentration measurements on four batches. A concentration of 95 is considered the end of the product's usefulness.

To perform the stability analysis, do the following steps:

1. Select **Help > Sample Data Library** and open *Reliability/Stability.jmp*.
2. Select **Analyze > Reliability and Survival > Degradation**.
3. Select the **Stability Test** tab.
4. Select Concentration (mg/Kg) and click **Y, Response**.
5. Select Time and click **Time**.
6. Select Batch Number and click **Label, System ID**.
7. Enter 95 for the Lower Spec Limit.

8. Click **OK**.

Figure 7.26 Stability Models



The test for equal slopes has a p -value of 0.8043. Because this is larger than a significance level of 0.25, the test is not rejected, and you conclude the degradation slopes are equal between batches.

The test for equal intercepts and slopes has a p -value of <.0001. Because this is smaller than a significance level of 0.25, the test is rejected, and you conclude that the intercepts are different between batches.

Because the test for equal slopes was not rejected, and the test for equal intercepts was rejected, the chosen model is the one with Different Intercepts and Common Slope. This model is the one selected in the report, and gives an estimated expiration date of 23.475.

Chapter 8

Destructive Degradation Model Product Deterioration over Time

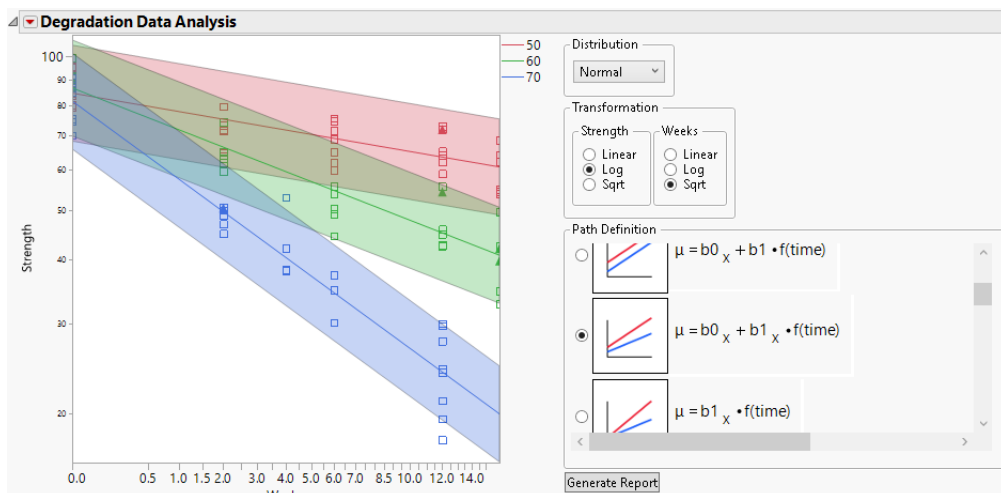
To measure a product characteristic, sometimes the product must be destroyed. For example, when measuring breaking strength, the product is stressed until it breaks. Because the test is destructive, there is only one observation per product unit. In such a situation, you can model product reliability using the Destructive Degradation platform.

The platform models how a (typically) nonnegative response changes over time. Observations are assumed to be independent and measure the value of the response and the time at failure. A large and flexible collection of pre-defined models is provided. The models include location-scale and log-location-scale distributions whose location parameters are functions of time. The models allow explanatory variables and additional parameters. When an explanatory variable is specified, the platform fits an accelerated destructive degradation model.

Note: If you require a model that is not represented among the models provided, you can use the Degradation platform. See “Destructive Degradation” on page 203 in the “Degradation” chapter.

For more information about destructive degradation and reliability, see Escobar et al. (2003) and Meeker and Escobar (1998).

Figure 8.1 Destructive Degradation Example of Model for Adhesive Bond.jmp



Contents

Example of the Destructive Degradation Platform 215

Launch the Destructive Degradation Platform..... 221

 Specify Two Y Columns 222

The Destructive Degradation Plot Options and Models 223

 Plot Options 224

 Models 225

The Destructive Degradation Report 229

 Model List..... 230

 Model Outlines 230

Destructive Degradation Platform Options 233

Statistical Details for the Destructive Degradation Platform 233

Example of the Destructive Degradation Platform

This example of an accelerated destructive degradation model is patterned after an example from Escobar et al. (2003). The data consist of measurements on the strength (measured in newtons) of an adhesive bond. Temperature is considered to be an acceleration factor. The product is stressed until the bond breaks and the required breaking stress is recorded. Because units at normal temperatures are unlikely to break, the units were tested at several levels over a wide range of temperatures. Strength less than 50 newtons is considered failure. You want to estimate the proportion of units with a strength below 50 newtons after 260 weeks (5 years) at a reference temperature of 35 degrees Celsius.

This example has three stages:

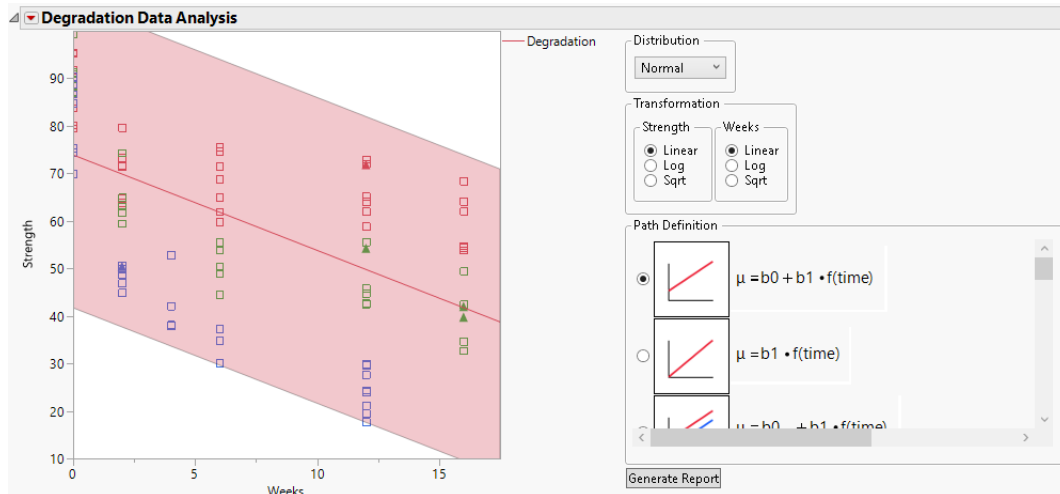
- “Perform the Initial Analysis” on page 215
- “Change the Model and Generate the Report” on page 216
- “Using the Profilers for Prediction” on page 219

Perform the Initial Analysis

1. Select **Help > Sample Data Library** and open Reliability/Adhesive Bond.jmp.
2. Select **Analyze > Reliability and Survival > Destructive Degradation**.
3. Select Strength and click **Y, Response**.
4. Select Weeks and click **Time**.
5. Select Degrees and click **X**.

The temperature is the accelerating factor in the experiment.

6. Select Censor and click **Censor**.
Notice that the **Censor Code** is set to Right.
7. Click **OK**.

Figure 8.2 Initial Degradation Plot


The platform specifies a default model. The default model assumes that the data are described by a single Normal distribution, whose location parameter is a linear function of time.

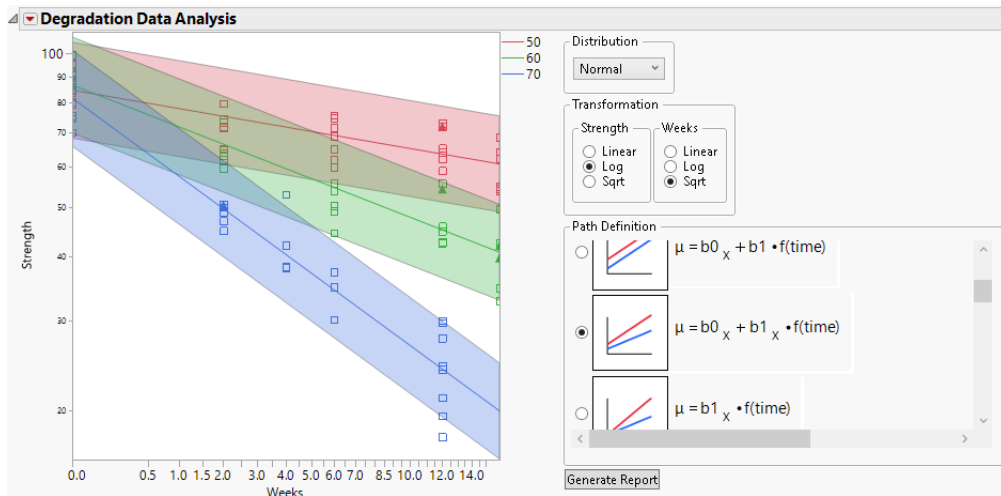
Change the Model and Generate the Report

1. Select **Log** for the Y (Strength) Transformation.
2. Select **Sqrt** for the Time (Weeks) Transformation.
3. Select $\mu = b0_x + b1_x * f(\text{time})$ for the Path Definition.

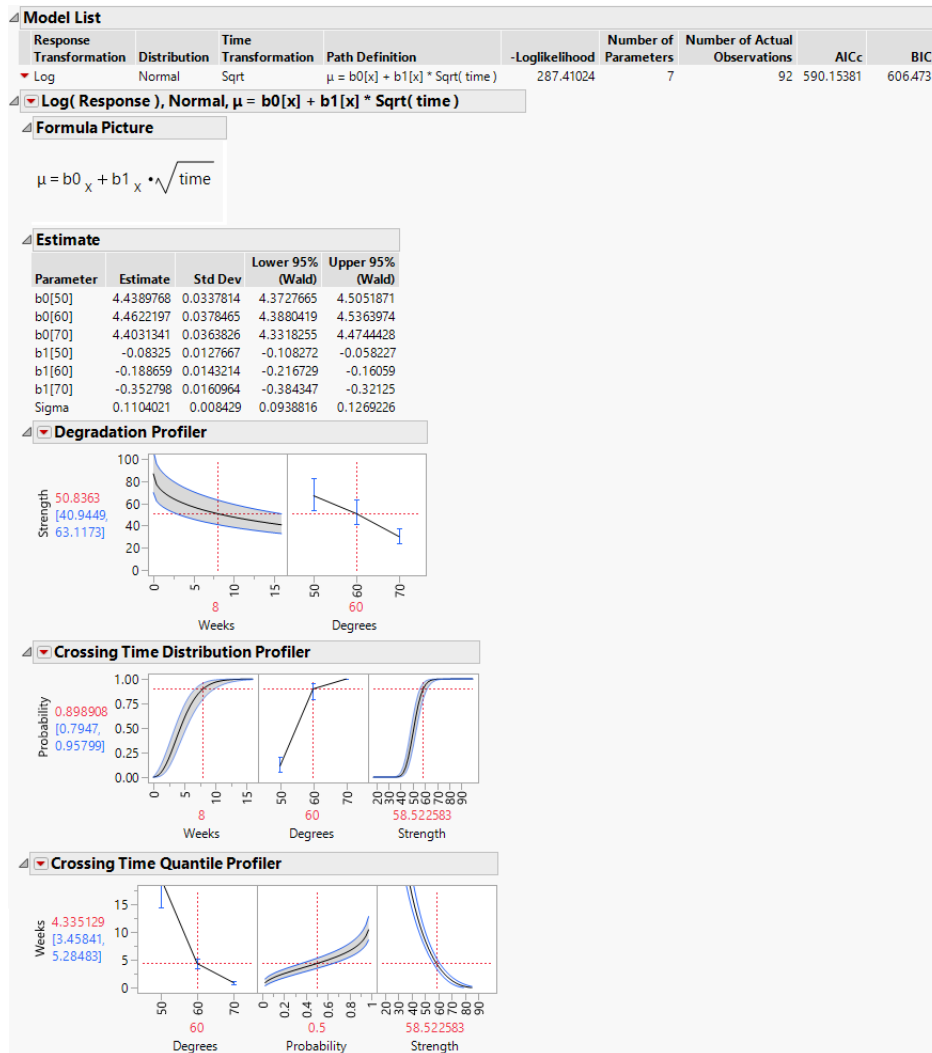
The subscript “x” denotes the accelerating variable, which is Degrees in this example.

Note: This model is linear in all parameters.

Figure 8.3 Plot Showing Model



4. Click **Generate Report**.

Figure 8.4 Report for Basic Model


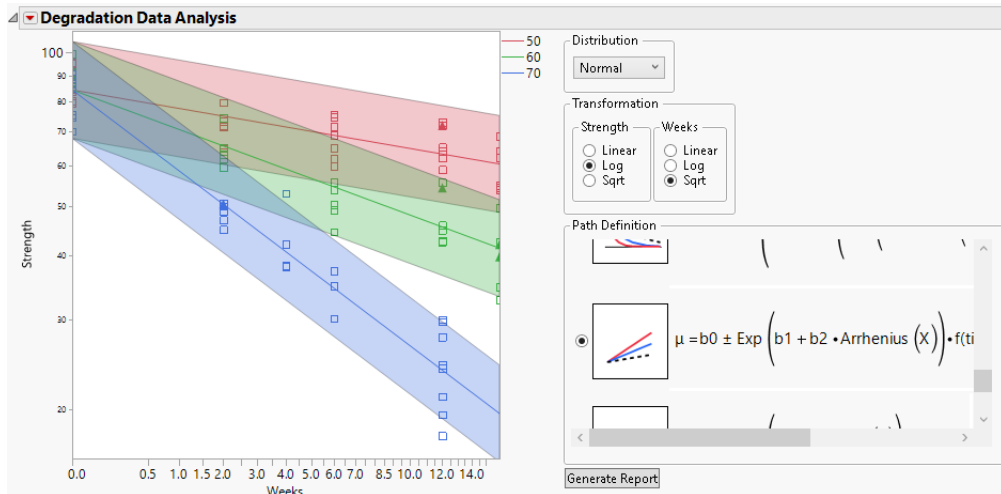
The estimates of the slope $b1$ at the three values of Degrees suggest that degradation occurs more quickly at higher temperatures. Failure mechanisms that depend on chemical processes are often well modeled using the Arrhenius model for temperature. For this reason, you now fit a model where an Arrhenius transformation is applied to Degrees, which is measured on a Celsius scale.

5. Select $\mu = b0 \pm \text{Exp}(b1 + b2 * \text{Arrhenius}(X)) * f(\text{time})$ for the Path Definition.

Note: This model is not linear in the parameters.

6. Select **Celsius** and click **OK**.

Figure 8.5 Plot Showing Model with Arrhenius Transformation



7. Click **Generate Report**.

Figure 8.6 Report Including Second Model with Arrhenius Transformation

Model List									
Response Transformation	Distribution	Time Transformation	Path Definition	-Loglikelihood	Number of Parameters	Number of Actual Observations	AICc	BIC	
▼ Log	Normal	Sqrt	$\mu = b_0[x] + b_1[x] \cdot \text{Sqrt}(\text{time})$	287.41024	7	92	590.15381	606.473	
▼ Log	Normal	Sqrt	$\mu = b_0 - \exp(b_1 + b_2 \cdot \text{Arrhenius}(X)) \cdot \text{Sqrt}(\text{time})$	288.09559	4	92	584.65095	594.27834	
▼ Log(Response), Normal, $\mu = b_0[x] + b_1[x] \cdot \text{Sqrt}(\text{time})$									
▼ Log(Response), Normal, $\mu = b_0 - \exp(b_1 + b_2 \cdot \text{Arrhenius}(X)) \cdot \text{Sqrt}(\text{time})$									
Formula Picture									
$\mu = b_0 - \exp\left(b_1 + b_2 \cdot \text{Arrhenius}(X)\right) \cdot \sqrt{\text{time}}$									
Estimate									
Parameter	Estimate	Std Dev	Lower 95% (Wald)	Upper 95% (Wald)					
b0	4.4354686	0.0205347	4.3952214	4.4757158					
b1	22.825503	1.4029332	20.075805	25.575202					
b2	-0.704728	0.0414658	-0.786	-0.623457					
Sigma	0.1114956	0.0085011	0.0948337	0.1281574					

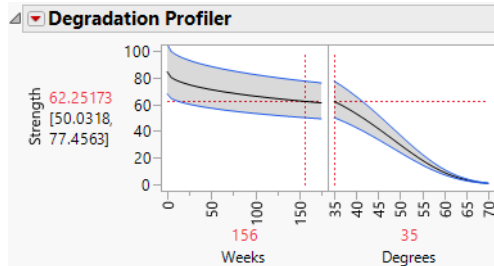
Using the Profilers for Prediction

Because the Arrhenius model shows a better fit, as indicated by its smaller AICc and BIC values (Figure 8.6), you continue your analysis using this model.

Recall that strength less than 50 newtons is considered failure. You are interested in units lasting 156 weeks (three years) at a reference temperature of 35 degrees Celsius. Change the settings in the profilers to reflect these values. Click the value in red beneath each plot's horizontal axis and enter the new value.

1. In the Degradation profiler for the Arrhenius model, set **Weeks** to 156 and **Degrees** to 35.

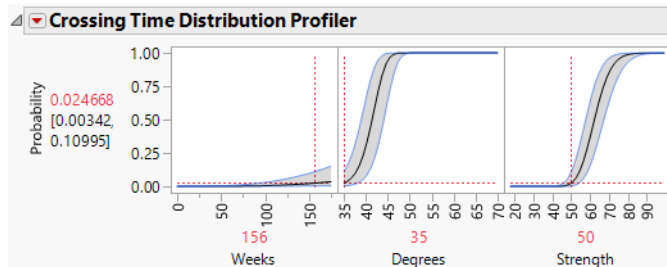
Figure 8.7 Degradation Profiler



The predicted **Strength** at these settings is 62.25173, with a 95% prediction interval ranging from 50.0318 to 77.4563. Failures are not very likely at these or less extreme settings.

2. In the Crossing Time Distribution Profiler, set **Weeks** to 156, **Degrees** to 35, and **Strength** to 50.

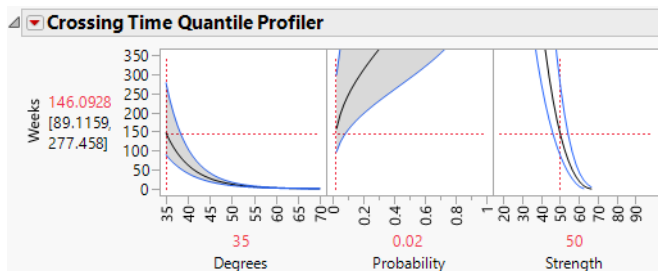
Figure 8.8 Crossing Time Distribution Profiler



At 156 weeks at a temperature of 35 degrees Celsius, the probability that the value of Strength is less than 50 is 0.024668. The 95% confidence interval ranges from 0.00342 to 0.10995. The probability of failure at these or less extreme conditions is about 2%.

3. In the Crossing Time Quantile Profiler, set **Degrees** to 35, **Probability** to 0.02, and **Strength** to 50.
4. Adjust the vertical axis of the Crossing Time Quantile Profiler so that the maximum value is about 350.

Figure 8.9 Crossing Time Quantile Profiler

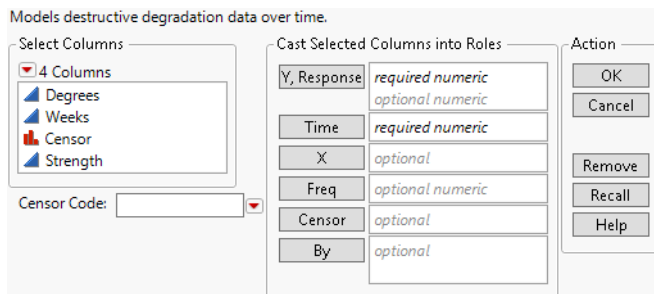


The number of weeks within which 2% of units fail at a temperature of 35 degrees Celsius is estimated to be 146.0928. The 95% confidence interval ranges from 89.1159 to 277.458.

Launch the Destructive Degradation Platform

Launch the Destructive Degradation platform by selecting **Analyze > Reliability and Survival > Destructive Degradation**. Figure 8.10 shows the Destructive Degradation launch window using the Adhesive Bond.jmp data table (located in the Reliability folder).

Figure 8.10 The Destructive Degradation Launch Window



For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The Destructive Degradation launch window contains the following options:

Y, Response Identifies the column containing the degradation measurements. When your response values are interval-censored, you can enter two columns. See [“Specify Two Y Columns”](#) on page 222.

Time Identifies the column containing the time values.

X Identifies an optional explanatory variable. If the distribution of Y changes not only over time, but is also impacted by some other variable, that additional variable can be supplied

as X. Use this role to specify the accelerating factor in an accelerated destructive degradation model.

Freq Identifies frequencies or observation counts when there are multiple units. If the value is 0 or a positive integer, then the value represents the frequencies or counts of observations with the given row's settings.

Censor Identifies an optional column that identifies censored response measurements. When there is only one column in the Y role, this column indicates whether the response Y in a given row is exact or right censored. A right-censored observation is one where the exact measurement is unknown, but is known to be larger than the Y value in the corresponding row.

Select the value that identifies right-censored observations from the Censor Code menu. Rows corresponding to other values in the Censor column are treated as uncensored. Rows with missing censor code values are excluded from the analysis. JMP attempts to detect the censor code and display it in the list.

By Identifies an optional By variable. A separate analysis is produced for each level of this variable.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

Specify Two Y Columns

You can specify two Y columns when some of the degradation measurements are interval censored or left censored. For a given row, the values in the two Y columns determine the type of censoring.

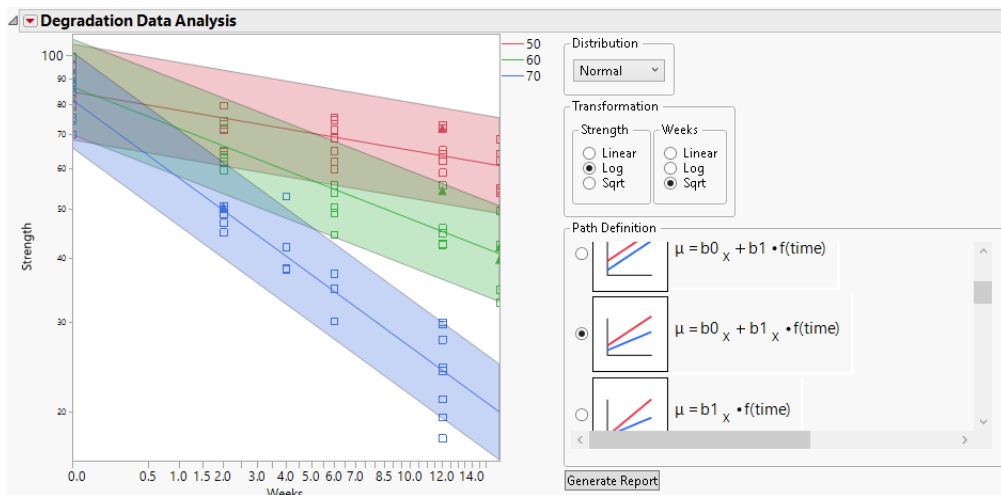
- If the two Y values are equal and neither is missing, then the common measurement is treated as exact.
- If the two Y values are not equal and neither is missing, then the measurement is interval censored and assumed to be between the two values.
- If only the first value is missing, then the measurement is left censored and assumed to be smaller than the second value.
- If only the second value is missing, then the measurement is right censored and assumed to be larger than the first value.

Note: The only way to fit left-censored measurements in the Destructive Degradation platform is through the use of two Y columns.

The Destructive Degradation Plot Options and Models

The Degradation Data Analysis plot shows the data and a graphical representation of the model that is currently specified based on selections of Distribution, Transformation, and Path Definition. The plot shown in Figure 8.11 is for the Adhesive Bond.jmp data table and represents a model that includes an optional X variable.

Figure 8.11 Destructive Degradation Plot and Options



Note the following about the plot:

- Data values are represented by markers.

Table 8.1 Marker Descriptions

Icon	Description
	An exact measurement.

Table 8.1 Marker Descriptions (*Continued*)

Icon	Description
▼	Left-censored observation, indicating that the censored measurement is below the triangle. Note: In the Degradation platform, left-censoring arises only when observations are interval censored. See “Specify Two Y Columns” on page 222.
▲	Right-censored observation, indicating that the censored measurement is above the triangle.
▼ ▲	Interval-censored observation, indicating that the censored value is within the specified interval.

Note: By default, the markers are not colored. To color them to match the model color scheme, select **Rows > Color or Mark by Column**. Select the X column. Deselect **Continuous Scale**. Select **JMP Default** from the Colors menu.

- For each level of X, a colored band appears. If the model does not include an X variable, then a single band appears. For a given value of Time, the upper and lower bounds of the band are the 0.025 and 0.975 percentiles of the fitted distribution of Y. The colors of the bands correspond to the values of X, as indicated by the legend to the upper right of the plot.

Note: Marker colors correspond to the color states assigned in the data table.

- The solid curve in the center of a band is the median of the fitted distribution of Y for the corresponding value of X over time. If the model does not include an X variable, then the curve plots the median of Y over time.

Plot Options

Distribution Choose a location-scale or a log-location-scale distribution.

Note: It is not recommended to fit a log-location-scale distribution model when you specify a Log transformation for the response column.

Transformation Choose a transformation function for the response Y and for the Time variable.

Note: If you apply the Log transformation to a column that contains nonpositive values, the rows with nonpositive values are omitted from the model fit. If you apply the Sqrt transformation to a column that contains negative values, the rows with negative values are omitted from the model fit.

Path Definition Choose a linear or a nonlinear path for the regression model. For more information about each model, see [“Models”](#) on page 225.

Generate Report Creates a report for the specified model. The first time you select Generate Report, a Model List outline is created. When you select Generate Report to fit other models, the Model List outline is updated and an outline is added for each model.

Models

The following table provides the equations for each model in the Path Definition list. For a description of each model, follow the link.

Note: The thumbnail sketch shown to the left of each equation shows a generic plot of the behavior of the location parameter, μ , over time. In the report’s main plot, the plot of the estimated median can differ from the thumbnail based on your selections for Distribution and Transformation.

Table 8.2 Model Equations

Model	Equation
Common Path with Intercept	$\mu = b_0 + b_1 * f(\text{time})$
Common Path without Intercept	$\mu = b_1 * f(\text{time})$
Common Slope	$\mu = b_{0X} + b_1 * f(\text{time})$
Individual Path with Intercept	$\mu = b_{0X} + b_{1X} * f(\text{time})$
Individual Path without Intercept	$\mu = b_{1X} * f(\text{time})$
Common Intercept	$\mu = b_0 + b_{1X} * f(\text{time})$

Table 8.2 Model Equations (*Continued*)

Model	Equation
First-Order Kinetics Type 1	$\mu = b_0 - b_1 * \text{Exp}[-b_2 * \text{Exp}[b_3 * [\text{Arrhenius}(X_0) - \text{Arrhenius}(X)]] * f(\text{time})]$
First-Order Kinetics Type 2	$\mu = b_0 * [1 - \text{Exp}[-b_1 * \text{Exp}[b_2 * [\text{Arrhenius}(X_0) - \text{Arrhenius}(X)]] * f(\text{time})]$
First-Order Kinetics Type 3	$\mu = b_0 + b_1 * \text{Exp}[-b_2 * \text{Exp}[b_3 * [\text{Arrhenius}(X_0) - \text{Arrhenius}(X)]] * f(\text{time})]$
First-Order Kinetics Type 4	$\mu = b_0 * \text{Exp}[-b_1 * \text{Exp}[b_2 * [\text{Arrhenius}(X_0) - \text{Arrhenius}(X)]] * f(\text{time})]$
Arrhenius Rate	$\mu = b_0 \pm \text{Exp}[b_1 + b_2 * \text{Arrhenius}(X)] * f(\text{time})$
Polynomial Rate	$\mu = b_0 \pm \text{Exp}[b_1 + b_2 * \text{Log}(X)] * f(\text{time})$
Exponential Rate	$\mu = b_0 \pm \text{Exp}[b_1 + b_2 * X] * f(\text{time})$

Models That Are Linear in Transformed Time

Common Path with Intercept

This model fits a single distribution whose location parameter changes linearly over transformed time. This model fits a common intercept and a common slope, regardless of whether there is an X variable.

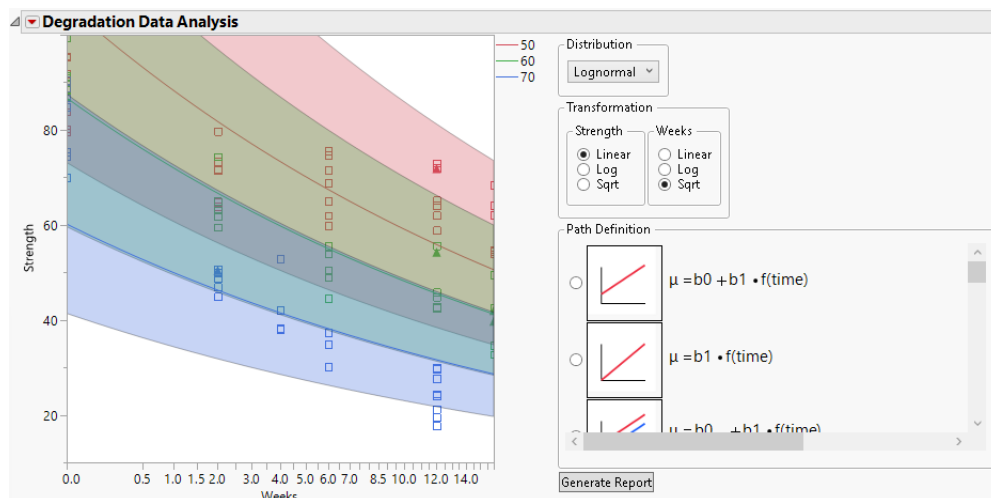
Common Path without Intercept

This model fits a single distribution whose location parameter changes linearly over transformed time but whose value at time zero is zero. Based on selections for Distribution and Transformation, the median curve might not be a straight line and might not pass through origin. This model fits a zero intercept and a common slope, regardless of whether there is an X variable.

Common Slope

In this model, the location parameters are linear functions of transformed time with separate intercepts for the values of X but a common slope. Based on selections for Distribution and Transformation, the model fits might appear as curves. For example, selecting a Lognormal distribution gives the plot in Figure 8.12.

Figure 8.12 Common Slope Model Using a Lognormal Distribution



Individual Path with Intercept

In this model, the location parameters are linear functions of transformed time with separate intercepts and separate slopes for the values of X .

Individual Path without Intercept

In this model, the location parameters are linear functions of transformed time with zero intercepts and separate slopes for the values of X .

Common Intercept

In this model, the location parameters are linear functions of transformed time with a common intercept and separate slopes for the values of X .

First-Order Kinetics Models

Four first-order kinetics models where the location parameter is a nonlinear function based on an Arrhenius transformation of temperature are provided. Each of these location models fit separate models for each value of the optional explanatory variable X .

When you first select any of these models, you must specify the measurement scale for temperature and a reference temperature value, X_0 . The specified reference temperature affects the interpretation of b_2 in these models. The b_2 parameter is the rate constant at X_0 . The value for X_0 is used to construct a time acceleration factor (Meeker and Escobar 1998). If you subsequently select another of the first-order kinetics models or the Arrhenius Rate model (see “Arrhenius Rate” on page 229), the platform remembers and uses these specifications.

First-Order Kinetics Type 1

In this model, $b1$ and $b2$ are positive. On a linear scale, the curves have an upper asymptote at $b0$ as time approaches infinity.

First-Order Kinetics Type 2

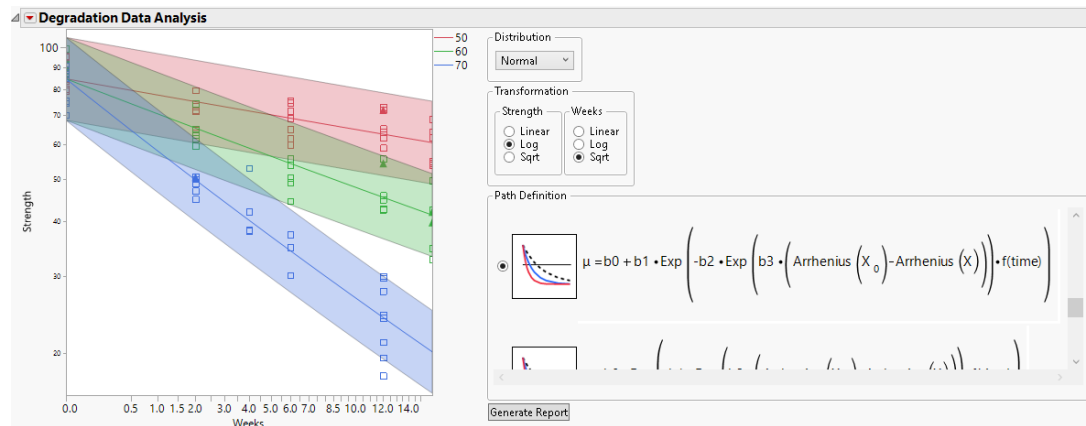
In this model, the location parameter is zero at time zero. Both $b0$ and $b1$ are positive. On a linear scale, the curves have an upper asymptote at $b0$ as time approaches infinity. You can think of the Type 2 model as a vertically shifted version of the Type 1 model.

First-Order Kinetics Type 3

In this model, both $b1$ and $b2$ are positive. Because of the sign preceding $b1$ is the opposite of the sign preceding $b1$ in the Type 1 model, this model is an inverted version of the Type 1 model. On a linear scale, it has a lower asymptote at $b0$ as time approaches infinity.

Given data exhibiting a negative slope over time, the fitted model can produce a plot similar to Figure 8.13. The figure is for Adhesive Bond.jmp. The selected temperature measurement scale is Celsius and the specified reference temperature under typical use conditions is 35 degrees.

Figure 8.13 Example of First-Order Kinetics Model Type 3



First-Order Kinetics Type 4

This model is a vertically shifted version of the Type 3 model. On a linear scale, the curves have a lower asymptote at 0 as time approaches infinity.

Three models where the location parameter is an exponential function of the transformed X variable are provided. Each of these location models fits common intercept and separate slope models for each value of the optional explanatory variable X. For each of these models, on a linear scale, the location parameter is linear in the transformed time values.

Arrhenius Rate

This model involves an exponential function of the Arrhenius transformation multiplied by transformed time. When you select this model, you are asked to specify the measurement scale for temperature, unless you have already supplied this information.

Polynomial Rate

This model involves an exponential function of a linear function of the log of X multiplied by transformed time.

Exponential Rate

This model involves an exponential function of a linear function of X multiplied by transformed time.

The Destructive Degradation Report

The Destructive Degradation report contains an outline for each model that you fit. When you fit a model, the Model List is updated with a row for that model.

Note: All models are fit using the maximum likelihood method.

Figure 8.14 Model List and Model Outlines

Model List								
Response Transformation	Distribution	Time Transformation	Path Definition	-Loglikelihood	Number of Parameters	Number of Actual Observations	AICc	BIC
Log	Normal	Sqrt	$\mu = b_0[x] + b_1[x] * \text{Sqrt}(\text{time})$	287.41024	7	92	590.15381	606.473
Log	Normal	Sqrt	$\mu = b_0 - \text{Exp}(b_1 + b_2 * \text{Arrhenius}(X)) * \text{Sqrt}(\text{time})$	288.09559	4	92	584.65095	594.27834
Log(Response), Normal, $\mu = b_0[x] + b_1[x] * \text{Sqrt}(\text{time})$ Log(Response), Normal, $\mu = b_0 - \text{Exp}(b_1 + b_2 * \text{Arrhenius}(X)) * \text{Sqrt}(\text{time})$								

Model List

The first four columns in the Model List reflect the choices that you made in the plot options. The -Loglikelihood, AICc, and BIC statistics are information-based measures that can be used for model comparisons. For descriptions of these measures, see *Fitting Linear Models*.

The three information-based measures in the Model List are comparable across models as long as the models being compared have the same Number of Actual Observations. If this is not the case, exercise caution because different models might use different subsets. The Number of Actual Observations might be reduced due to the choice of distribution or the choice of transformation. Choosing a log-location-scale distribution excludes all non-positive Y values. Also, the Log and Sqrt transformations exclude all non-positive values.

Each row of the Model List table has a red triangle menu with the following options:

- Scroll To** Scrolls the report window to the corresponding model outline.
- Remove** Removes the model from the Model List and removes the corresponding model outline from the report.

Model Outlines

The outline for each model contains a red triangle menu with the following option:

- Remove** Removes the model outline from the report and removes the model from the Model List.

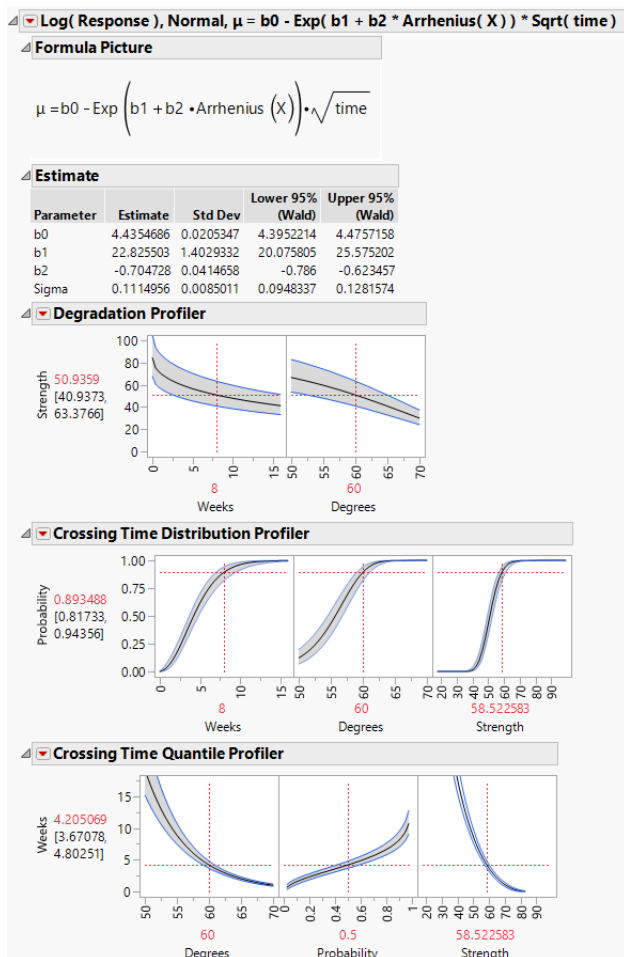
The outline for each model contains the following reports:

- Formula Picture** Shows the equation for the location parameter.
- Estimate** Shows parameter estimates and their standard deviations. This report also contains 95% Wald-based confidence intervals for the parameters.

Note: If you specify a reference temperature (X_0), this value is also displayed at the top of the Estimate report.

- Distribution, Quantile, and Inverse Prediction** For more information about the profilers, see [“Profilers”](#) on page 231.
- Residual Plots** For more information about the residual plots, see [“Residual Plots”](#) on page 232.

Figure 8.15 Reports within a Model Outline



Profilers

Three profilers appear in the model report.

Degradation Profiler

The Degradation Profiler shows a profiler view of the Degradation Data Analysis plot for the given model. The response is the degradation response Y. The profiler includes a plot against the Time variable and a plot against the optional explanatory variable X (if you have specified one). The plot against Time shows the median of the fitted distribution of Y as a solid curve. The dashed curves show the 0.025 and 0.975 percentiles of the fitted distribution of Y.

Crossing Time Distribution Profiler

Use the Crossing Time Distribution Profiler to determine the probability that the degradation measurement falls below a given threshold at some point in time.

The profiler plots the estimated cumulative distribution function of the response Y as a function of Time, the optional X variable, and Y. The plots for Time and Y show Wald confidence intervals.

Crossing Time Quantile Profiler

Use the Crossing Time Quantile Profiler to determine the time at which a specified proportion of measurements falls below a given threshold value.

The profiler plots the estimated Time as a function of the optional X variable, quantile values for Y (Probability), and Y. The plots for Probability and Y show Wald confidence intervals.

Residual Plots

Four residual plots appear in the model report. Use these plots to validate the distributional assumption for the model. For more information about the standardized residuals, see [“Statistical Details for the Destructive Degradation Platform”](#) on page 233.

Cox-Snell Residual P-P Plot

If the points deviate far from the diagonal, then the distributional assumption might be violated. The Cox-Snell Residual P-P Plot red triangle menu has an option called Save Residuals that enables you to save the residual data to the data table. See Meeker and Escobar (1998, sec. 17.6.1) for a discussion of Cox-Snell residuals.

Probability Plot of Standardized Residuals

If the points deviate far from the diagonal, then the distributional assumption might be violated. The Probability Plot of Standardized Residuals red triangle menu has an option called Save Residuals that enables you to save the residual data to the data table.

Standardized Residuals versus Time

Use the Standardized Residuals versus Time plot to examine differences in variability over time. The Standardized Residuals versus Time red triangle menu has an option called Save Residuals that enables you to save the residual data to the data table.

Standardized Residuals versus Predicted

Use the Standardized Residuals versus Predicted plot to examine differences in variability over the range of predicted values. The Standardized Residuals versus Predicted red triangle menu has an option called Save Residuals that enables you to save the residual data to the data table.

Destructive Degradation Platform Options

The Degradation Data Analysis red triangle menu contains the following options:

Graph Options Provides options for modifying the platform graphs.

Shade Turns the shading on or off. By default, the upper and lower bounds of each shaded band at a given value of time correspond to the 0.025 and 0.975 quantiles of the distribution of Y .

Shade Coverage If shading is on, you can enter a proportion to increase or decrease the amount of shading coverage.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Statistical Details for the Destructive Degradation Platform

The destructive degradation model can be expressed as follows:

$$g(Y) \sim F(\mu, \sigma)$$

$$\mu = h(f(\text{Time}), X)$$

where:

- $g(Y)$ is the transformed Y variable
- F is the selected probability distribution
- μ is the location parameter, defined by h
- h is a function that relates the transformed Time variable and the explanatory variable X
- σ is the scale parameter of the distribution
- $f(\text{Time})$ is the transformed Time variable
- X is an optional explanatory variable

The standardized residuals are obtained as follows:

- For location-scale distributions, the standardized residuals are $(y-\mu)/\sigma$.
- For log-location-scale distributions, the standardized residuals are $(\log(y)-\mu)/\sigma$.

Chapter 9

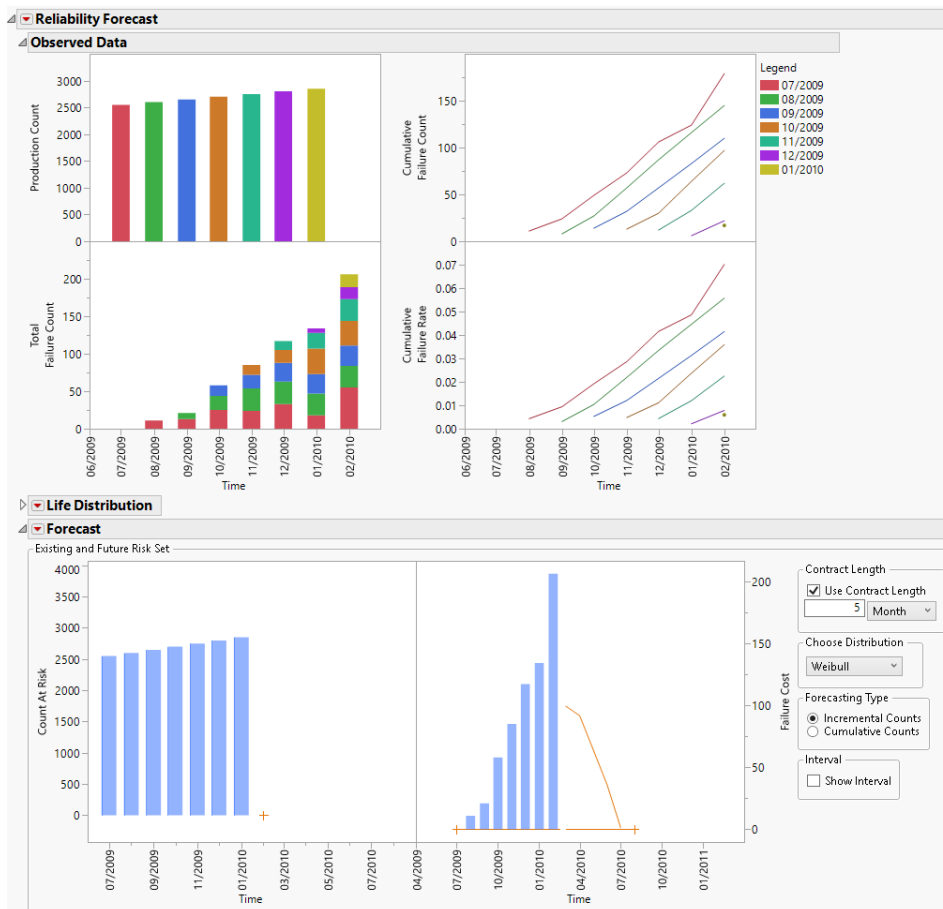
Reliability Forecast

Forecast Product Failure Using Production and Failure Data

The Reliability Forecast platform helps you predict the number of future failures. JMP estimates the parameters for a life distribution using production dates, failure dates, and production volume.

Using the interactive graphs, you can adjust factors such as future production volumes and contract length to estimate future failures. Repair costs can be incorporated into the analysis to forecast the total cost of repairs across all failed units.

Figure 9.1 Example of a Reliability Forecast



Contents

Overview of the Reliability Forecast Platform 237

Example Using the Reliability Forecast Platform..... 237

Launch the Reliability Forecast Platform..... 241

The Reliability Forecast Report..... 244

 Observed Data Report 244

 Life Distribution Report..... 244

 Forecast Report 245

Reliability Forecast Platform Options 249

Additional Example Using the Reliability Forecast Platform..... 250

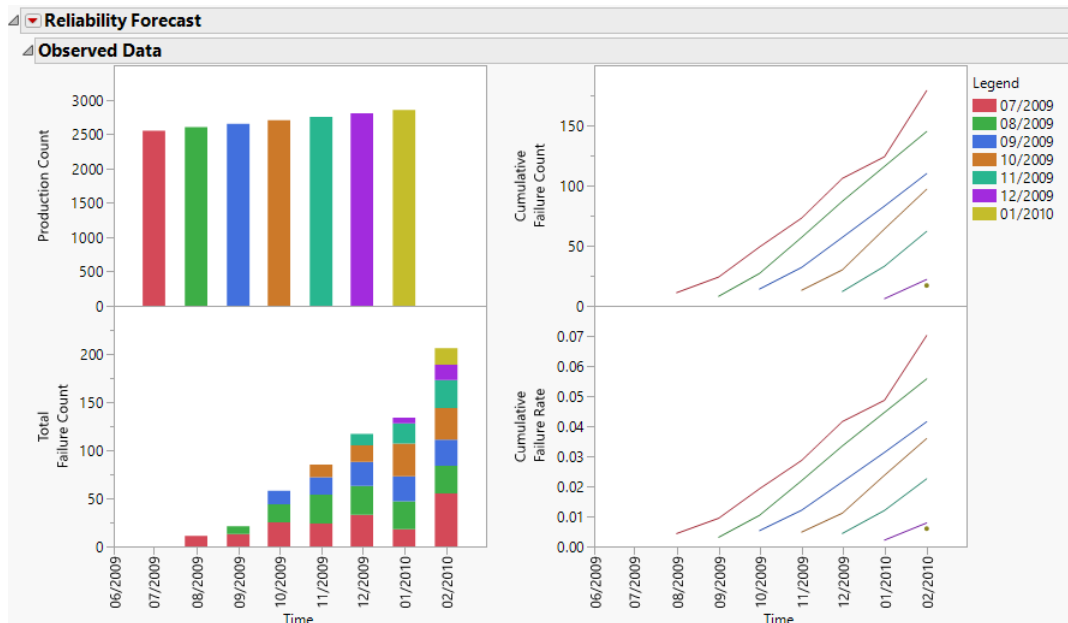
Overview of the Reliability Forecast Platform

The Reliability Forecast platform helps you predict the number of future failures. JMP estimates the parameters for a life distribution using production dates, failure dates, and production volume. Using the interactive graphs, you can adjust factors such as future production volumes and contract length to estimate future failures. Repair costs can be incorporated into the analysis to forecast the total cost of repairs across all failed units.

Example Using the Reliability Forecast Platform

You have data on seven months of production and returns. You want to use this information to forecast the total number of units that will be returned for repair through February 2011. The product has a 12-month contract.

1. Select **Help > Sample Data Library** and open Reliability/Small Production.jmp.
2. Select **Analyze > Reliability and Survival > Reliability Forecast**.
3. On the **Nevada Format** tab, select Sold Quantity and click **Production Count**.
4. Select Sold Month and click **Timestamp**.
5. Select the other columns and click **Failure Count**.
6. Click **OK**.

Figure 9.2 Observed Data Report


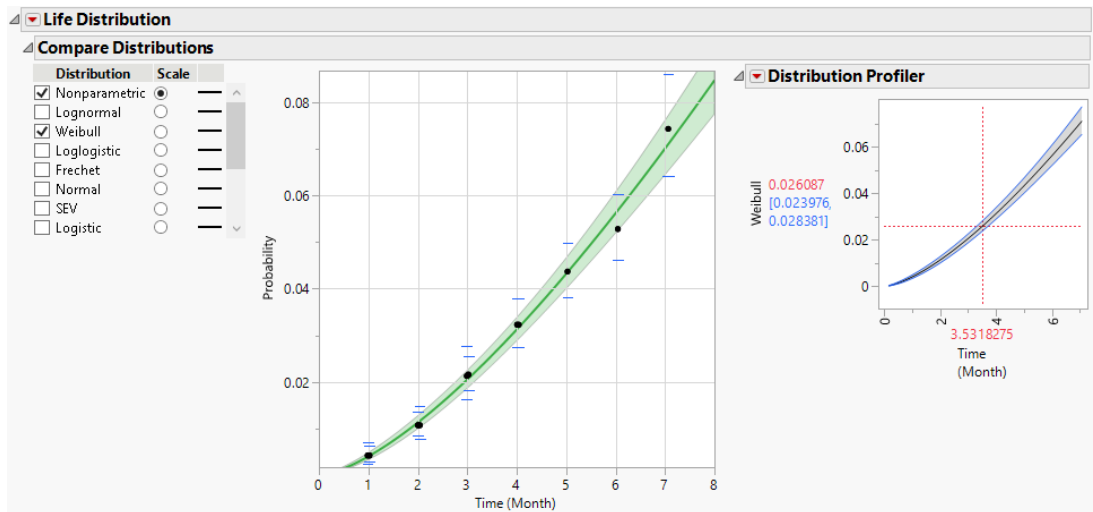
In the Observed Data report, the bottom left shows bar charts of previous failures.

Cumulative failures are shown in the line graphs on the right. Note that production levels are fairly consistent. As production accumulates over time, more units are at risk of failure, so the cumulative failure rate gradually increases. The consistent production levels also result in similar cumulative failure rates and counts from month to month.

7. Click the **Life Distribution** disclosure icon.

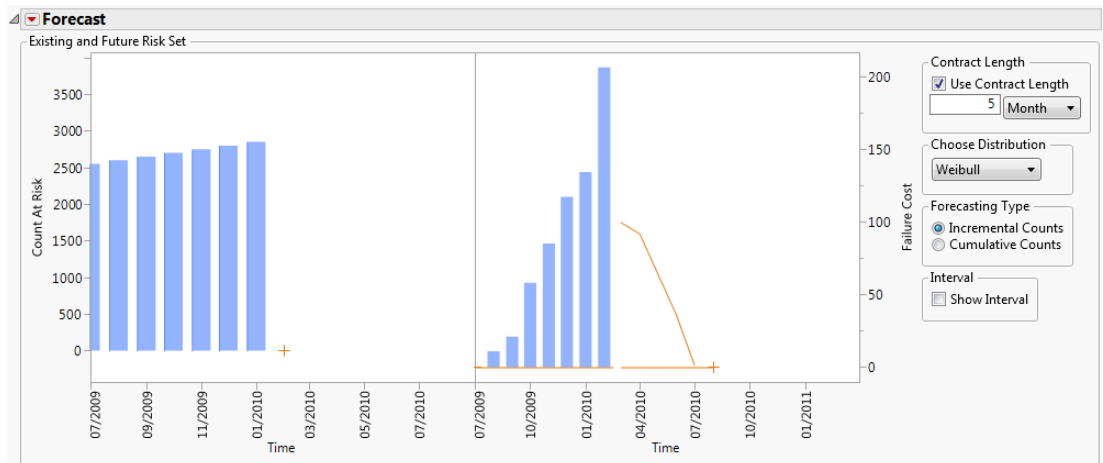
JMP fits production and failure data with a Weibull distribution using the Life Distribution platform (Figure 9.3). JMP then uses the fitted Weibull distribution to forecast returns for the next five months (Figure 9.4).

Figure 9.3 Life Distribution Report



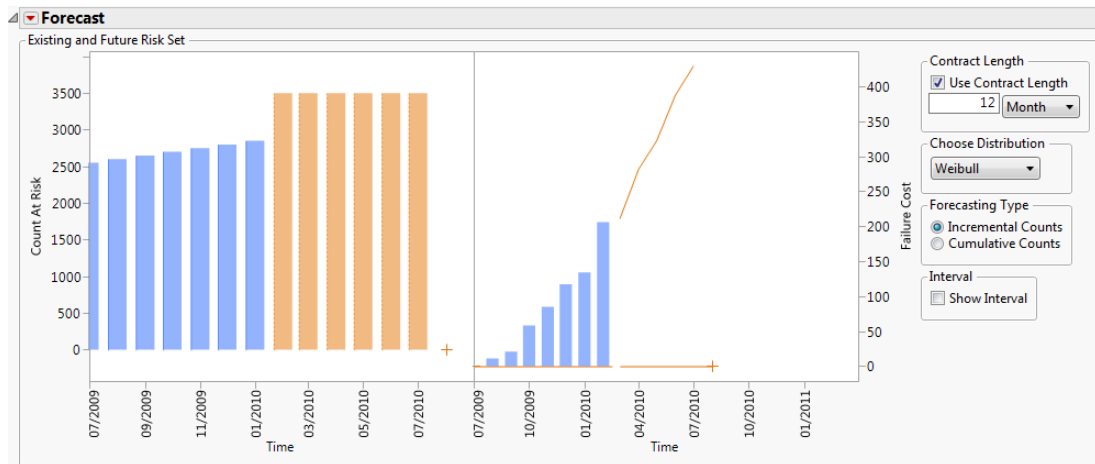
The Forecast report shows previous production in the left graph (Figure 9.4). In the right graph, you see that the number of previous failures increased steadily over time.

Figure 9.4 Forecast Report



8. In the Forecast report, type 12 for the Contract Length.
9. In the left Forecast graph, drag the animated hotspot over to July 2010 and upward to approximately 3500.

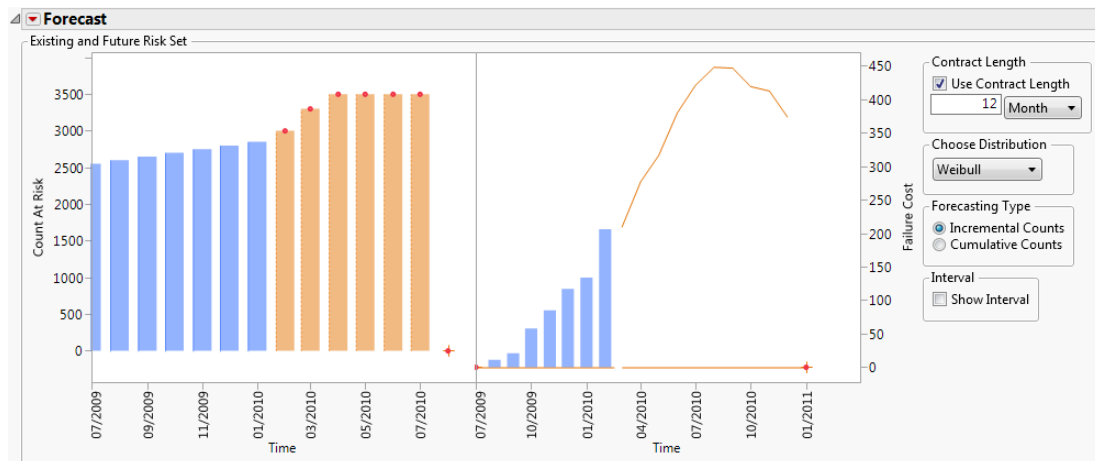
Orange bars appear in the left graph to represent future production. The monthly returned failures in the right graph increase gradually through August 2010.

Figure 9.5 Production and Failure Estimates


10. In the left graph, drag the February 2010 hotspot to approximately 3000 and then drag the March 2010 hotspot to approximately 3300.

11. In the right graph, drag the right hotspot to February 2011.

JMP estimates that the number of returns will gradually increase through August 2010 and decrease by February 2011.

Figure 9.6 Future Production Counts and Forecasted Failures


Launch the Reliability Forecast Platform

Launch the Reliability Forecast platform by selecting **Analyze > Reliability and Survival > Reliability Forecast**.

Figure 9.7 The Reliability Forecast Launch Window

For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The launch window includes a tab for each contract data format: Nevada, Dates, and Time to Event (for time to failure data). The following sections describe these formats.

Nevada Format

Contract data is commonly stored in the Nevada format: shipment or production dates and failure counts within specified periods are shaped like the state of Nevada. Figure 9.8 shows the Small Production.jmp sample data table.

Figure 9.8 Example of the Nevada Format

		Sold Quantity	Sold Month	08/2009	09/2009	10/2009	11/2009	12/2009	01/2010	02/2010
1		2550	07/2009	11	13	25	24	33	18	55
2		2600	08/2009	0	8	19	30	30	29	29
3		2650	09/2009	0	0	14	18	25	26	27
4		2700	10/2009	0	0	0	13	17	34	33
5		2750	11/2009	0	0	0	0	12	21	29
6		2800	12/2009	0	0	0	0	0	6	16
7		2850	01/2010	0	0	0	0	0	0	17

The Nevada Format tab contains the following options:

- Interval Censored Failure** Considers the returned quantity to be interval censored. The interval is between the last recorded time and the time that the failure was observed. This option is selected by default.
- Life Time Unit** Physical date-time format of all time stamps, including column titles for return counts. The platform uses this setting in forecasting step increments.
- Production Count** Column that identifies the number of units produced.
- Timestamp** Column that contains the production date.
- Failure Count** Column that contains the number of failed units.
- Group ID** Column by which observations are grouped. Each group has its own distribution fit and forecast. A combined forecast is also included.

Dates Format

The Dates format focuses on production and failure dates. One data table specifies the production counts for each time period. The other table provides failure dates, failure counts, and the corresponding production times of the failures.

Figure 9.9 shows the Small Production part1.jmp and Small Production part2.jmp sample data tables.

Figure 9.9 Example of the Dates Format

production data

	Sold Quantity	Sold Month
1	2550	07/2009
2	2600	08/2009
3	2650	09/2009
4	2700	10/2009
5	2750	11/2009
6	2800	12/2009
7	2850	01/2010

failure data

	Return Month	Return Quantity	Sold Month
1	08/2009	11	07/2009
2	09/2009	13	07/2009
3	10/2009	25	07/2009
4	11/2009	24	07/2009
5	12/2009	33	07/2009
6	01/2010	18	07/2009
7	02/2010	55	07/2009

The Dates Format tab is divided into Production Data and Failure Data sections.

Production Data

Select Table Select the table that contains the number of units and the production dates.

Failure Data

Select Table Select the table that contains failure data, such as the number of failed units, production dates, and failure dates.

Left Censor Column that identifies censored observations.

Timestamp Column that links production observations to failure observations, indicating which batch a failed unit came from.

Censor Code Value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

Other options are identical to those on the Nevada Format tab. See [“Nevada Format”](#) on page 241.

Time to Event Format

The Time to Event format shows production and failure data. Unlike the Nevada and Dates formats, Time to Event data does not include date-time information, such as production or failure dates. This format can also accommodate arbitrary censoring schemes in the data. See [“Defining Risk Sets with Time to Event Data”](#) on page 247.

Figure 9.10 shows the Small Production Time to Event.jmp sample data table. For an example that uses time-to-event data, see [“Additional Example Using the Reliability Forecast Platform”](#) on page 250.

Figure 9.10 Example of the Time to Event Format

	start time	end time	failure counts	
	Time (Month)	Time Right	Freq	Timestamp
1	0	1.0184804928	11	07/2009
2	1.0184804928	2.0369609856	13	07/2009
3	2.0369609856	3.022587269	25	07/2009
4	3.022587269	4.0410677618	24	07/2009
5	4.0410677618	5.0266940452	33	07/2009
6	5.0266940452	6.045174538	18	07/2009
7	6.045174538	7.0636550308	55	07/2009
8	7.0636550308	•	2371	07/2009

The Time to Event Format tab contains the following options:

Forecast Start Time Time at which the forecast begins. Enter the first value that you want on the horizontal axis. To enable forecasting, you must select Numeric for the Life Time Unit and set the Forecast Start Time to zero.

Censor Code Code for censored observations. Available only when you assign a **Censor** variable.

Other options are identical to those on the Nevada Format tab. See [“Nevada Format”](#) on page 241.

The Reliability Forecast Report

The Reliability Forecast report contains the Observed Data report, Life Distribution report, and Forecast report. Use these reports to view the current data, compare distributions to find the right fit, and then adjust factors that affect forecasting. You can also save the forecast to a data table for use in financial forecasts. See [Figure 9.1 “Example of a Reliability Forecast”](#) on page 235 for an example of the report.

The Reliability Forecast red triangle menu contains options for filtering the observed data by date and saving the data in another format. See [“Reliability Forecast Platform Options”](#) on page 249.

Observed Data Report

The Observed Data report gives you a quick view of Nevada and Dates data ([Figure 9.11](#)).

- Bar charts show the production and failure counts for the specified production periods.
- Line charts show the cumulative forecast by production period.

Note: Time-to-event data does not include date-time information, such as production or failure dates, so the platform does not create an Observed Data report for this format.

Life Distribution Report

The Life Distribution report lets you compare distributions and work with profilers to find the right fit. When you select a distribution in the Forecast report, the Life Distribution report is updated. For more information about this report, see the [“Life Distribution”](#) chapter on page 29. (Some of the options described in the Life Distribution chapter are not available in the Life Distribution report in the Reliability Forecast platform.)

Forecast Report

The Forecast report provides interactive graphs that help you forecast failures. By dragging hotspots, you can add anticipated production counts and see how they affect the forecast.

Adjusting Risk Sets

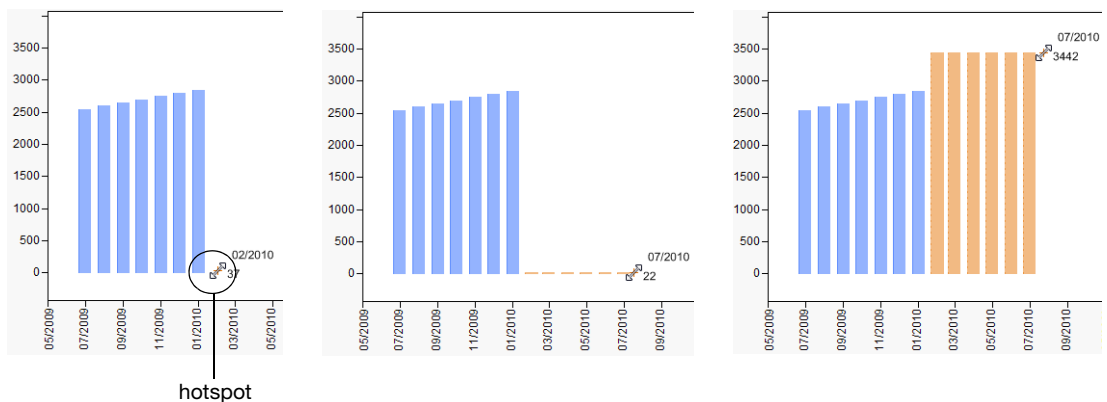
Future Production

In the left graph, the blue bars represent previous production counts. To add anticipated production, follow these steps:

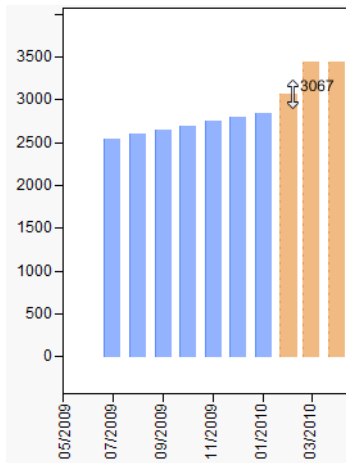
1. Drag a hotspot to the right to add one or more production periods.

The orange bars represent future production.

Figure 9.11 Add Production Periods



2. Drag each bar upward or downward to change the production count for each period.

Figure 9.12 Adjust Production Counts

Tip: To adjust future production counts, click the Forecast red triangle and select **Spreadsheet Configuration of Risk Sets**. See [“Spreadsheet Configuration of Risk Sets”](#) on page 248.

Existing Production

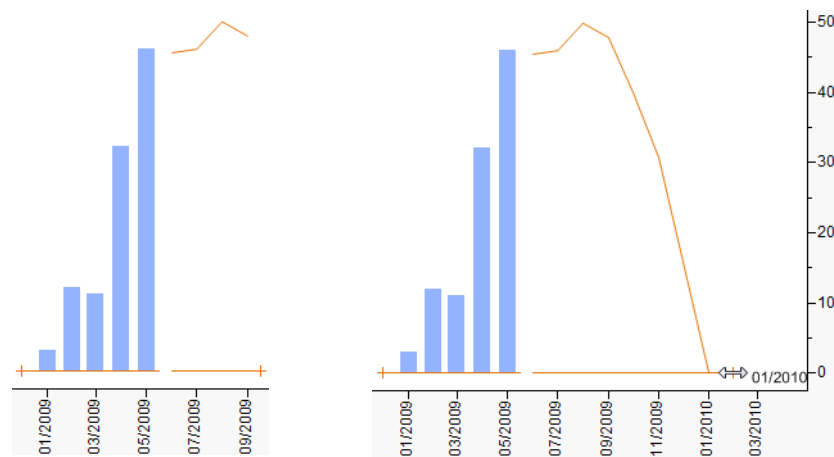
To remove a risk set from the forecast results, right-click a blue bar and select **Exclude**. To return that data to the risk set, right-click and select **Include**.

Tip: To adjust existing production counts, click the Forecast red triangle and select **Spreadsheet Configuration of Risk Sets**. See [“Spreadsheet Configuration of Risk Sets”](#) on page 248.

Forecasting Failures

When you adjust production in the left graph, the right graph is updated to estimate future failures (Figure 9.13). Dragging a hotspot lets you change the forecast period. The orange line then shortens or lengthens to show the estimated failure counts.

Figure 9.13 Adjust the Forecast Period



Defining Risk Sets with Time to Event Data

You can obtain forecasts for arbitrary risk sets using the Time to Event Format tab. In the launch window's Time to Event Format tab, you must select Numeric for the Life Time Unit and set the Forecast Start Time to zero. Enter appropriate columns for Time to Event, Censor, Freq, and Group ID.

The plot on the left provides locations for bars representing existing production. Drag the hotspot at 0 to the left to create existing risk sets. Drag the hotspot at 1 to the right to create future production risk sets.

You can specify existing and future risk sets using the Spreadsheet Configuration of Risk Sets option. For time-to-event data, you must enter negative time values for the existing risk set in the Future Risk area. See ["Spreadsheet Configuration of Risk Sets"](#) on page 248.

For an example that uses time-to-event data, see ["Additional Example Using the Reliability Forecast Platform"](#) on page 250.

Forecast Graph Options

To explore your data further, you can change the contract length, distribution type, and other options.

- To forecast failures for a different contract period, enter the number next to **Use Contract Length**. Change the time unit if necessary.
- To change the distribution fit, select a distribution from the **Choose Distribution** list. The distribution is then fit to the future graph of future risk. The distribution fit appears in the Life Distribution report plot, and a new profiler is added. Changing the distribution fit in the Life Distribution report does not change the fit in the Forecast graph.

- If you are more interested in the total number of failures over time, select **Cumulative Counts**. Otherwise, JMP shows failures incrementally, which can make trends easier to identify.
- To show 95% confidence limits for the anticipated number of failures, select **Show Interval**.

Forecast Report Options

The Forecast red triangle menu contains the following options:

Animation Controls the flashing of the hotspots. You can also right-click a blue bar in the existing risk set and select or deselect **Animation**.

Interactive Configuration of Risk Sets Determines whether you can drag hotspots in the graphs.

Spreadsheet Configuration of Risk Sets Lets you specify production counts and periods (rather than adding them to the interactive graphs). You can also exclude production periods from the analysis.

- To remove an existing time period from analysis, highlight the period in the Existing Risk area, click, and then select **Exclude**. Or select **Include** to return the period to the forecast.
- To edit production, double-click in the appropriate Future Risk field and enter the new values.
- To add a production period to the forecast, right-click in the Future Risk area and select one of the **Append** options. (**Append Rows** adds one row; **Append N Rows** lets you specify the number of rows.)

As you change these values, the graphs update accordingly.

Note: If you launched the platform using the Time to Event Format tab, values for the existing risk set must be entered in the Future Risk area using negative time values.

Import Future Risk Set Lets you import future production data from another open data table. The new predictions then appear in the future risk graph. The imported data must have a column for timestamps and for quantities.

Show Interval Shows or hides 95% confidence limits in the graph. This option works the same as selecting **Show Interval** next to the graphs.

After you select **Show Interval**, the **Forecast Interval Type** option appears in the menu. Select one of the following interval types:

Plugin Interval Considers only forecasting errors given a fixed distribution.

Prediction Interval Considers forecasting errors when a distribution is estimated with estimation errors (for example, with a non-fixed distribution).

If you select Prediction Interval, the **Prediction Interval Settings** option appears in the menu. Approximate intervals are initially shown in the graph. Select **Monte Carlo Sample Size** or **Random Seed** to specify those values instead. To use the system clock, enter a missing number.

Use Contract Length Determines whether the specified contract length is considered in the forecast. This option works the same as selecting **Use Contract Length** next to the graphs.

Use Failure Cost Shows failure cost instead of the failure count in the future risk graph.

After you select **Use Failure Cost**, the **Set Failure Cost** option appears in the menu. This option enables you to set a cost for each failure. If you specified a Group variable in the launch window, the Set Failure Cost window enables you to specify separate costs for failures in each group.

Save Forecast Data Table Saves the cumulative and incremental number of returns in a new data table, along with the variables that you selected in the launch window. For grouped analyses, table names include the group ID and the word “Aggregated”. Existing returns are also included in the aggregated data tables.

Reliability Forecast Platform Options

The Reliability Forecast red triangle menu contains the following options:

Save Data in Time to Event Format Saves Nevada or Dates data in a Time to Event formatted table.

Show Legend Shows or hides a legend for the Observed Data report. Unavailable for Time to Event data.

Show Graph Filter Shows or hides the Graph Filter so that you can select which production periods to show in the Observed Data graphs. Bars fade for deselected periods. Deselect the periods to show the graph in its original state. Unavailable for Time to Event data.

See *Using JMP* for more information about the following options:

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Additional Example Using the Reliability Forecast Platform

You have seven months of production and returns data that are stored in time-to-event format. You want to use this information to forecast the total number of units that will be returned for repair in the next 6 months. The product has a 4-month contract.

1. Select **Help > Sample Data Library** and open Reliability/Small Production Time to Event.jmp.
2. Select **Analyze > Reliability and Survival > Reliability Forecast**.
3. On the **Time to Event Format** tab, select Time (Month) and Time Right. Click **Time to Event**.
4. Select Freq and click **Freq**.
5. Click **OK**.
6. In the Forecast report, type 4 for the Contract Length.
7. (Optional.) Click the **Life Distribution** disclosure icon.

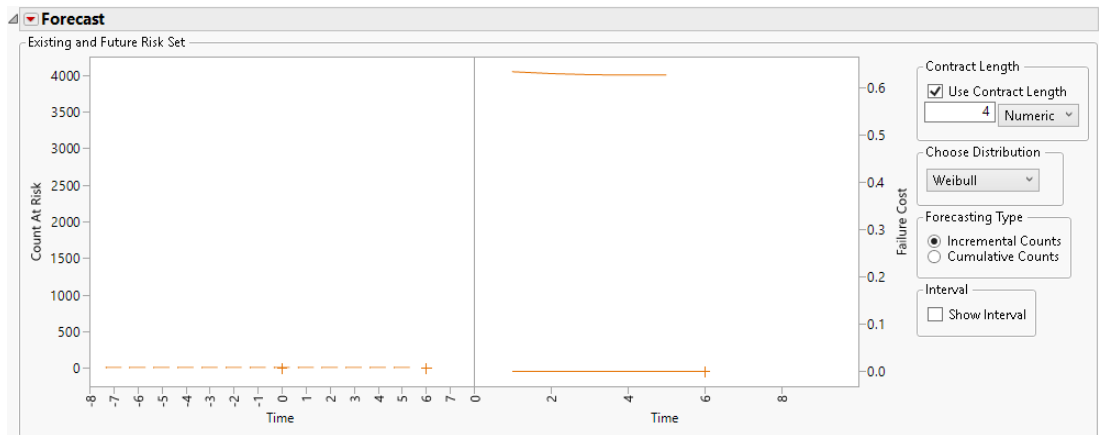
The Life Distribution report shows a Weibull fit for the production and failure data. The Reliability Forecast platform then uses the fitted Weibull distribution to forecast returns.

Specify the Current and Future Production Counts

8. Set the vertical axis settings to prepare for past and future production counts. Right-click the vertical axis in the left Forecast graph and select **Axis Settings**. Specify the following axis settings:
 1. Type -250 for the Minimum.
 2. Type 4250 for the Maximum.
 3. Type 500 for the Increment.
 4. Click **OK**.
9. Set the horizontal axis to cover the previous 8 and next 8 months. Right-click the horizontal axis in the left Forecast graph and select **Axis Settings**. Specify the following axis settings:
 1. Type -8 for the Minimum.
 2. Type 8 for the Maximum.
 3. Click **OK**.

10. Set the past 7 months of production counts. In the left Forecast graph, drag the leftmost animated hotspot left to -7.
11. Set the future 5 months of production estimates. In the left Forecast graph, drag the rightmost animated hotspot right to 5.

Figure 9.14 Risk Set after Dragging Animated Hotspots



12. In the Forecast red triangle menu, select **Spreadsheet Configuration of Risk Sets**.
13. Drag the bottom of the Future Risk panel down so that you can see all 13 rows of the Future Risk table.

Note: When using time-to-event formatted data, the existing production quantities must be specified in the Future Risk table using negative numeric time values.

14. Type in the Count values that are shown in Figure 9.15.

Figure 9.15 Future Risk Count Specifications

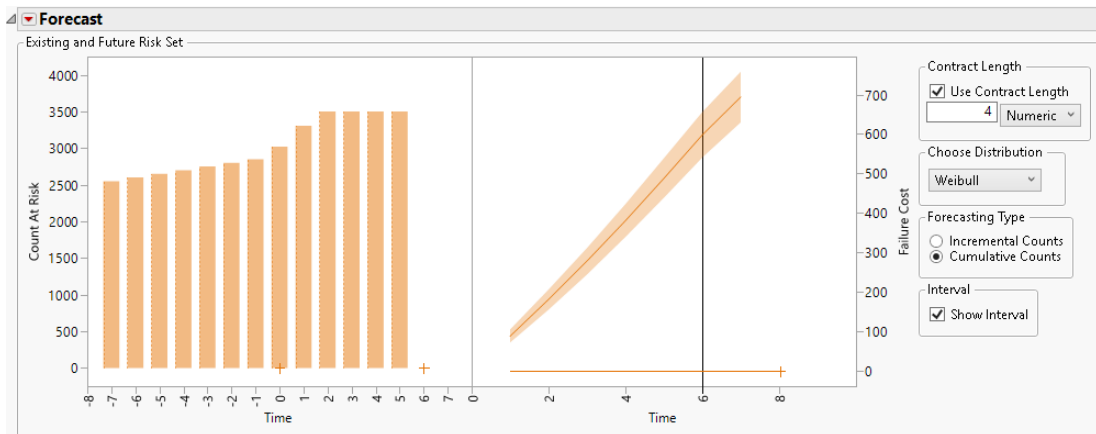
Future Risk		
	Time	Count
1	-7	2550
2	-6	2600
3	-5	2650
4	-4	2700
5	-3	2750
6	-2	2800
7	-1	2850
8	0	3022
9	1	3307
10	2	3502
11	3	3502
12	4	3502
13	5	3502
optional		

15. In the Forecasting Type panel, select **Cumulative Counts**.
16. In the Interval panel, select **Show Interval**.
17. Right-click the horizontal axis in the right Forecast graph and select **Axis Settings**. Add a reference line at month 6 using the following steps:
 1. Type 6 for the Value.
 2. Click **Add**.
 3. Click **OK**.

The forecast for the total number of units that will be returned for repair in the next 6 months is approximately 600 with confidence interval about 550 to 650.

Tip: You can use the **Save Forecast Data Table** option in the Forecast red triangle menu to see exact values for the forecasts and intervals.

Figure 9.16 Forecast Results



Chapter 10

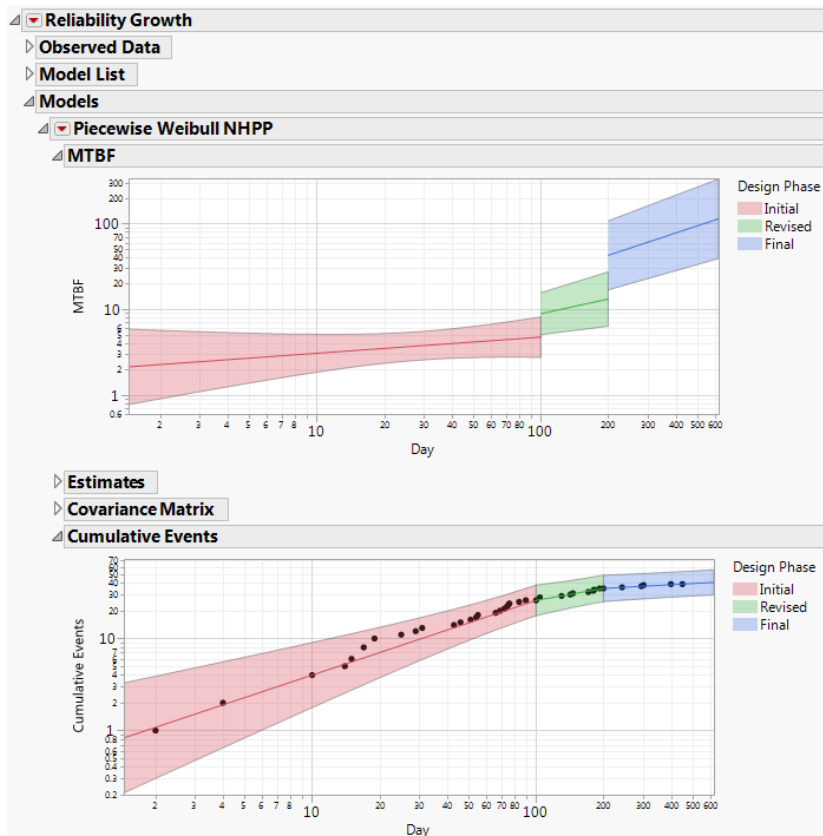
Reliability Growth

Model System Reliability as Changes Are Implemented

The Reliability Growth platform models the change in reliability of a single repairable system over time as improvements are incorporated into its design. A reliability growth testing program attempts to increase the system's mean time between failures (MTBF) by integrating design improvements as failures are discovered.

The Reliability Growth platform fits Crow-AMSAA models. These are non-homogeneous Poisson processes with Weibull intensity functions. Separate models can accommodate various phases of a reliability growth program. The platform also fits models for multiple systems.

Figure 10.1 Example of Plots for a Three-Phase Reliability Growth Model



Contents

Overview of the Reliability Growth Platform	255
Example Using the Reliability Growth Platform	256
Launch the Reliability Growth Platform	260
Launch Window Roles	261
Data Table Structure	262
The Reliability Growth Report	265
Model Descriptions and Feasibilities	265
Observed Data Report	265
Reliability Growth Platform Options	268
Model Reports	269
Crow AMSAA	270
Crow AMSAA with Modified MLE	276
Fixed Parameter Crow AMSAA	277
Piecewise Weibull NHPP	278
Reinitialized Weibull NHPP	282
Piecewise Weibull NHPP Change Point Detection	285
Additional Examples of the Reliability Growth Platform	286
Piecewise NHPP Weibull Model Fitting with Interval-Censored Data	286
Piecewise Weibull NHPP Change Point Detection with Time in Dates Format	289
Statistical Details for the Reliability Growth Platform	291
Statistical Details for the Crow-AMSAA Report	291
Statistical Details for the Piecewise Weibull NHPP Change Point Detection Report	292

Overview of the Reliability Growth Platform

The Reliability Growth platform fits Crow-AMSAA models, described in MIL-HDBK-189 (1981). A Crow-AMSAA model is a non-homogeneous Poisson process (NHPP) model with Weibull intensity; it is also known as a power law process. Such a model allows the failure intensity to vary over time. The failure intensity is defined by the two parameters β and λ .

For single-prototype data, the platform fits four classes of models and performs automatic change-point detection. The following reports are available:

- Simple Crow-AMSAA model, where both parameters are estimated using maximum likelihood.
- Crow-AMSAA with Modified MLE, where the maximum likelihood estimate for β is corrected for bias.
- Fixed Parameter Crow-AMSAA model, where the user is allowed to fix either or both parameters.
- Piecewise Weibull NHPP model, where the parameters are estimated for each test phase, taking failure history from previous phases into account.
- Reinitialized Weibull NHPP model, where both parameters are estimated for each test phase in a manner that ignores failure history from previous phases.
- Automatic estimation of a change-point and the associated piecewise Weibull NHPP model, for reliability growth situations where different failure intensities can define two distinct test phases.

For multiple-prototype data, the platform fits the following classes of models:

- Piecewise Weibull NHPP model, where each system in a multi-phase study follows the same piecewise Weibull NHPP model. The differences among the systems are assumed to be due to the randomness of individual realizations of the same model. This model contains one β parameter for each phase and one λ parameter.
- Piecewise Weibull NHPP with Different Intercepts model, where each system in a multi-phase study follows a separate piecewise Weibull NHPP model. This model contains one β parameter for each phase and one λ parameter for each system.
- Distinct Phase Weibull NHPP model, where each system in a multi-phase study follows the same Crow-AMSAA model in each phase. This model contains one β parameter and one λ parameter for each phase.
- Distinct Weibull NHPP model, where each system in a multi-phase study follows a separate Crow-AMSAA model in each phase. This model contains one β parameter and one λ parameter for each combination of system and phase in the study.

- Distinct System Weibull NHPP model, where each system in the study follows a separate Crow-AMSAA model with different parameters.
- Identical System Weibull NHPP model, where each system in the study follows a single Crow-AMSAA model. The differences among the systems are assumed to be due to the randomness of individual realizations of the same model.

Interactive profilers enable you to explore changes in MTBF, failure intensity, and cumulative failures over time. When you suspect a change in intensity over the testing period, you can use the change-point detection option to estimate a change-point and its corresponding model.

Example Using the Reliability Growth Platform

Suppose that you are testing a prototype for a new class of turbine engines. The testing program has been ongoing for over a year and has been through three phases.

The data are given in the `TurbineEngineDesign1.jmp` sample data table, found in the Reliability subfolder. For each failure that occurred, the number of days since test initiation was recorded in the column `Day`. The number of failures on a given day, or equivalently, the number of required fixes, was recorded in the column `Fixes`.

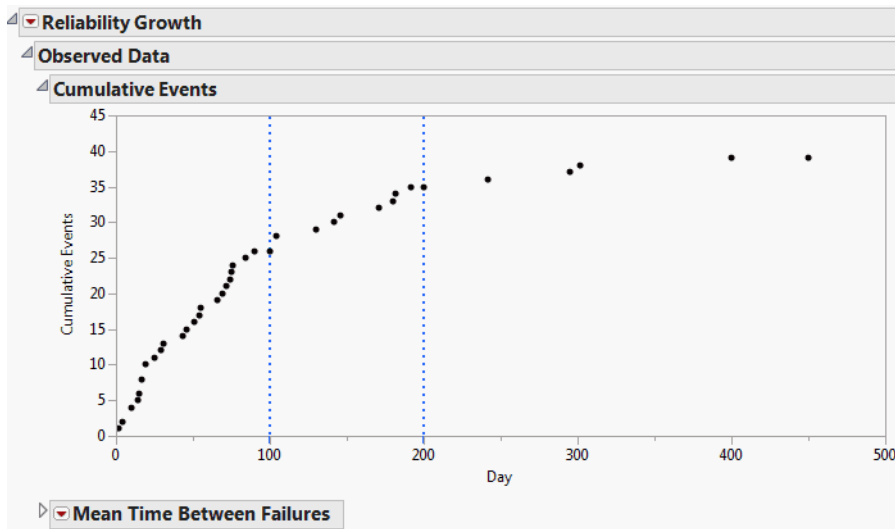
The first 100-day phase of the program was considered the initial testing phase. Failures were addressed with aggressive design changes, resulting in a substantially revised design. This was followed by another 100-day phase, during which failures in the revised design were addressed with design changes to subsystems. The third and final testing phase ran for 250 days. During this final phase, failures were addressed with local design changes.

Each phase of the testing was terminated based on the designated number of days, so that the phases are time terminated. Specifically, a given phase is time terminated at the start time of the next phase. However, the failure times are exact (not censored).

1. Select **Help > Sample Data Library** and open `Reliability/TurbineEngineDesign1.jmp`.
2. Select **Analyze > Reliability and Survival > Reliability Growth**.
3. On the **Time to Event Format** tab, select `Day` and click **Time to Event**.
4. Select `Fixes` and click **Event Count**.
5. Select `Design Phase` and click **Phase**.
6. Click **OK**.

The Reliability Growth report appears. The Cumulative Events plot shows the cumulative number of failures by day. Vertical dashed blue lines show the two transitions between the three phases.

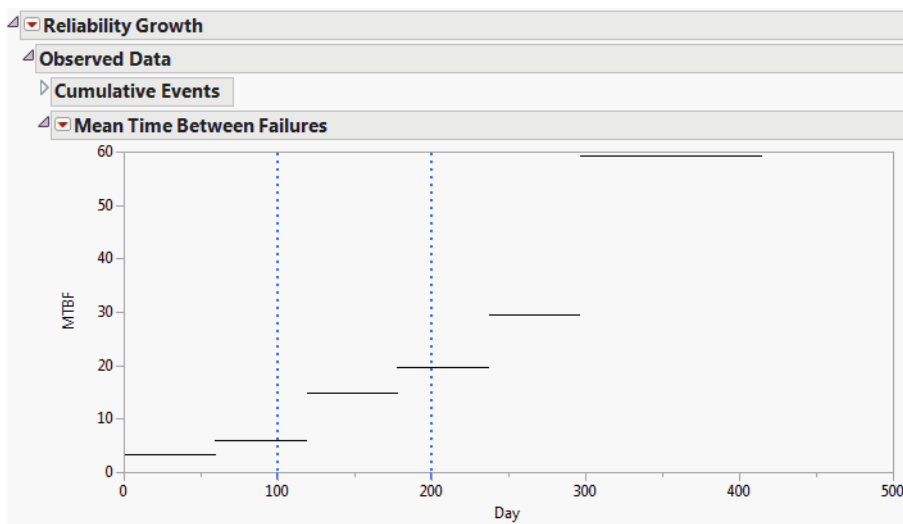
Figure 10.2 Observed Data Report



7. Click the **Mean Time Between Failures** disclosure icon.

This provides a plot with horizontal lines at the mean times between failures computed over intervals of a predetermined size. An option in the red triangle menu enables you to specify the interval size.

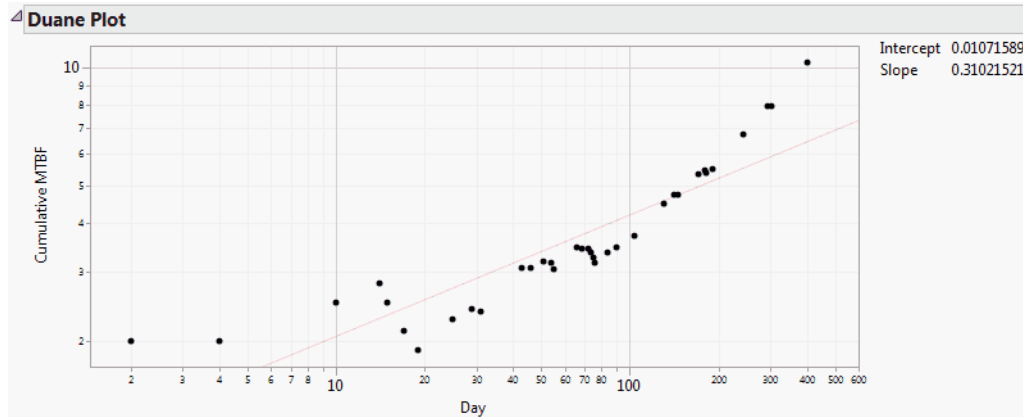
Figure 10.3 Mean Time between Failures Plot



8. Click the **Duane Plot** disclosure icon.

This provides a plot that displays the Cumulative MTBF estimates on the vertical axis versus the time to event variable on the horizontal axis. If the data follow the Crow-AMSAA model, the points should fall along a line when plotted on log-log paper.

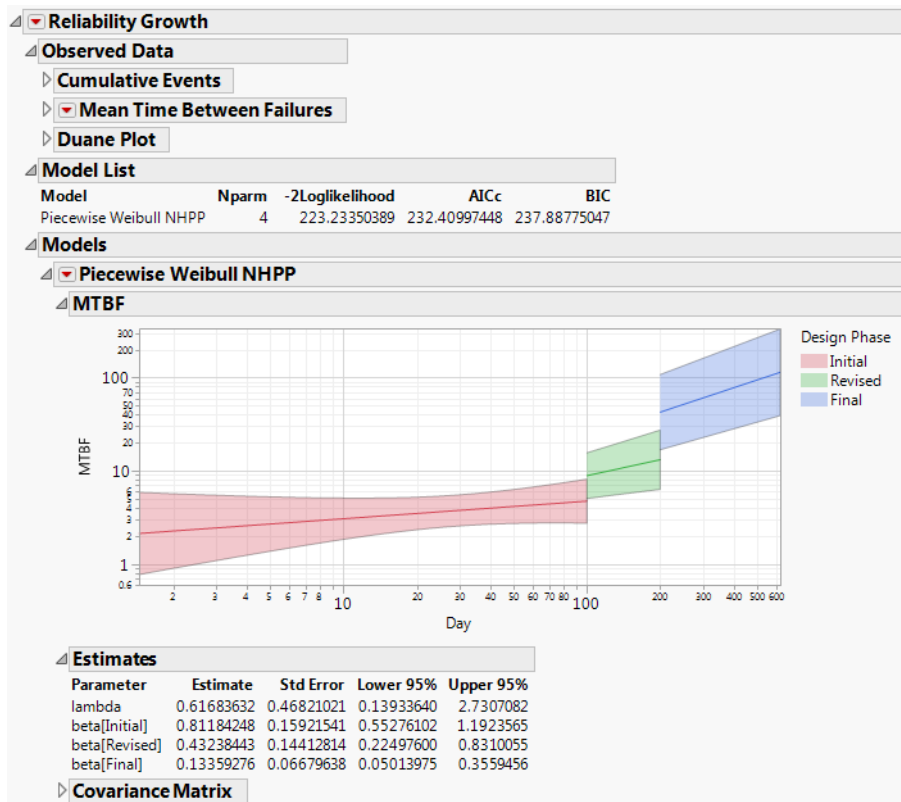
Figure 10.4 Duane Plot



9. Click the Reliability Growth red triangle and select **Fit Model > Piecewise Weibull NHPP**.

This fits Weibull NHPP models to the three phases of the testing program, treating these phases as multiple stages of a single reliability growth program (Figure 10.5). Options in the Piecewise Weibull NHPP red triangle menu provide various other plots and reports.

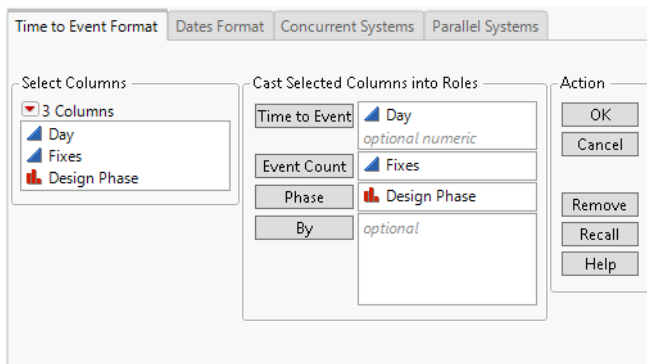
Figure 10.5 Piecewise Weibull NHPP Report



Launch the Reliability Growth Platform

Launch the Reliability Growth platform by selecting **Analyze > Reliability and Survival > Reliability Growth**. The launch window, using data from TurbineEngineDesign1.jmp, is shown in Figure 10.6.

Figure 10.6 Reliability Growth Launch Window



For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The launch window includes a tab for each of four data formats: Time to Event Format, Dates Format, Concurrent Systems, and Parallel Systems.

- The Dates Format assumes that time is recorded in a date/time format, indicating an absolute date or time. On the Dates Format tab, the time column or columns are given the Timestamp role.
- The other three data formats assume that time is recorded as the number of time units; for example, days or hours, since initial start-up of the system. The test start time is assumed to be time zero. On the Time to Event Format tab, the time column or columns are given the Time to Event role.

Rows with missing data in the columns Time to Event, Timestamp, Event Count, Phase, or System ID are not included in the analysis.

Note: All data formats for the Reliability Growth platform require that times or time intervals be specified in non-decreasing order.

Launch Window Roles

The roles that are available for variables in the Reliability Growth launch window are determined by the specified data format tab. The roles are explained in this section.

Time to Event

Time to Event is the number of time units that elapse between the start of the test and the occurrence of an event (a failure or test termination). The test start time is assumed to be time zero. Note that the Time to Event role is available only on the Time to Event Format tab.

Two conventions are allowed (See [“Exact Failure Times versus Interval Censoring”](#) on page 262.):

- A single column can be entered. In this case, it is assumed that the column gives the exact elapsed times at which events occurred.
- Two columns can be entered, giving interval start and end times. If an interval's start and end times differ, it is assumed that the corresponding events given in the Event Count column occurred at some unknown time within that interval. We say that the data are interval-censored. If the interval start and end times are identical, it is assumed that the corresponding events occurred at that specific time point, so that the times are exact (not censored).

The platform requires that the time columns be sorted in non-decreasing order. When two columns giving interval start and end times are provided, these intervals must not overlap (except at their endpoints). Intervals with zero event counts that fall strictly within a phase can be omitted, as they do not affect the likelihood function.

Timestamp

Timestamp is an absolute time (for example, a date). As with Time to Event, Timestamp allows times to be entered using either a single column or two columns. Note that the Timestamp role is available only on the Dates Format tab.

For times entered as Timestamp, the first row of the table is considered to give the test start time:

- When a single column is entered, the timestamp corresponding to the test start, with an event count of 0, should appear in the first row.
- When two columns, giving time intervals, are entered, the first entry in the first column should be the test start timestamp. (See [“Phase”](#) on page 262.)

Other details are analogous to those described for Time to Event Format in the section [“Time to Event”](#) on page 261. See also [“Exact Failure Times versus Interval Censoring”](#) on page 262.

Event Count

Event Count is the number of events, usually failures addressed by corrective actions (fixes), occurring at the specified time or within the specified time interval. If no column is entered as Event Count, it is assumed that the Event Count for each row is one.

System ID

System ID identifies the prototype for an observation for multiple-prototype data. The System ID role is available only in the Concurrent Systems and Parallel Systems data formats. It is a required role for both data formats.

Phase

Reliability growth programs often involve several periods, or phases, of active testing. These testing phases can be specified in the optional Phase column. For the Time to Event Format and Dates Format data formats, the Phase variable can be of any data or modeling type. For the Parallel Systems data format, the data type of the Phase variable must be Numeric. For more information about structuring multi-phase data, see [“Test Phases”](#) on page 263. For an example, see [“Piecewise NHPP Weibull Model Fitting with Interval-Censored Data”](#) on page 286.

By

This produces a separate analysis for each value that appears in the column.

Data Table Structure

The Time to Event Format and the Dates Format enable you to enter either a single column *or* two columns as Time to Event or Timestamp, respectively. The Concurrent Systems and the Parallel Systems enable you to enter one column corresponding to each level of the System ID variable. There are examples of multiple-prototype data tables in the Reliability folder of the Sample Data folder: Concurrent Systems.jmp for Concurrent Systems and four tables with the prefix Parallel Systems for Parallel Systems.

This section describes how to use these two approaches to specify the testing structure.

Exact Failure Times versus Interval Censoring

In some testing situations, the system being tested is checked periodically for failures. In this case, failures are known to have occurred within time intervals, but the precise time of a failure is not known. We say that the failure times are *interval-censored*.

The Reliability Growth platform accommodates both exact, non-censored failure-time data, and interval-censored data. When a single column is entered as Time to Event or Timestamp, the times are considered exact times (not censored).

When two columns are entered, the platform views these as defining the start and end points of time intervals. If an interval's start and end times differ, then the times for failures occurring within that interval are considered to be interval-censored. If the end points are identical, then the times for the corresponding failures are assumed to be exact and equal to that common time value. So, you can represent both exact and interval-censored failure times by using two time columns.

In particular, *exact failures times* can be represented in one of two ways: As times given by a single time column, or as intervals with identical endpoints, given by two time columns.

Model-fitting in the Reliability Growth platform relies on the likelihood function. The likelihood function takes into account whether interval-censoring is present or not. So, mixing interval-censored with exact failure times is permitted.

Failure and Time Termination

A test plan can call for test termination once a specific number of failures has been achieved or once a certain span of time has elapsed. For example, a test plan might terminate testing after 50 failures occur. Another plan might terminate testing after a six-month period.

If testing terminates based on a specified number of failures, we say that the test is *failure terminated*. If testing is terminated based on a specified time interval, we say that the test is *time terminated*. The likelihood functions used in the Reliability Growth platform reflect whether the test phases are failure or time terminated.

Test Phases

Reliability growth testing often involves several *phases* of testing. For example, the system being developed or the testing program might experience substantial changes at specific time points. The data table conveys the start time for each phase and whether each phase is failure or time terminated, as described below.

Single Test Phase

When there is a single test phase, the platform infers whether the test is failure or time terminated from the time and event count entries in the last row of the data table.

- If the last row contains an exact failure time with a nonzero event count, the test is considered failure terminated.
- If the last row contains an exact failure time with a zero event count, the test is considered time terminated.

- If the last row contains an interval with nonzero width, the test is considered time terminated with termination time equal to the right endpoint of that interval.

Note: To indicate that a test has been time terminated, be sure to include a last row in your data table showing the test termination time. If you are entering a single column as Time to Event or Timestamp, the last row should show a zero event count. If you are entering two columns as Time to Event or Timestamp, the right endpoint of the last interval should be the test-termination time. In this case, if there were no failures during the last interval, you should enter a zero event count.

Multiple Test Phases

When using Time to Event Format, the start time for any phase other than the first should be included in the time column(s). When using Dates Format, the start times for all phases should be included in the time column(s). If no events occurred at a phase start time, the corresponding entry in the Event Count column should be zero. For times given in two columns, it might be necessary to reflect the phase start time using an interval with identical endpoints and an event count of zero.

In a multi-phase testing situation, the platform infers whether each phase, other than the last, is failure or time terminated from the entries in the last row preceding a phase change. Suppose that Phase A ends and that Phase B begins at time t_B . In this case, the first row corresponding to Phase B contains an entry for time t_B .

- If the failure time for the last failure in Phase A is exact and if that time differs from t_B , then Phase A is considered to be time terminated. The termination time is equal to t_B .
- If the failure time for the last failure in Phase A is exact and is equal to t_B , then Phase A is considered to be failure terminated.
- If the last failure in Phase A is interval-censored, then Phase A is considered to be time terminated with termination time equal to t_B .

The platform infers whether the final phase is failure or time terminated from the entry in the last row of the data table.

- If the last row contains an exact failure time with a nonzero event count, the test is considered failure terminated.
- If the last row contains an exact failure time with a zero event count, or an interval with nonzero width, the test is considered time terminated. In the case of an interval, the termination time is taken as the right endpoint.

The Reliability Growth Report

The Observed Data report appears by default. If the Parallel Systems data format is specified in the launch window, the Model Descriptions and Feasibilities report also appears by default.

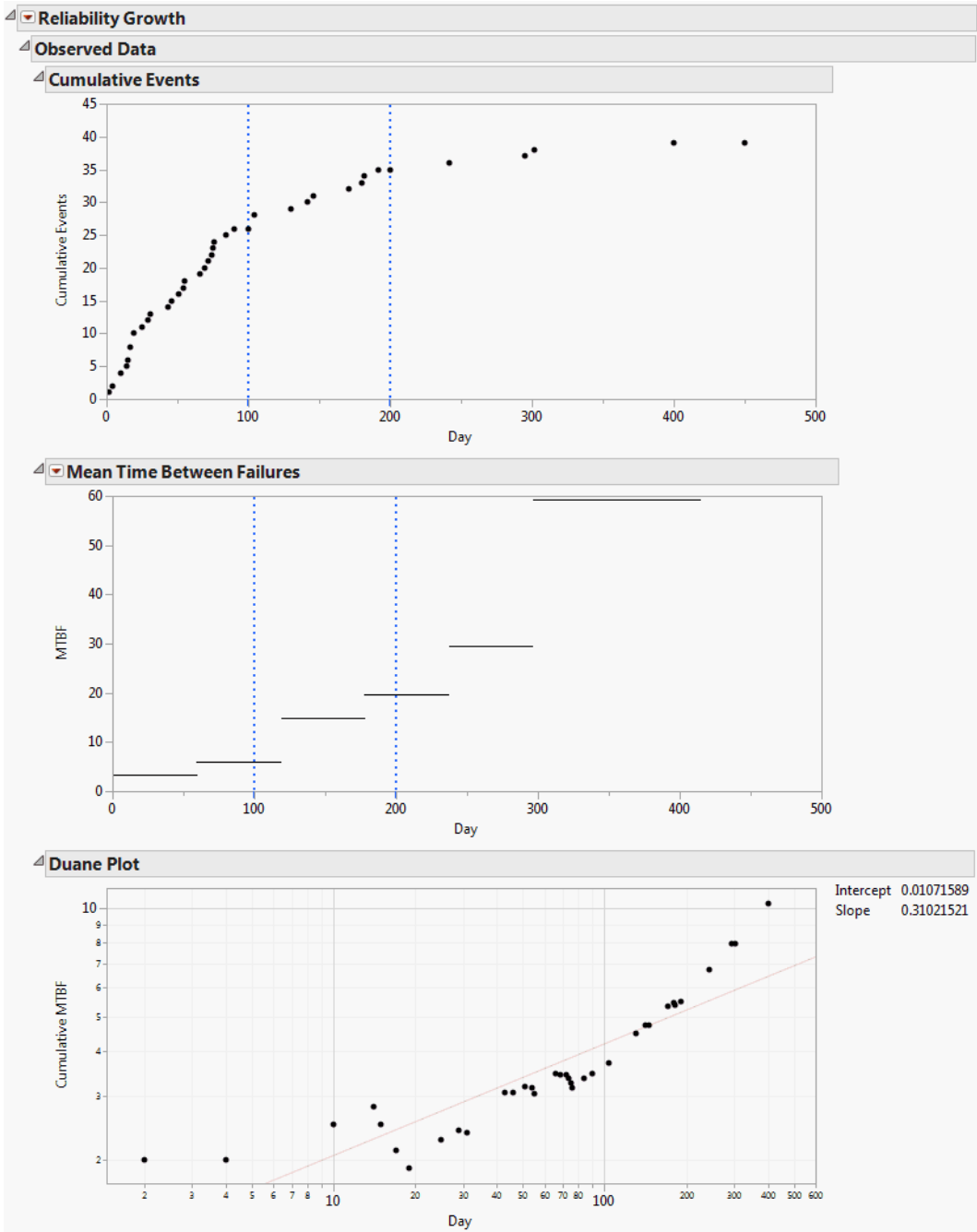
Model Descriptions and Feasibilities

The Model Description and Feasibilities report lists the name and descriptions of the parallel systems models. It also contains a column that shows if a model is available for the current data. If a model is listed as not available, the reason is reported in the right-most column of this table.

Observed Data Report

The Observed Data report contains the Cumulative Events plot, the Mean Time Between Failures plot, and the Duane plot. These are shown in Figure 10.7, where we have opened the Mean Time Between Failures and Duane Plot reports. To produce this report, follow the instructions in [“Example Using the Reliability Growth Platform”](#) on page 256.

Figure 10.7 Observed Data Report



Cumulative Events Plot

This plot shows how events are accumulating over time. The vertical coordinate for each point on the Cumulative Events plot equals the total number of events that have occurred by the time given by the point's horizontal coordinate.

Whenever a model is fit, that model is represented on the Cumulative Events plot. Specifically, the cumulative events estimated by the model are shown by a curve, and 95% confidence intervals are shown by a solid band. Check boxes to the right of the plot enable you to control which models are displayed.

For the Parallel Systems data format, the Cumulative Events plot has a separate panel in the plot for each level of the System ID variable. It also contains a panel at the bottom of the plot that overlays the cumulative events for all levels of the System ID variable. The levels of the System ID variable are designated by colors.

Mean Time Between Failures Plot

The Mean Time Between Failures plot shows mean times between failures averaged over small time intervals of equal length. These are not tied to the Phases. The default number of equal length intervals is based on the number of rows.

Mean Time Between Failures Plot Options

Click the Mean Time Between Failures red triangle and select Options to open a window that enables you to specify intervals over which to average.

Two types of averaging are offered:

- Equal Interval Average MTBF (Mean Time Between Failures) enables you to specify a common interval size.
- Customized Average MTBF enables you to specify cut-off points for time intervals.
 - Double-click within a table cell to change its value.
 - Right-click in the table to open a menu that enables you to add and remove rows.

Duane Plot

The Duane Plot displays Cumulative MTBF estimates plotted against the Time to Event variable, with both axes on a log₁₀ scale. If the data follow the Crow-AMSAA model, the points should follow a line when plotted on log-log paper.

Note: The Duane Plot is available only if failure times are exact and Time to Event Format is used. The plot is not available for interval-censored data or data entered in Dates Format.

The line displayed on the plot is the least squares regression line for the regression of $\log_{10}$ of Cumulative MTBF on $\log_{10}$ of the Time to Event variable.

Note: The Duane Plot does not reflect Phase variables. Rows where the Time to Event variable defines Phase changes are ignored in constructing the plot and fitting the regression line.

Intercept and Slope

Intercept and Slope values are displayed to the right of the plot.

- The Intercept value is given, for historical reasons, as the intercept for a fit in the *natural* logarithmic scale. Specifically, the natural logarithm of Cumulative MTBF is regressed on the natural logarithm of Time to Event. The value of Intercept is the value predicted by this regression equation at $\log(1) = 0$, where *log* is the natural logarithm. To obtain the intercept for a fit in terms of base 10 logarithms, divide the Intercept value by $\log(10)$. See Tobias and Trindade (2012, ch. 13).
- The Slope value is the slope for a fit in either the natural or the logarithmic scale. This follows from the properties of logarithms.

Reliability Growth Platform Options

The Reliability Growth red triangle menu has the following options:

Fit Model If the Time to Event Format, Dates Format, or Concurrent Systems data formats are specified in the launch window, this menu contains options to fit various Non-Homogeneous Poisson Process (NHPP) models. These options are described in detail below. Depending on the choices made in the launch window, the possible options are:

- “Crow AMSAA” on page 270
- “Crow AMSAA with Modified MLE” on page 276
- “Fixed Parameter Crow AMSAA” on page 277
- “Piecewise Weibull NHPP” on page 278
- “Reinitialized Weibull NHPP” on page 282
- “Piecewise Weibull NHPP Change Point Detection” on page 285

Fit Parallel System Model If the Parallel Systems data format is specified in the launch window, this menu contains options to fit various models for multiple-prototype data. The possible options in this menu are dependent on choices made in the launch window.

See *Using JMP* for more information about the following options:

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Model List

Once a model is fit, the Model List report appears. This report provides various statistical measures that describe the fit of the model. As additional models are fit, they are added to the Model List, which provides a convenient summary for model comparison. The models are sorted in ascending order based on AICc. The Model List report contains the following statistics:

Nparm The number of parameters in the model.

-2Loglikelihood The likelihood function is a measure of how probable the observed data are, given the estimated model parameters. In a general sense, the higher the likelihood, the better the model fit. It follows that smaller values of twice the negative of the log-likelihood (-2Loglikelihood) indicate better model fits.

AICc The Corrected Akaike's Information Criterion.

BIC The Bayesian Information Criterion.

See *Fitting Linear Models* for more information about -2Loglikelihood, AICc, and BIC.

Model Reports

This section describes the reports generated when the available models are fit as well as the options available in those model reports.

- "Crow AMSAA"
- "Crow AMSAA with Modified MLE"
- "Fixed Parameter Crow AMSAA"
- "Piecewise Weibull NHPP"
- "Reinitialized Weibull NHPP"
- "Piecewise Weibull NHPP Change Point Detection"

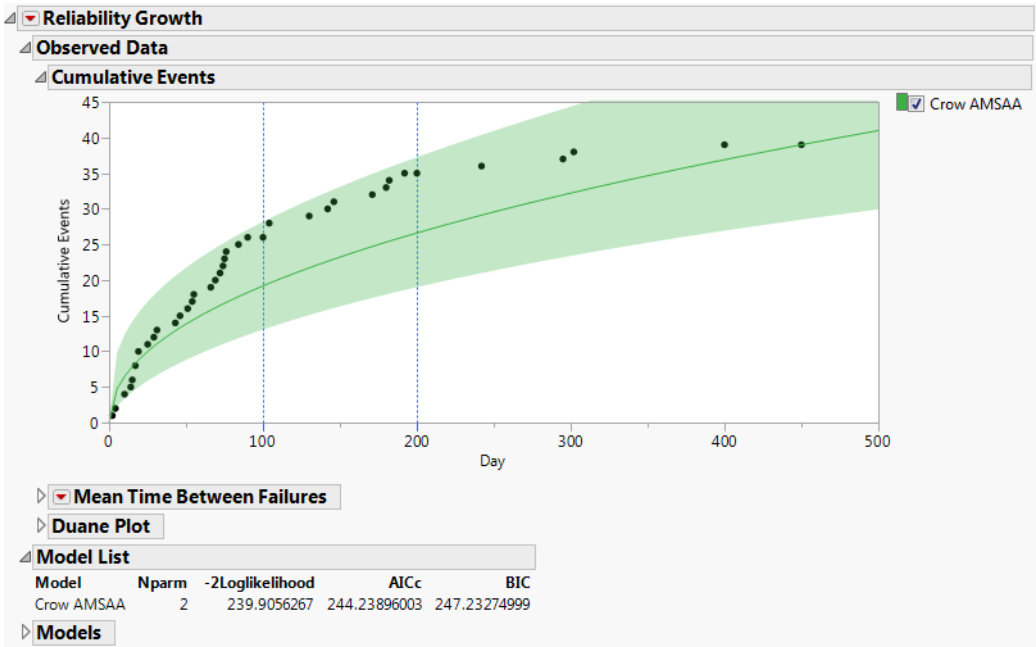
Crow AMSAA

This option fits a Crow-AMSAA model (MIL-HDBK-189 1981). A Crow-AMSAA model is a nonhomogeneous Poisson process with failure intensity as a function of time t given by $\rho(t) = \lambda\beta t^{\beta-1}$. Here, λ is a scale parameter and β is a growth parameter. This function is also called a Weibull intensity, and the process itself is also called a power law process (Rigdon and Basu 2000; Meeker and Escobar 1998). Note that the Recurrence platform fits the Power Nonhomogeneous Poisson Process. The Power Nonhomogeneous Poisson Process is equivalent to the Crow-AMSAA model, although it uses a different parameterization. See “Fit Model” on page 164 in the “Recurrence Analysis” chapter.

The intensity function is a concept applied to repairable systems. Its value at time t is the limiting value of the probability of a failure in a small interval around t , divided by the length of this interval; the limit is taken as the interval length goes to zero. You can think of the intensity function as measuring the likelihood of the system failing at a given time. If $\beta < 1$, the system is improving over time. If $\beta > 1$, the system is deteriorating over time. If $\beta = 1$, the rate of occurrence of failures is constant.

When the Crow AMSAA option is selected, the Cumulative Events plot updates to show the cumulative events curve estimated by the model. For each time point, the shaded band around this curve defines a 95% confidence interval for the true cumulative number of events at that time. The Model List report also updates. Figure 10.8 shows the Observed Data report for the data in TurbineEngineDesign1.jmp.

Figure 10.8 Crow AMSAA Cumulative Events Plot and Model List Report



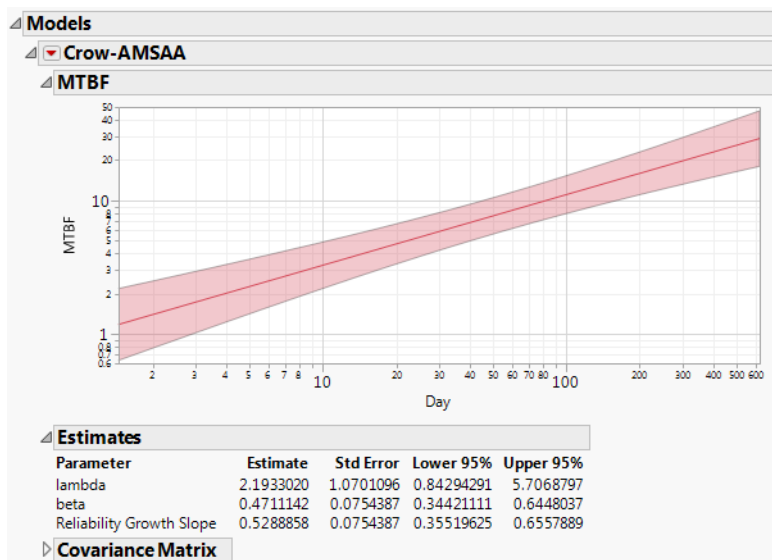
Crow-AMSAA Report

A Crow-AMSAA report opens within the Models report. If Time to Event Format is used, the Crow-AMSAA report shows an MTBF plot with both axes scaled logarithmically. See [“MTBF Plot”](#) on page 271.

MTBF Plot

The mean time between failures (MTBF) plot is displayed by default (Figure 10.9). For each time point, the shaded band around the MTBF plot defines a 95% confidence interval for the true MTBF at time t . The plot is shown with both axes logarithmically scaled. With this scaling, the MTBF plot is linear.

Figure 10.9 MTBF Plot



To see why the MTBF plot is linear when logarithmic scaling is used, consider the following. The mean time between failures is the reciprocal of the intensity function. For the Weibull intensity function, the MTBF is $1/(\lambda\beta t^{\beta-1})$, where t represents the time since testing initiation. It follows that the logarithm of the MTBF is a linear function of $\log(t)$, with slope $1 - \beta$. The estimated MTBF is defined by replacing the parameters λ and β by their estimates. So the log of the estimated MTBF is a linear function of $\log(t)$.

Estimates

Maximum likelihood estimates for lambda (λ), beta (β), and the Reliability Growth Slope ($1 - \beta$), appear in the Estimates report below the plot (Figure 10.9). Standard errors and 95% confidence intervals for λ , β , and $1 - \beta$ are given. For more information about the calculations, see [“Parameter Estimates for Crow-AMSAA Models”](#) on page 291.

Covariance Matrix

Estimated covariance matrix for the estimates of the parameters of the fitted model. This report is closed by default.

Crow-AMSAA Options

This section describes the options that are available in the Crow-AMSAA red triangle menu when a Crow AMSAA model is fit.

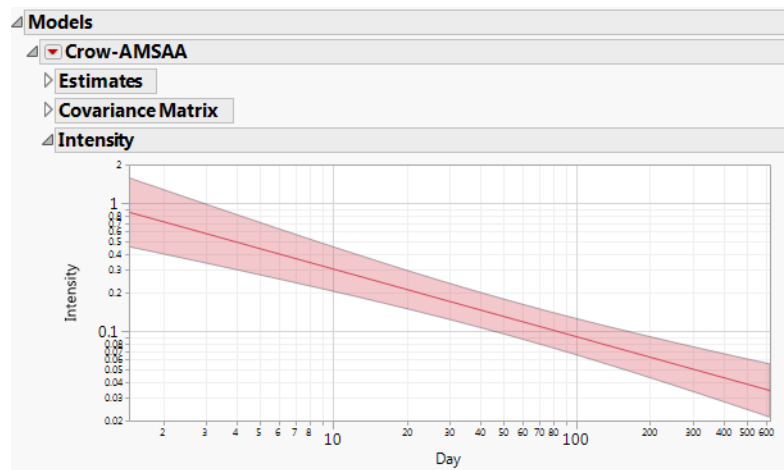
Show MTBF Plot

This option shows or hides the MTBF Plot. See [“MTBF Plot”](#) on page 271.

Show Intensity Plot

This plot shows the estimated intensity function. The Weibull intensity function is given by $\rho(t) = \lambda\beta t^{\beta-1}$, so it follows that $\log(\text{Intensity})$ is a linear function of $\log(t)$. Both axes are scaled logarithmically.

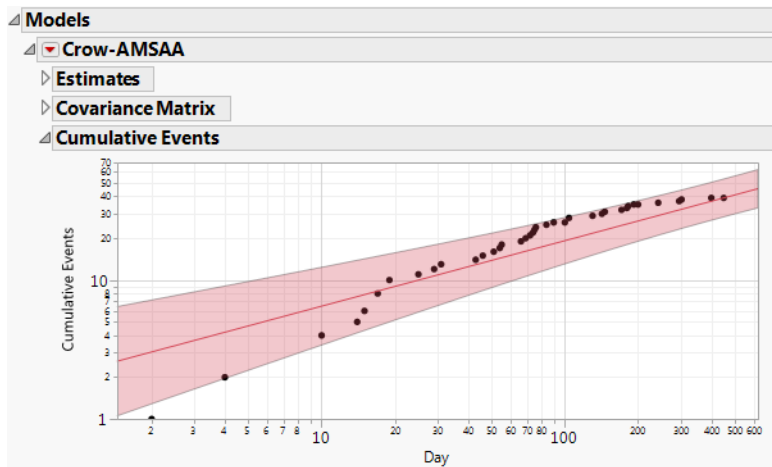
Figure 10.10 Intensity Plot



Show Cumulative Events Plot

This plot shows the estimated cumulative number of events. The observed cumulative numbers of events are also displayed on this plot. Both axes are scaled logarithmically.

Figure 10.11 Cumulative Events Plot



For the Crow-AMSAA model, the cumulative number of events at time t is given by λt^β . It follows that the logarithm of the cumulative number of events is a linear function of $\log(t)$. So, the plot of the estimated Cumulative Events is linear when plotted against logarithmically scaled axes.

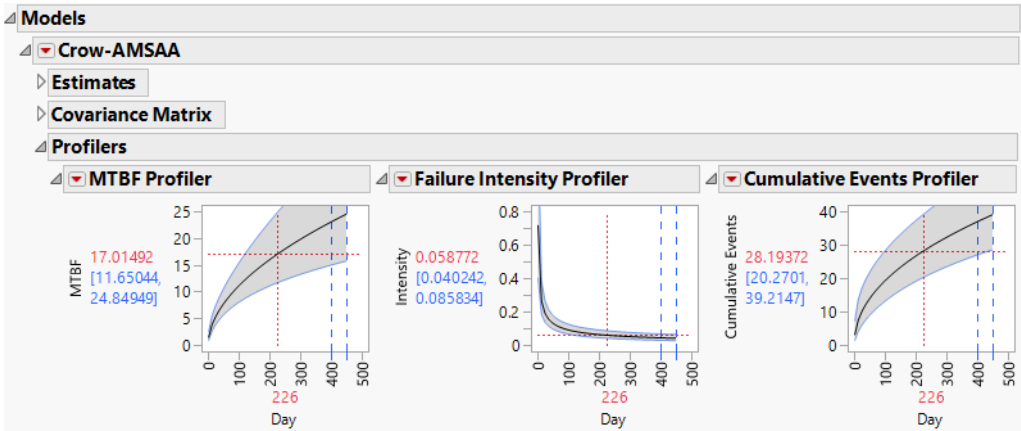
Show Profilers

Three profilers are displayed, showing estimated MTBF, Failure Intensity, and Cumulative Events. These profilers do not use logarithmic scaling. By dragging the red vertical dashed line in any profiler, you can explore model estimates at various time points; the value of the selected time point is shown in red beneath the plot. Also, you can set the time axis to a specific value by pressing Ctrl while you click in the plot. A blue vertical dashed line denotes the time point of the last observed failure.

The profilers also display 95% confidence bands for the estimated quantities. For the specified time setting, the estimated quantity (in red) and 95% confidence limits (in black) are shown to the left of the profiler. See [“Profilers”](#) on page 292.

Note that you can link these profilers by selecting **Factor Settings > Link Profilers** from any of the profiler red triangle menus. For more information about the use and interpretation of profilers, see *Fitting Linear Models*. See also *Profilers*.

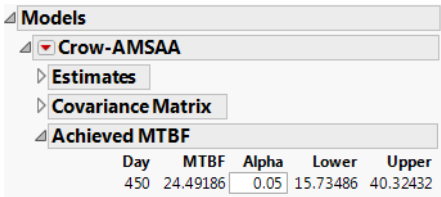
Figure 10.12 Profilers



Achieved MTBF

A confidence interval for the MTBF at the point when testing concludes is often of interest. For uncensored failure time data, this report gives an estimate of the Achieved MTBF and a 95% confidence interval for the Achieved MTBF. You can specify a $100 \cdot (1 - \alpha)\%$ confidence interval by entering a value for Alpha. The report is shown in Figure 10.13. For censored data, only the estimated MTBF at test termination is reported.

Figure 10.13 Achieved MTBF Report



There are infinitely many possible failure-time sequences from an NHPP; the observed data represent only one of these. Suppose that the test is failure terminated at the n^{th} failure. The confidence interval computed in the Achieved MTBF report takes into account the fact that the n failure times are random. If the test is time terminated, then the number of failures as well as their failure times are random. Because of this, the confidence interval for the Achieved MTBF differs from the confidence interval provided by the MTBF Profiler at the last observed failure time. See Crow (1982) and Lee and Lee (1978).

When the test is failure terminated, the confidence interval for the Achieved MTBF is exact. However, when the test is time terminated, an exact interval cannot be obtained. In this case, the limits are conservative in the sense that the interval contains the Achieved MTBF with probability at least $1 - \alpha$.

Goodness of Fit

The Goodness of Fit report tests the null hypothesis that the data follow a Crow-AMSAA model. Depending on whether one or two time columns are entered, either a Cramér-von Mises (see “[Cramér-von Mises Test for Data with Uncensored Failure Times](#)” on page 275) or a chi-squared test (see “[Chi-Squared Goodness of Fit Test for Interval-Censored Failure Times](#)” on page 275) is performed.

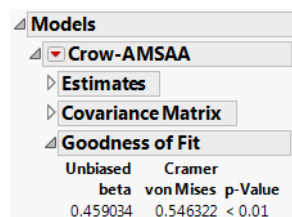
Cramér-von Mises Test for Data with Uncensored Failure Times

When the data are entered in the launch window as a single Time to Event or Timestamp column, the goodness of fit test is a Cramér-von Mises test. For the Cramér-von Mises test, large values of the test statistic lead to rejection of the null hypothesis and the conclusion that the model does not fit adequately. The test uses an unbiased estimate of beta, given in the report. The value of the test statistic is found below the Cramér-von Mises heading.

The entry below the p-Value heading indicates how unlikely it is for the test statistic to be as large as what is observed if the data come from a Crow-AMSAA model. The platform computes *p*-values up to 0.25. If the test statistic is smaller than the value that corresponds to a *p*-value of 0.25, the report indicates that its *p*-value is ≥ 0.25 . For more information about this test, see Crow (1975).

Figure 10.14 shows the goodness-of-fit test for the fit of a Crow-AMSAA model to the data in TurbineEngineDesign1.jmp. The computed test statistic corresponds to a *p*-value that is less than 0.01. We conclude that the Crow-AMSAA model does not provide an adequate fit to the data.

Figure 10.14 Goodness of Fit Report - Cramér-von Mises Test



Models		
Crow-AMSAA		
Estimates		
Covariance Matrix		
Goodness of Fit		
Unbiased	Cramer	
beta	von Mises	p-Value
0.459034	0.546322	< 0.01

Chi-Squared Goodness of Fit Test for Interval-Censored Failure Times

When the data are entered in the launch window as two Time to Event or Timestamp columns, a chi-squared goodness of fit test is performed. The chi-squared test is based on comparing observed to expected numbers of failures in the time intervals defined. Large values of the test statistic lead to rejection of the null hypothesis and the conclusion that the model does not fit.

In the Reliability Growth platform, the chi-squared goodness of fit test is intended for interval-censored data where the time intervals specified in the data table cover the entire time period of the test. This means that the start time of an interval is the end time of the preceding interval. In particular, intervals where no failures occurred should be included in the data table. If some intervals are not consecutive, or if some intervals have identical start and end times, the algorithm makes appropriate accommodations. But the resulting test is only approximately correct.

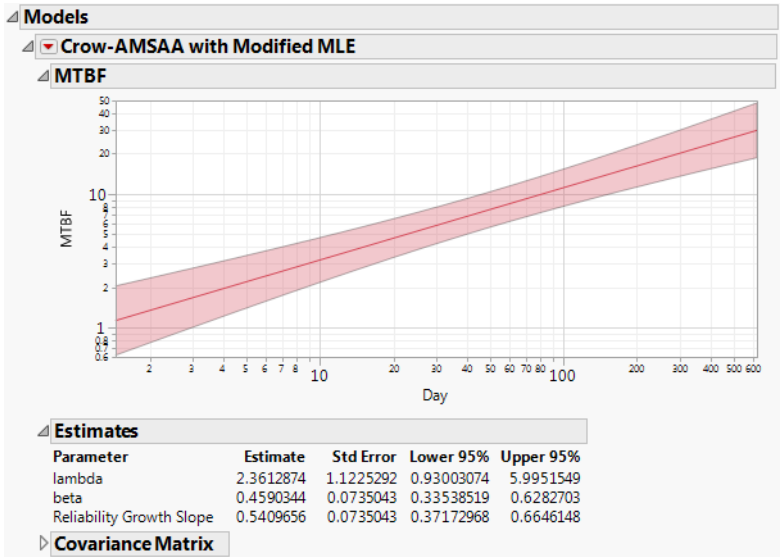
Crow AMSAA with Modified MLE

In the Crow-AMSAA model, the maximum likelihood estimate (MLE) of β is biased. This option fits a Crow AMSAA model where β is adjusted for bias.

Note: This option is available only when the data are entered in the launch window as a single Time to Event or Timestamp column. It is not available for interval-censored data.

Figure 10.15 shows a Crow-AMSAA with Modified MLE fit to the data in TurbineEngineDesign1.jmp.

Figure 10.15 Crow-AMSAA with Modified MLE Report



The formula for the bias-corrected estimate of β depends on whether the test is failure terminated or time terminated. See [“Parameter Estimates for Crow-AMSAA with Modified MLE”](#) on page 291.

When the Crow-AMSAA with Modified MLE option is selected, the Cumulative Events Plot updates to display this model. The Model List also updates. The Crow-AMSAA with Modified MLE report opens to show the MTBF plot, Estimates, and Covariance Matrix for the Crow-AMSAA with Modified MLE fit; this plot is described in the section [“MTBF Plot”](#) on page 283.

In addition to Show MTBF plot, available options are Show Intensity Plot, Show Cumulative Events Plot, Show Profilers, Achieved MTBF, and Goodness of Fit. These reports are described under [“Crow AMSAA”](#) on page 270. For more information about how the modified MLEs are used to construct these reports, see [“Parameter Estimates for Crow-AMSAA with Modified MLE”](#) on page 291. Details about the Goodness of Fit and Achieved MTBF reports specific to the modified MLE option are given below.

Goodness of Fit

Because the Crow-AMSAA with Modified MLE option is available only when the data are entered as a single Time to Event or Timestamp column, the Goodness of Fit test is a Cramér-von Mises test. Because the estimate of β is used in this test is bias-corrected, the test results are identical to those of the Goodness of Fit test for the Crow-AMSAA model.

Achieved MTBF

The achieved MTBF is estimated using the modified MLEs. However, the confidence interval for the achieved MTBF uses the true MLE and is identical to the interval given by the Crow AMSAA model.

Fixed Parameter Crow AMSAA

This option enables you to specify parameter values for a Crow-AMSAA fit. If a Crow-AMSAA report has not been obtained before choosing the Fixed Parameter Crow-AMSAA option, then both a Crow-AMSAA report and a Fixed Parameter Crow-AMSAA report are provided.

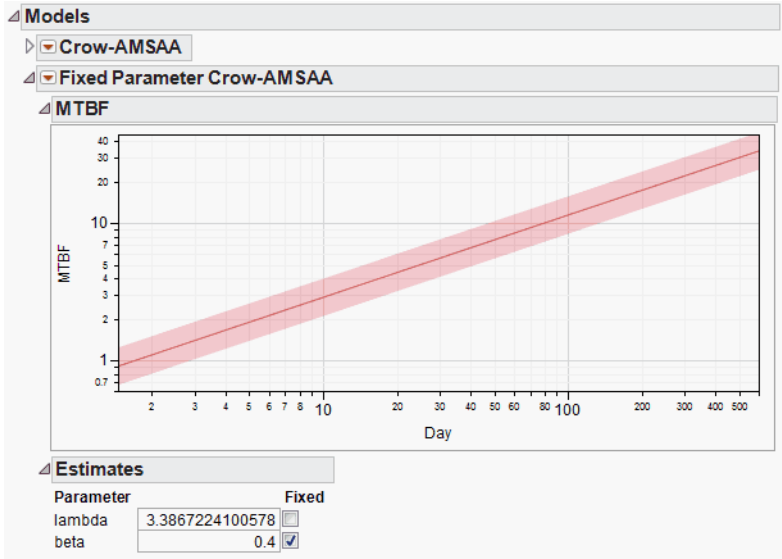
When the Fixed Parameter Crow-AMSAA option is selected, the Cumulative Events Plot updates to display this model. The Model List also updates. The Fixed Parameter Crow-AMSAA report opens to show the MTBF plot for the Crow-AMSAA fit; this plot is described in the section [“MTBF Plot”](#) on page 283.

In addition to Show MTBF plot, available options are Show Intensity Plot, Show Cumulative Events Plot, and Show Profilers. The construction and interpretation of these plots is described under [“Crow AMSAA”](#) on page 270.

Estimates

The initial parameter estimates are the MLEs from the Crow-ASMAA fit. Either parameter can be fixed by checking the box next to the desired parameter and then entering the desired value. The model is re-estimated and the MTBF plot updates to describe this model. Figure 10.16 shows a fixed-parameter Crow-AMSAA fit to the data in TurbineEngineDesign1.jmp, with the value of beta set at 0.4.

Figure 10.16 Fixed Parameter Crow-AMSAA Report

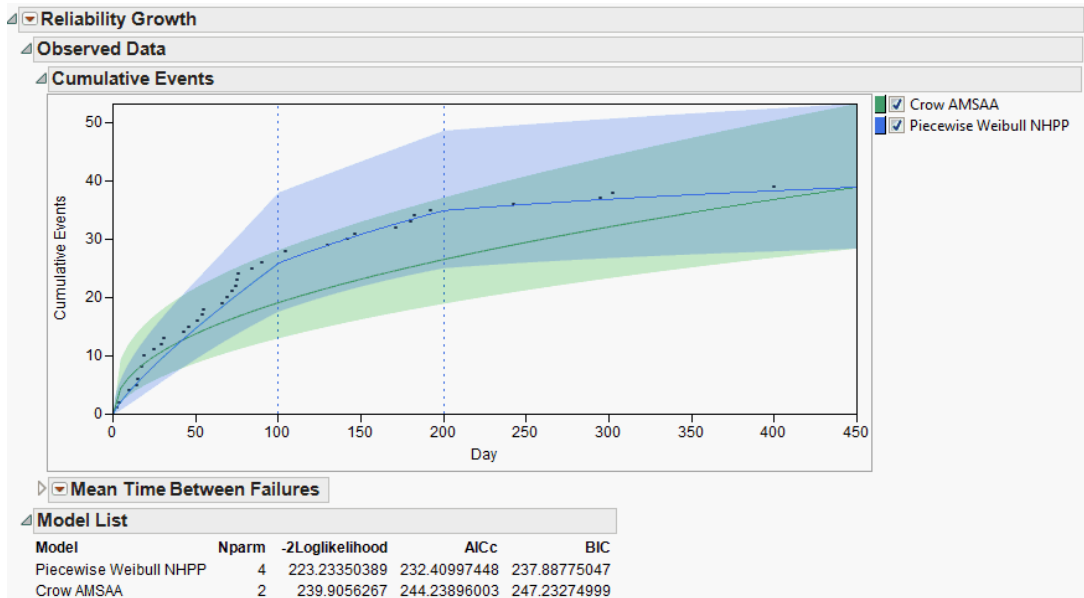


Piecewise Weibull NHPP

The Piecewise Weibull NHPP model can be fit when a Phase column specifying at least two values has been entered in the launch window. Crow-AMSAA models are fit to each of the phases under the constraint that the cumulative number of events at the start of a phase matches that number at the end of the preceding phase. For proper display of phase transition times, the first row for every Phase other than the first must give that phase's start time. See ["Multiple Test Phases"](#) on page 264.

When the report is run, the Cumulative Events plot updates to show the piecewise model. Blue vertical dashed lines show the transition times for each of the phases. The Model List also updates. See Figure 10.17, where both a Crow-AMSAA model and a piecewise Weibull NHPP model have been fit to the data in TurbineEngineDesign1.jmp. Note that both models are compared in the Model List report.

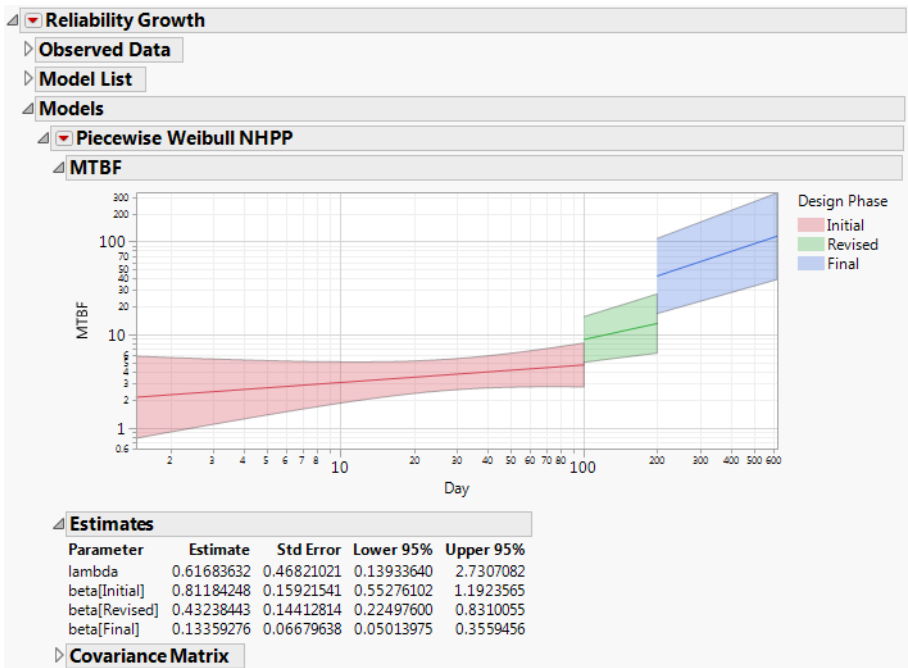
Figure 10.17 Cumulative Events Plot and Model List Report



Piecewise Weibull NHPP Report

By default, the Piecewise Weibull NHPP report shows the estimated MTBF plot. The phases are differentiated with different colors. The Estimates and Covariance Matrix reports are shown below the plot.

Figure 10.18 Piecewise Weibull NHPP Report



MTBF Plot

The MTBF plot and the Estimates report appear by default when the Piecewise Weibull NHPP option is chosen (Figure 10.18). When Time to Event Format is used, the axes are logarithmically scaled. For more information about the plot, see “MTBF Plot” on page 271.

Estimates

The Estimates report gives estimates of the model parameters. Note that only the estimate for the value of λ corresponding to the first phase is given. In the piecewise model, the cumulative events at the end of one phase must match the number at the beginning of the subsequent phase. Because of these constraints, the estimate of λ for the first phase and the estimates of the β s determine the remaining λ s.

The method used to calculate the estimates, their standard errors, and the confidence limits is similar to that used for the simple Crow-AMSAA model. See “Parameter Estimates for Crow-AMSAA Models” on page 291. The likelihood function reflects the additional parameters and the constraints on the cumulative numbers of events.

Covariance Matrix

Estimated covariance matrix for the estimates of the parameters of the fitted model. This report is closed by default.

Piecewise Weibull NHPP Options

This section describes the options available in the Piecewise Weibull NHPP red triangle menu.

Show MTBF Plot

This option shows or hides the MTBF plot. See [“MTBF Plot”](#) on page 271.

Show Intensity Plot

The Intensity plot shows the estimated intensity function and confidence bounds over the design phases. The intensity function is generally discontinuous at a phase transition. Color coding facilitates differentiation of phases. If Time to Event Format is used, the axes are logarithmically scaled. See [“Show Intensity Plot”](#) on page 272.

Show Cumulative Events Plot

The Cumulative Events plot shows the estimated cumulative number of events, along with confidence bounds, over the design phases. The model requires that the cumulative events at the end of one phase match the number at the beginning of the subsequent phase. Color coding facilitates differentiation of phases. If Time to Event Format is used, the axes are logarithmically scaled. See [“Show Cumulative Events Plot”](#) on page 272.

Show Profilers

Three profilers are displayed, showing estimated MTBF, Failure Intensity, and Cumulative Events. These profilers do not use logarithmic scaling. For more information about interpreting and using these profilers, see the section [“Show Profilers”](#) on page 273.

It is important to note that, due to the default resolution of the profiler plot, discontinuities are not displayed clearly in the MTBF or Failure Intensity Profilers. In the neighborhood of a phase transition, the profiler trace shows a nearly vertical, but slightly sloped, line; this line represents a discontinuity (Figure 10.19). Such a line at a phase transition should not be used for estimation. Obtain a higher-resolution display to make these lines appear more vertical:

1. Press Ctrl while clicking in the profiler plot.
2. Enter a larger value for Number of Plotted Points in the window. (See Figure 10.20, where we have specified 500 as the Number of Plotted Points.)

Figure 10.19 Profilers

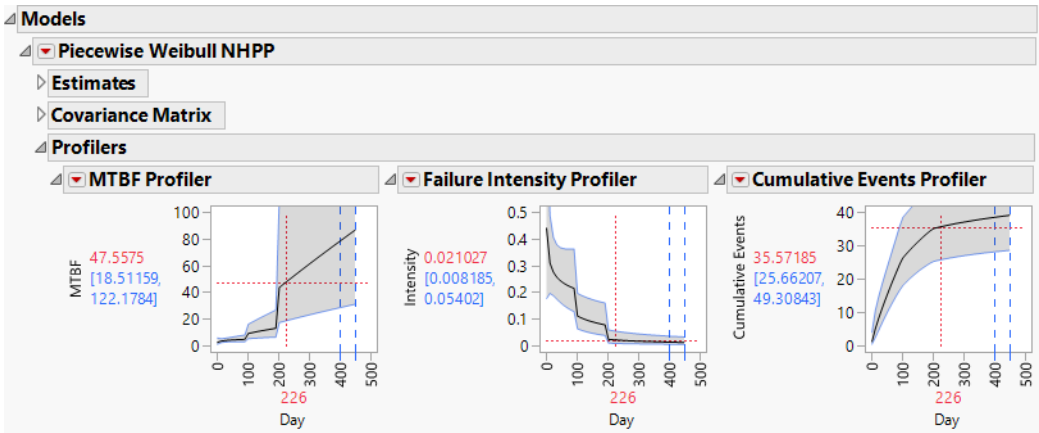


Figure 10.20 Factor Settings Window

The figure shows a 'Factor Settings' window for the 'Day' factor. The window has a title bar 'Factor Settings' and a close button. The 'Factor' is set to 'Day'. The 'Current Value' is 226. The 'Minimum Setting' is 2. The 'Maximum Setting' is 450. The 'Number of Plotted Points' is 500. The 'Show' checkbox is checked. The 'Lock Factor Setting' checkbox is unchecked. There are 'OK' and 'Cancel' buttons at the bottom.

Reinitialized Weibull NHPP

This option fits an independent growth model to the data from each test phase. Fitting models in this fashion can be useful when the factors influencing the growth rate, either in terms of testing or engineering, have changed substantially between phases. In such a situation, you might want to compare the test phases independently. The Reinitialized Weibull NHPP option is available when a Phase column specifying at least two phases has been entered in the launch window.

For the algorithm to fit this model, each row that contains the first occurrence of a new phase must contain the start date.

- Suppose that a single column is entered as Time to Event or Timestamp. Then the start time for a new phase, with a zero Event Count, must appear in the first row for that phase. See the sample data table `ProductionEquipment.jmp`, in the Reliability subfolder, for an example.
- If two columns are entered, then an interval whose left endpoint is that start time must appear in the first row, with the appropriate event count. The sample data table

TurbineEngineDesign2.jmp, found in the Reliability subfolder, provides an example. Also, see [“Piecewise NHPP Weibull Model Fitting with Interval-Censored Data”](#) on page 286.

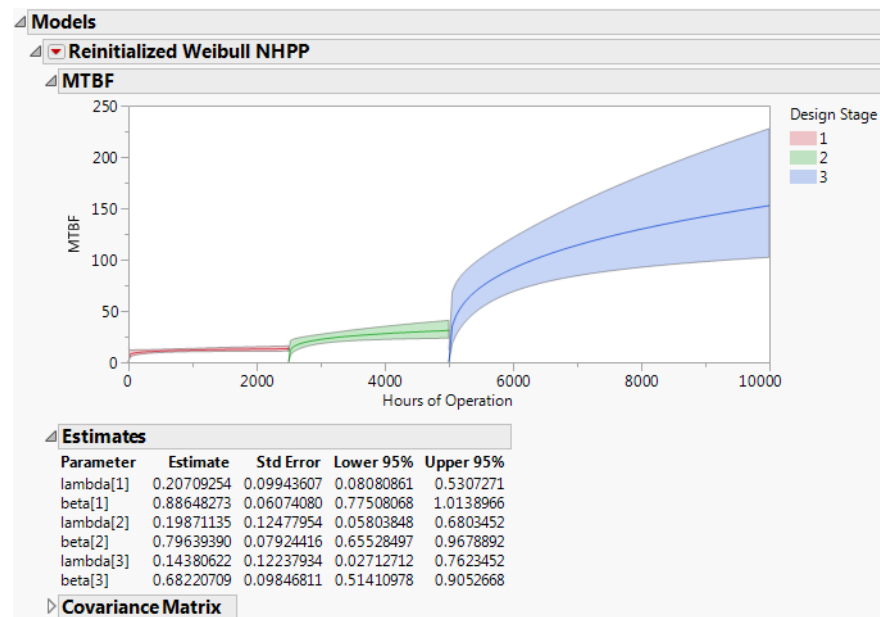
See [“Multiple Test Phases”](#) on page 264.

Independent Crow-AMSAA models are fit to the data from each of the phases. When the report is run, the Cumulative Events plot updates to show the reinitialized models. Blue vertical dashed lines show the transition points for each of the phases. The Model List also updates.

Reinitialized Weibull NHPP Report

By default, the Reinitialized Weibull NHPP report shows the estimated MTBF plot. The phases are differentiated with different colors. The Estimates and Covariance Matrix reports are shown below the plot. (See Figure 10.21, which uses the ProductionEquipment.jmp sample data file from the Reliability subfolder.)

Figure 10.21 Reinitialized Weibull NHPP Report



MTBF Plot

The MTBF plot opens by default when the Reinitialized Weibull NHPP option is chosen. For more information s about the plot, see [“MTBF Plot”](#) on page 271.

Estimates

The Estimates report gives estimates of λ and β for each of the phases. For a given phase, λ and β are estimated using only the data from that phase. The calculations assume that the phase begins at time 0 and reflect whether the phase is failure or time terminated, as defined by the data table structure. (See [“Test Phases”](#) on page 263.) Also shown are standard errors and 95% confidence limits. These values are computed as described in [“Parameter Estimates for Crow-AMSAA Models”](#) on page 291.

Covariance Matrix

Estimated covariance matrix for the estimates of the parameters of the fitted model. This report is closed by default.

Reinitialized Weibull NHPP Options

This section describes the options available in the Reinitialized Weibull NHPP red triangle menu.

Show MTBF Plot

This option shows or hides the MTBF plot. See [“MTBF Plot”](#) on page 271.

Show Intensity Plot

The Intensity plot shows the estimated intensity functions for the phases, along with confidence bands. Because the intensity functions are computed based only on the data within a phase, they are discontinuous at phase transitions. Color coding facilitates differentiation of phases. See [“Show Intensity Plot”](#) on page 272.

Show Cumulative Events Plot

The Cumulative Events plot for the Reinitialized Weibull NHPP model portrays the estimated cumulative number of events, with confidence bounds, over the design phases in the following way. Let t represent the time since the first phase of testing began. The model for the phase that is in effect at time t is evaluated at time t . In particular, the model for the phase that is in effect is not evaluated at the time since the beginning of the specific phase; rather it is evaluated at the time since the beginning of the *first* phase of testing.

At phase transitions, the cumulative events functions are discontinuous. The Cumulative Events plot matches the estimated cumulative number of events at the beginning of one phase to the cumulative number at the end of the previous phase. This matching allows the user to compare the observed cumulative events to the estimated cumulative events functions. Color coding facilitates differentiation of phases.

Show Profilers

Three profilers are displayed, showing estimated MTBF, Failure Intensity, and Cumulative Events. Note that the Cumulative Events Profiler is constructed as described in the Cumulative Events Plot section. In particular, the cumulative number of events at the beginning of one phase is matched to the number at the end of the previous phase. See [“Show Profilers”](#) on page 281.

Piecewise Weibull NHPP Change Point Detection

The Piecewise Weibull NHPP Change Point Detection option attempts to find a time point where the reliability model changes. This might be useful if you suspect that a change in reliability growth has occurred over the testing period. Note that detection seeks only a single change point, corresponding to two potential phases.

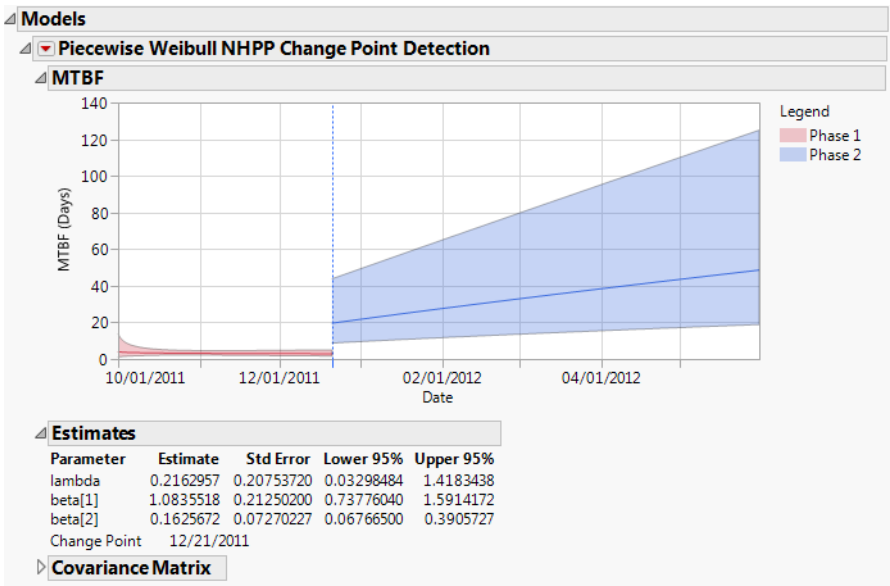
This option is available only when a Phase has *not* been entered in the launch window and one of the following is true:

- a single column has been entered as Time to Event or Timestamp in the launch window (indicating that failure times are exact), and
- two columns have been entered as Time to Event in the Concurrent Systems panel of the launch window.

When the Piecewise Weibull NHPP Change Point Detection option is selected, the estimated model plot and confidence bands are added to the Cumulative Events report under Observed Data. The Model List updates, giving statistics that are conditioned on the estimated change point. Under Models, a Piecewise Weibull NHPP Change Point Detection report is provided.

The default Piecewise Weibull NHPP Change Point Detection report shows the MTBF plot, Estimates, and Covariance Matrix. (See Figure 10.22, which uses the data in `BrakeReliability.jmp`, found in the Reliability subfolder.) Note that the Change Point, shown at the bottom of the Estimates report, is estimated as 12/21/2011. The standard errors and confidence intervals consider the change point to be known. The plot and the Estimates report are described in the section [“Piecewise Weibull NHPP”](#) on page 278.

Figure 10.22 Piecewise Weibull NHPP Change Point Detection Report



Available options are: Show MTBF Plot, Show Intensity Plot, Show Cumulative Events Plot, Show Profilers. These options are described in the section [“Reinitialized Weibull NHPP Options”](#) on page 284.

The procedure used in estimating the change point is described in [“Statistical Details for the Piecewise Weibull NHPP Change Point Detection Report”](#) on page 292.

Additional Examples of the Reliability Growth Platform

- [“Piecewise NHPP Weibull Model Fitting with Interval-Censored Data”](#)
- [“Piecewise Weibull NHPP Change Point Detection with Time in Dates Format”](#)

Piecewise NHPP Weibull Model Fitting with Interval-Censored Data

The sample data file TurbineEngineDesign2.jmp, found in the Reliability subfolder, contains data on failures for a turbine engine design over three phases of a testing program. The first two columns give time intervals during which failures occurred. These intervals are recorded as days since the start of testing. The exact failure times are not known; it is known only that failures occurred within these intervals.

The reports of failures are provided generally at weekly intervals. Intervals during which there were no failures and which fell strictly within a phase are not included in the data table (for example, the interval 106 to 112 is not represented in the table). Because these make no contribution to the likelihood function, they are not needed for estimation of model parameters.

However, to fit a Piecewise Weibull NHPP or Reinitialized Weibull NHPP model, it is important that the start times for all phases be provided in the Time to Event or Timestamp columns.

Here, the three phases began at days 0 (Initial phase), 91 (Revised phase), and 200 (Final phase). There were failures during the weeks that began the Initial and Revised phases. However, no failures were reported between days 196 and 231. For this reason, an interval with beginning and ending days equal to 200 was included in the table (row 23), reflecting 0 failures. This is necessary so that JMP knows the start time of the Final phase.

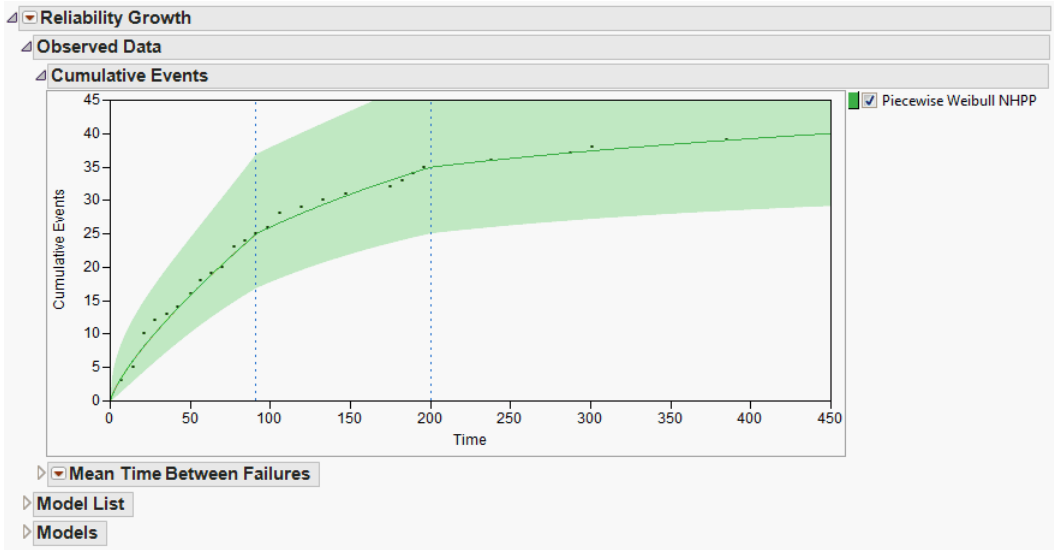
The test was terminated at 385 days. This is an example of interval-censored failure times with time terminated phases.

Note: The phase start times are required for proper display of the transition times for the Piecewise Weibull NHPP model; they are required for estimation of the Reinitialized Weibull NHPP model. For interval-censored data, the algorithm defines the beginning time for a phase as the start date recorded in the row containing the first occurrence of that phase designation. In our example, if row 23 were not in the table, the beginning time of the Final phase would be taken as 231.

1. Select **Help > Sample Data Library** and open Reliability/TurbineEngineDesign2.jmp.
2. Select **Analyze > Reliability and Survival > Reliability Growth**.
3. On the **Time to Event Format** tab, select the columns Interval Start and Interval End, and click **Time to Event**.
4. Select Fixes and click **Event Count**.
5. Select Design Phase and click **Phase**.
6. Click **OK**.
7. Click the Reliability Growth red triangle and select **Fit Model > Piecewise Weibull NHPP**.

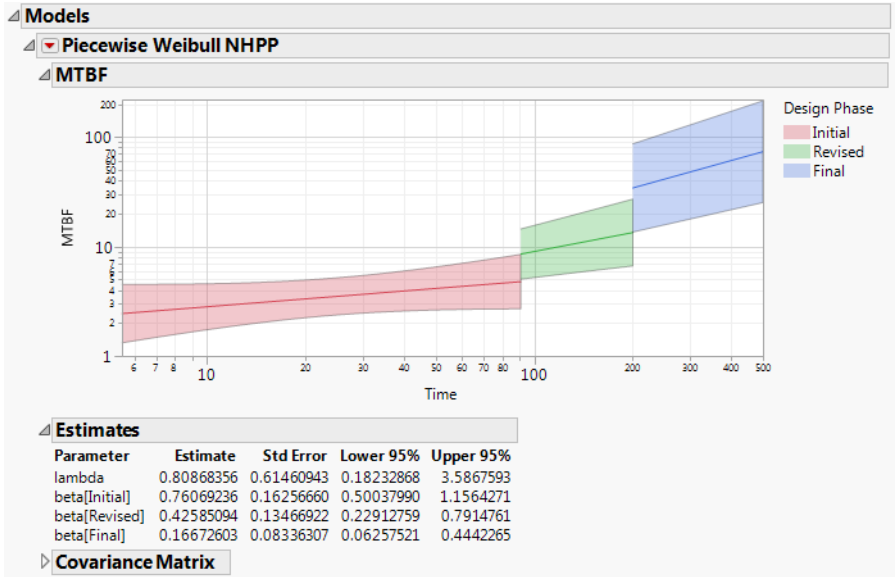
The Cumulative Events plot from the Observed Data report is shown in Figure 10.23. The vertical dashed blue lines indicate the phase transition points. The first occurrence of Revised in the column Design Phase is in row 14. So, the start of the Revised phase is taken to be the Interval Start value in row 14, namely, day 91. Similarly, the first occurrence of Final in the column Design Phase is in row 23. So, the start of the Final phase is taken to be the Interval Start value in row 23, namely, day 200.

Figure 10.23 Cumulative Events Plot



The Piecewise Weibull NHPP report is found under the Models outline node. Here we see the mean time between failures increasing over the three phases. From the Estimates report, we see that the estimates of beta decrease over the three testing phases.

Figure 10.24 MTBF Plot



Piecewise Weibull NHPP Change Point Detection with Time in Dates Format

The file BrakeReliability.jmp, found in the Reliability subfolder, contains data on fixes to a braking system. The Date column gives the dates when Fixes, given in the second column, were implemented. For this data, the failure times are known. Note that the Date column must be in ascending order.

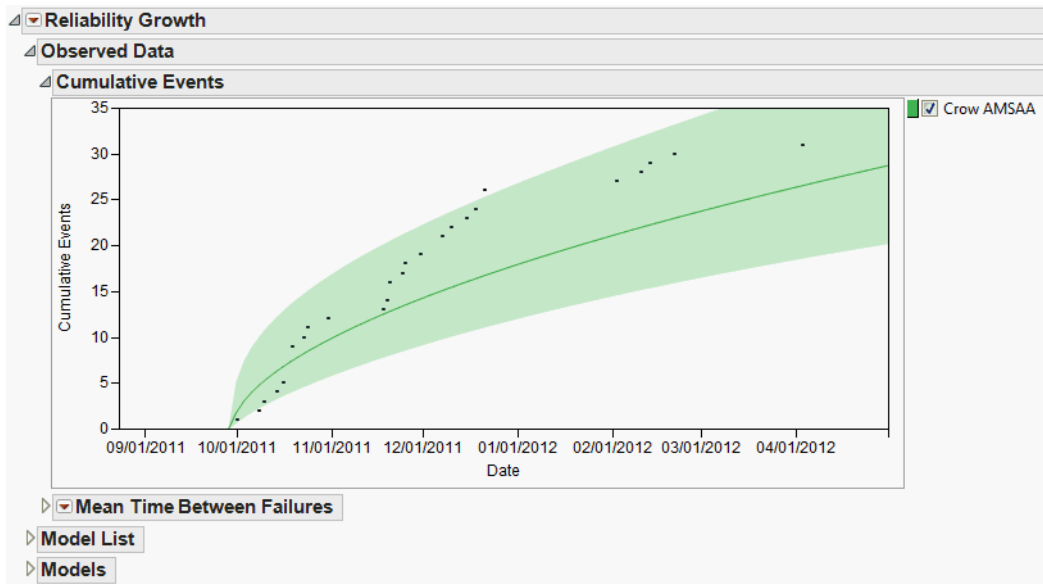
The test start time is the first entry in the Date column, 09/29/2011, and the corresponding value for Fixes is set at 0. This is needed in order to convey the start time for testing. If there had been a nonzero value for Fixes in this first row, the corresponding date would have been treated as the test start time. However, the value of Fixes would have been treated as 0 in the analysis.

The test termination time is given in the last row as 05/31/2012. Because the value in Fixes in the last row is 0, the test is considered to be time terminated on 5/31/2012. If there had been a nonzero value for Fixes in this last row, the test would have been considered failure terminated.

1. Select **Help > Sample Data Library** and open Reliability/BrakeReliability.jmp.
2. Select **Analyze > Reliability and Survival > Reliability Growth**.
3. Select the **Dates Format** tab.
4. Select Date and click **Timestamp**.
5. Select Fixes and click **Event Count**.
6. Click **OK**.
7. Click the Reliability Growth red triangle and select **Fit Model > Crow AMSAA**.

The Cumulative Events plot in the Observed Data report updates to show the model. The model does not seem to fit the data very well.

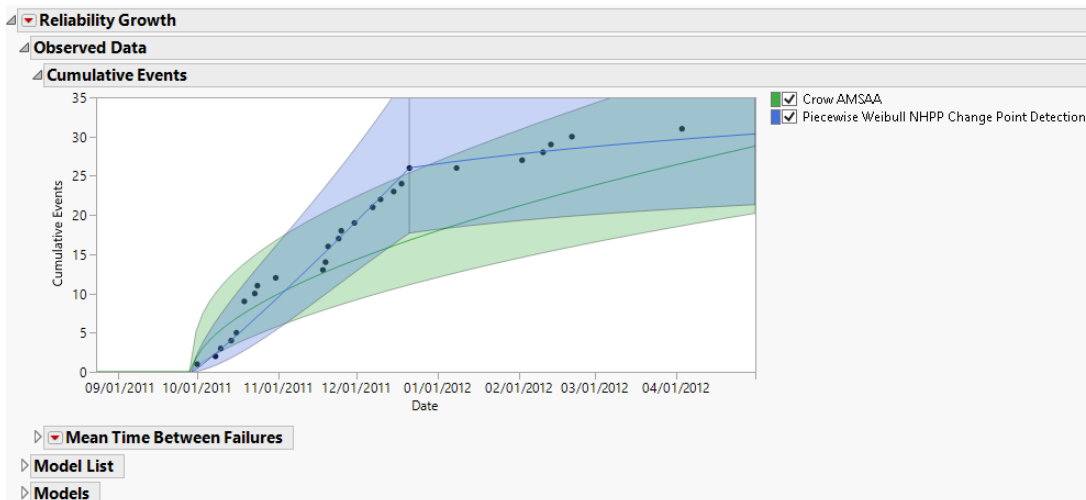
Figure 10.25 Cumulative Events Plot with Crow AMSAA Model



- Click the Reliability Growth red triangle and select **Fit Model > Piecewise Weibull NHPP Change Point Detection**.

The Cumulative Events plot in the Observed Data report updates to show the piecewise model fit using change-point detection. Both models are shown in Figure 10.26. Though the data are rather sparse, the piecewise model seems to provide a better fit to the data.

Figure 10.26 Cumulative Events Plot with Two Models



Statistical Details for the Reliability Growth Platform

- [“Statistical Details for the Crow-AMSAA Report”](#)
- [“Statistical Details for the Piecewise Weibull NHPP Change Point Detection Report”](#)

Statistical Details for the Crow-AMSAA Report

This section contains details for the parameter estimates and profilers that appear in the Crow-AMSAA Report.

Parameter Estimates for Crow-AMSAA Models

With the exception of the Crow-AMSAA with Modified MLE option, the estimates for λ and β are maximum likelihood estimates, which are computed as follows. The likelihood function is derived using the methodology in Meeker and Escobar (1998). It is reparametrized in terms of $\text{param}_1 = \log(\lambda)$ and $\text{param}_2 = \log(\beta)$. This is done to enable the use of an unconstrained optimization algorithm, namely, an algorithm that searches from $-\infty$ to $+\infty$. The MLEs for param_1 and param_2 are obtained.

The standard errors for λ and β are obtained from the Fisher information matrix. Confidence limits for param_1 and param_2 are calculated based on the asymptotic distribution of the MLEs, using the Wald statistic. These estimates and their confidence limits are then transformed back to the original units using the exponential function.

Parameter Estimates for Crow-AMSAA with Modified MLE

For the Crow AMSAA with Modified MLE option, the estimate for β is corrected for bias. The formula for the bias-corrected estimate of β depends on whether the test is failure terminated or time terminated.

Denote the MLE for β by $\hat{\beta}$, let n be the number of observations, and let T be the total test time.

The bias-corrected estimate (*modified MLE*) of β is $\bar{\beta}$, where:

$$\bar{\beta} = \left(\frac{n-2}{n} \right) \hat{\beta} \text{ for a failure-terminated test}$$

$$\bar{\beta} = \left(\frac{n-1}{n} \right) \hat{\beta} \text{ for a time-terminated test}$$

The modified MLE for λ , denoted $\bar{\lambda}$, is calculated according to the expression given by the likelihood function, but based on the adjusted value of beta:

$$\bar{\lambda} = n/T^{\bar{\beta}}$$

The covariance matrix for the parameters is estimated using the Fisher information matrix. (See “[Parameter Estimates for Crow-AMSAA Models](#)” on page 291.) However, the bias-corrected estimates for λ and β are substituted for the MLEs in the resulting formulas. All confidence bands in plots and confidence intervals in reports are based on this procedure.

Profilers

For the Crow-AMSAA models, the estimates for the MTBF, Intensity, and Cumulative Events given in the profilers are obtained by replacing the parameters λ and β in their theoretical expressions by their MLEs. In the case of the Crow-AMSAA with Modified MLE option, the modified MLEs are used. Confidence limits are obtained by applying the delta method to the log of the expression of interest.

For example, consider the cumulative events function. The cumulative number of events at time t since testing initiation is given by $N(t) = \lambda t^{\beta}$. It follows that $\log(N(t)) = \log(\lambda) + \beta \log(t)$. The parameters λ and β in $\log(N(t))$ are replaced by their MLEs (or modified MLEs) to estimate $\log(N(t))$. The delta method is applied to this expression to obtain an estimate of its variance. This estimate is used to construct a 95% Wald-based confidence interval. The resulting confidence limits are then transformed using the exponential function to give confidence limits for the estimated cumulative number of events at time t .

Statistical Details for the Piecewise Weibull NHPP Change Point Detection Report

The change point is estimated as follows:

- Using consecutive event times, disjoint intervals are defined.
- Each point within such an interval can be considered to be a change point defining a piecewise Weibull NHPP model with two phases. So long as the two phases defined by that point each consist of at least two events, the algorithm can compute MLEs for the parameters of that model. The log-likelihood for that model can also be computed.
- Within each of the disjoint intervals, a constrained optimization routine is used to find a local optimum for the log-likelihood function.
- These local optima are compared, and the point corresponding to the largest is chosen as the estimated change point.

Note that this procedure differs from the grid-based approach described in Guo et al. (2010).

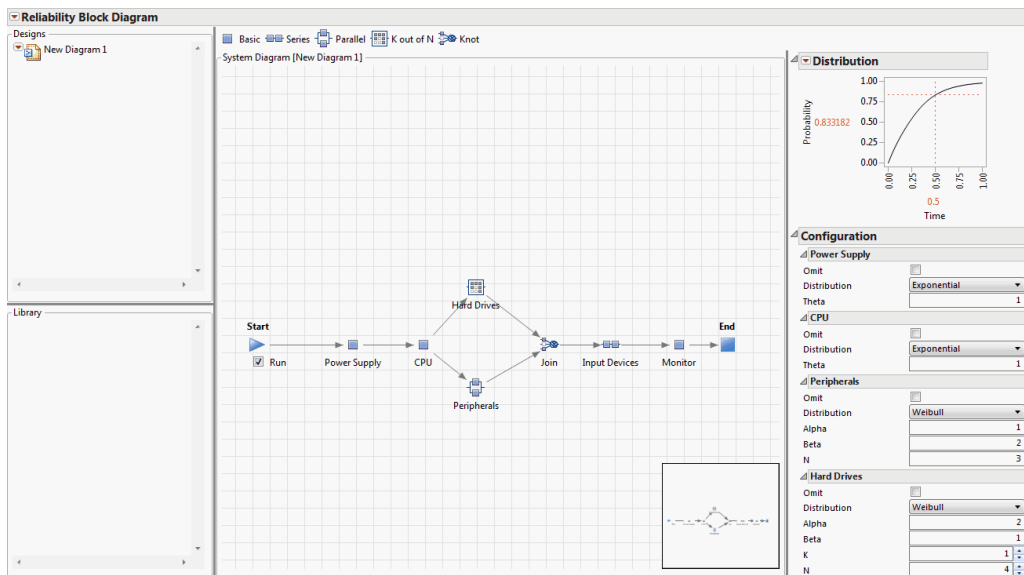
Chapter 11

JMP[®] PRO Reliability Block Diagram Engineer System Reliabilities

The Reliability Block Diagram platform is available only in JMP Pro.

The Reliability Block Diagram platform displays the reliability relationships among a system's components. If reliability distributions are assigned to the components, the platform analytically models the reliability behavior. A Reliability Block Diagram (also known as a dependence diagram) illustrates how component reliability contributes to the success or failure of a system.

Figure 11.1 Example of a Reliability Block Diagram



Contents

Overview of the Reliability Block Diagram Platform	295
Example Using the Reliability Block Diagram Platform	295
Add Components	297
Align Shapes	299
Connect Shapes	300
Configure Components	301
The Reliability Block Diagram Window	303
Preview Window	304
Reliability Block Diagram Platform Options	305
Workspace Options	308
Configuration Settings	309
Distribution Configurations	309
Specify a Nonparametric Distribution	311
Options for Design and Library Items	312
Profilers	313
Component Importance and Time to Failure	315
Component Plots	317
Print Algebraic Reliability Formula	320
Generate Algebraic Expression Data Table	321
Clone and Delete	321

Overview of the Reliability Block Diagram Platform

The Reliability Block Diagram platform enables you to diagram a system and its related components to show how component reliability affects the success of the whole system. Each block in a Reliability Block Diagram represents a component in the system and is connected to other components in the system.

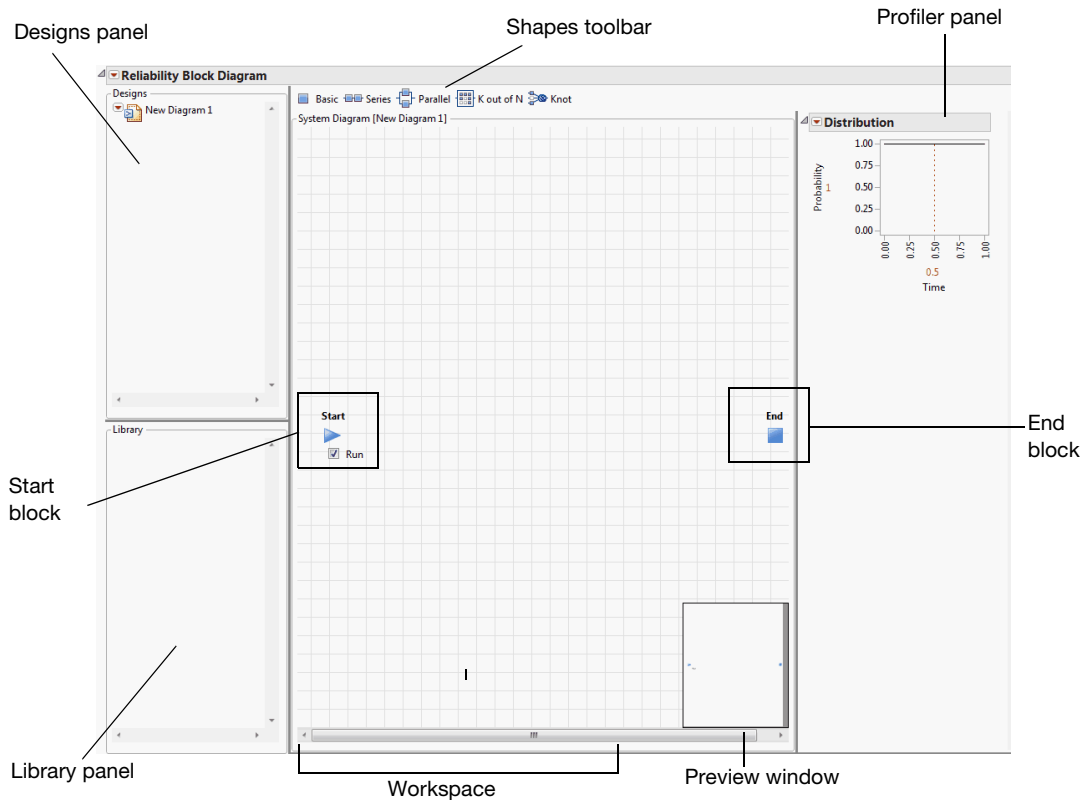
A Reliability Block Diagram can include components connected either in series or in parallel. A failure in any series component causes the entire system to fail. In a parallel-connected (or redundant) system, all parallel components must fail for the entire system to fail. In addition, the diagram can include K out of N components where k -out-of- n components must function for the system to function.

The Reliability Block Diagram template places a Start block on the left side and an End block on the right side of the diagram. Using the Shape tools, you diagram the system to be analyzed beginning at the Start block and connecting components to reach the End block.

Example Using the Reliability Block Diagram Platform

In this example, you learn how to create a new Reliability Block Diagram.

1. Select **Analyze > Reliability and Survival > Reliability Block Diagram**.
A blank Reliability Block Diagram window appears.

Figure 11.2 New Reliability Block Diagram


Note: The Distribution profiler appears by default.

2. In the Designs panel, select and rename New Diagram 1 to Computer.

The Workspace is now named System Diagram [Computer].

3. Deselect **Run** in the Start block.

With **Run** selected, the platform updates the diagram's reliability calculations after each change to the system diagram. These changes can include adding or deleting components, changing a component's configuration, and adding or deleting a connection.

With **Run** deselected, the platform does not update the reliability calculations after any changes.

Tip: Deselect **Run** when you are diagramming large systems. Select **Run** when the diagram is complete.

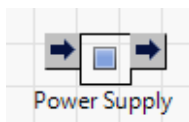
4. Proceed with ["Add Components"](#) on page 297.

JMP[®] PRO Add Components

The Reliability Block Diagram drawing elements that are located in the toolbar are called *shapes*. The term *component* refers to a shape that represents a constituent part of the system.

1. Click the Basic icon  **Basic** on the Shape toolbar and drag the shape to the System Diagram to the right of the Start block.
2. Select the label, replace New Basic 1 by Power Supply, and press Enter.

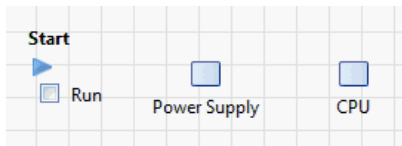
Figure 11.3 Basic Shape



When you click the label or the shape, connection arrows appear. The arrows disappear when you click elsewhere in the template.

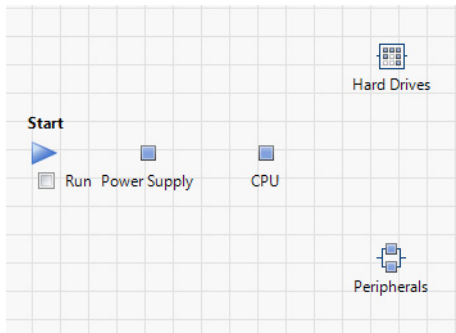
3. Drag a second Basic shape to the right of the Power Supply shape.
4. Select the label and type CPU.

Figure 11.4 Example System Diagram

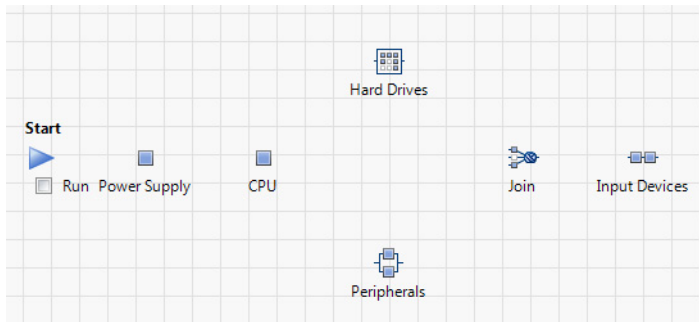


Note: You will align the shapes later, in the section [“Align Shapes”](#) on page 299.

5. Drag a Parallel shape  **Parallel** to a position to the right and below the CPU shape.
6. Select the label and type Peripherals.
7. Drag a K out of N shape  **K out of N** to a position to the right and above the CPU shape.
8. Select the label and type Hard Drives.

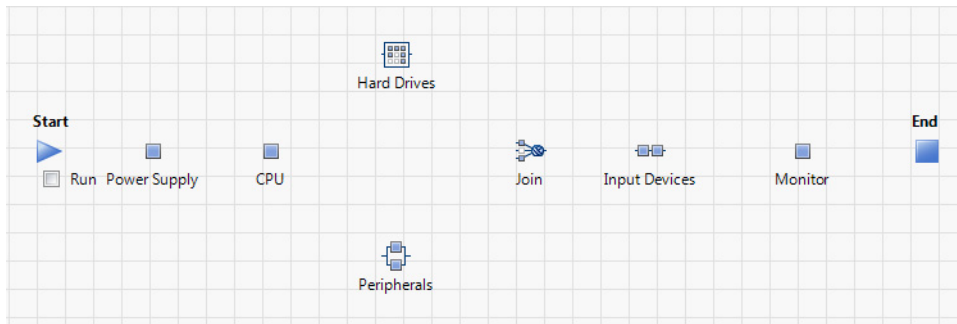
Figure 11.5 Partial System Diagram

9. Drag a Knot shape  to the right of the previous shapes.
10. Select the label and type Join.
11. Drag a Series shape  to a position to the right of the Knot shape.
12. Select the label and type Input Devices.

Figure 11.6 Partial System Diagram

13. Drag a Basic shape to a position to the right of the Input Devices shape.
14. Select the label and type Monitor.

Figure 11.7 System Diagram Showing All Shapes



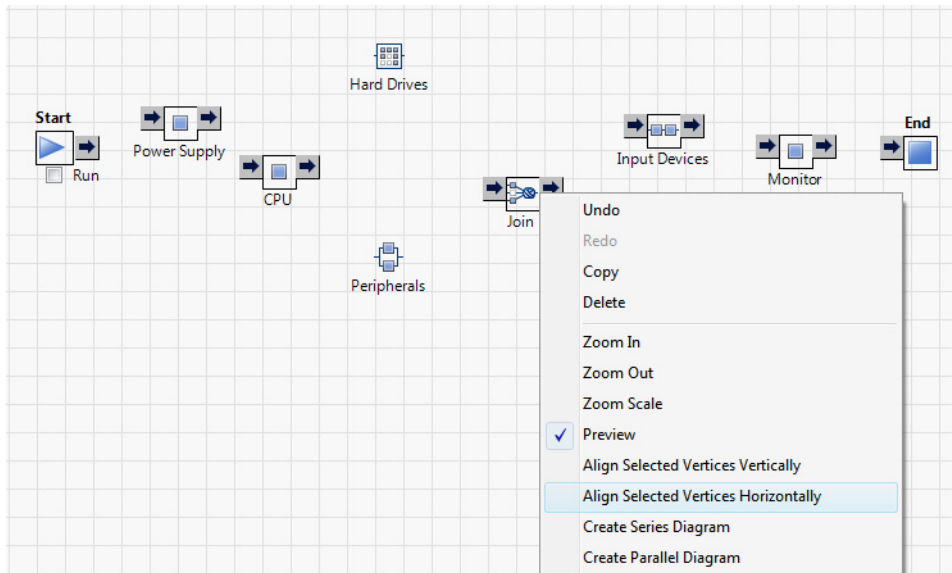
15. Proceed with [“Align Shapes”](#) on page 299.

JMP[®] PRO Align Shapes

1. To vertically align the shapes for Hard Drives and Peripherals, select the components:
 - Hard Drives
 - Peripherals

Tip: To select shapes, drag the cursor around the shapes or press Shift and click each shape.

2. With the shapes selected, right-click one of the shapes and select **Align Selected Vertices Vertically**.
3. To horizontally align the remaining shapes, select the following components:
 - Start
 - Power Supply
 - CPU
 - Join
 - Input Devices
 - Monitor
 - End
4. With the shapes selected, right-click one of the shapes and select **Align Selected Vertices Horizontally**.

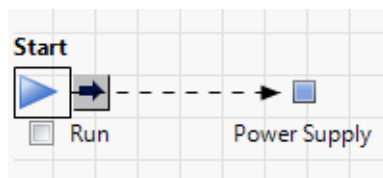
Figure 11.8 Align Shapes Horizontally


5. Proceed with [“Connect Shapes”](#) on page 300.

JMP PRO Connect Shapes

To connect shapes, select a shape to display its connection arrows. Suppose you want to connect shape A to shape B. Select shape A. Drag the right arrow to shape B to indicate that shape A precedes shape B. Drag the left arrow to shape B to indicate that shape B precedes shape A. To connect the shapes in your diagram, select the right arrows to connect to the next shape in the sequence.

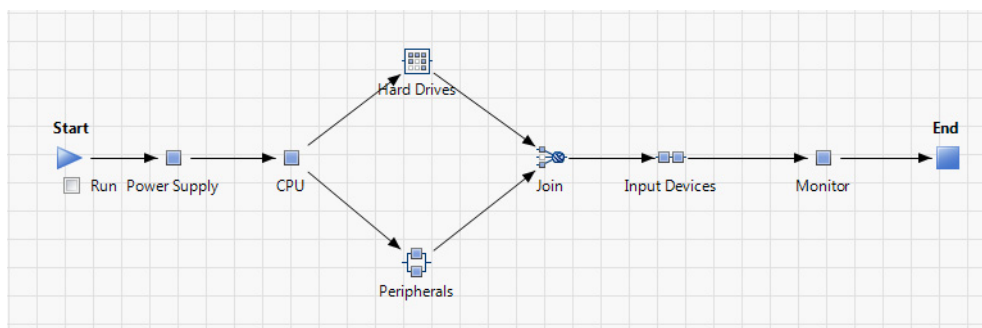
1. Select the Start block (blue arrow) to display the connection arrow.
2. Select the single connection arrow  and drag it to the Power Supply component.

Figure 11.9 Connecting Shapes


3. For each of the following components, click the first component, select its right connection arrow, and drag the arrow to the second component:
 1. Power Supply → CPU

2. CPU → Hard Drive
3. CPU → Peripherals
4. Hard Drives → Join
5. Peripherals → Join
6. Join → Input Devices
7. Input Devices → Monitor
8. Monitor → End block

Figure 11.10 Completed System Diagram



4. Proceed with “Configure Components” on page 301.

JMP PRO Configure Components

1. In the Configuration panel, enter the Configuration settings for the components. See “Configuration Settings” on page 309.

Table 11.1 Configuration Settings

Component	Settings
Power Supply	• Distribution—Exponential • Theta—1
CPU	• Distribution—Exponential • Theta—1

Table 11.1 Configuration Settings (Continued)

Component	Settings
Peripherals	- Distribution—Weibull - Alpha—1 - Beta—2 - N—3
Hard Drives	- Distribution—Weibull - Alpha—2 - Beta—1 - K—1 - N—4
Join	Minimum available—1
Input Devices	- Distribution—Fréchet - Location—0 - Scale—1 - N—2
Monitor	- Distribution—Exponential - Theta—1

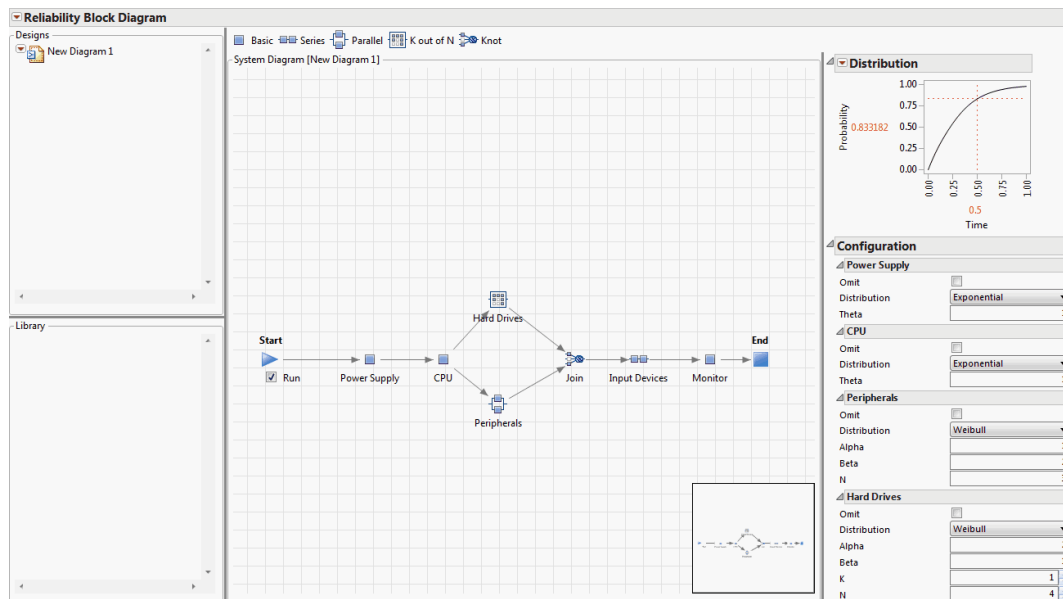
The Reliability Block Diagram is complete (Figure 11.11).

2. Select **Run**.

The system’s reliability information is updated. This is reflected in the Distribution plot in the Profiler pane.

3. To save the Reliability Block Diagram as a JMP Scripting Language (JSL) file, select **File > Save** and name it `exampleRBDcomplete.jsl`.

Figure 11.11 Example Reliability Block Diagram



JMP PRO The Reliability Block Diagram Window

The Reliability Block Diagram window is divided into the following panels:

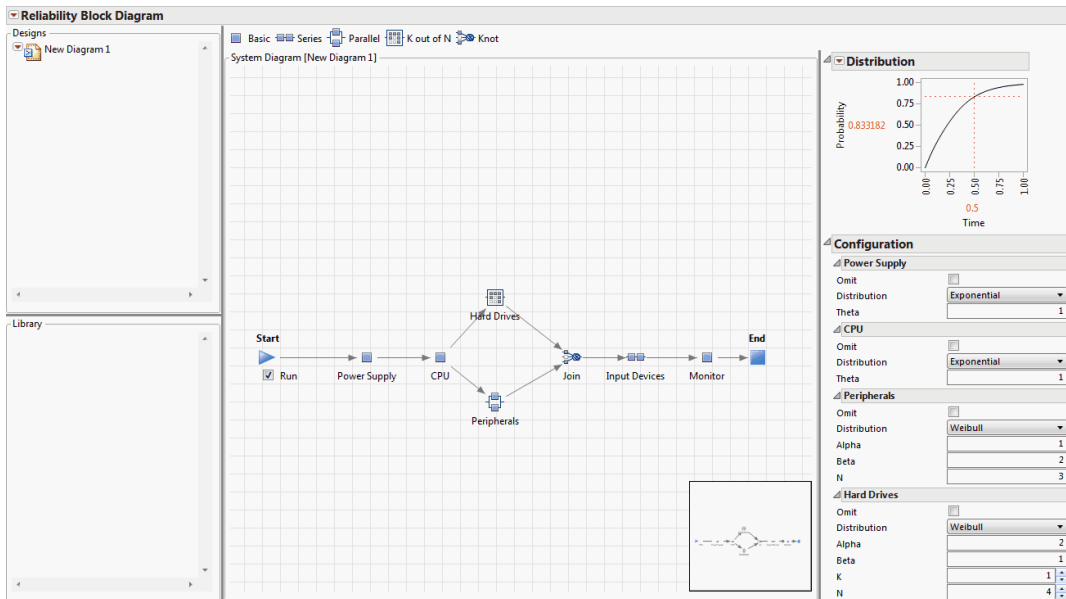
- **Designs:** Lists system diagrams that were created using the Reliability Block Diagram platform.
- **Library:** Lists sub-system designs that are available for reuse in creating new system diagrams.

Note: Each system diagram that appears in the Designs and the Library panels contains its own red triangle menu. The options in each of those red triangle menus are described in [“Options for Design and Library Items”](#) on page 312.

- **Workspace:** Displays the Shape toolbar, the System Diagram, and the Preview window.

Tip: To hide the Preview window, right-click in the Workspace and deselect Preview.

- **Profiler Panel:** Displays the Distribution profiler, Configuration settings, and various system and component profilers and plots that you select from the Diagram red triangle menu.

Figure 11.12 Example of a Reliability Block Diagram


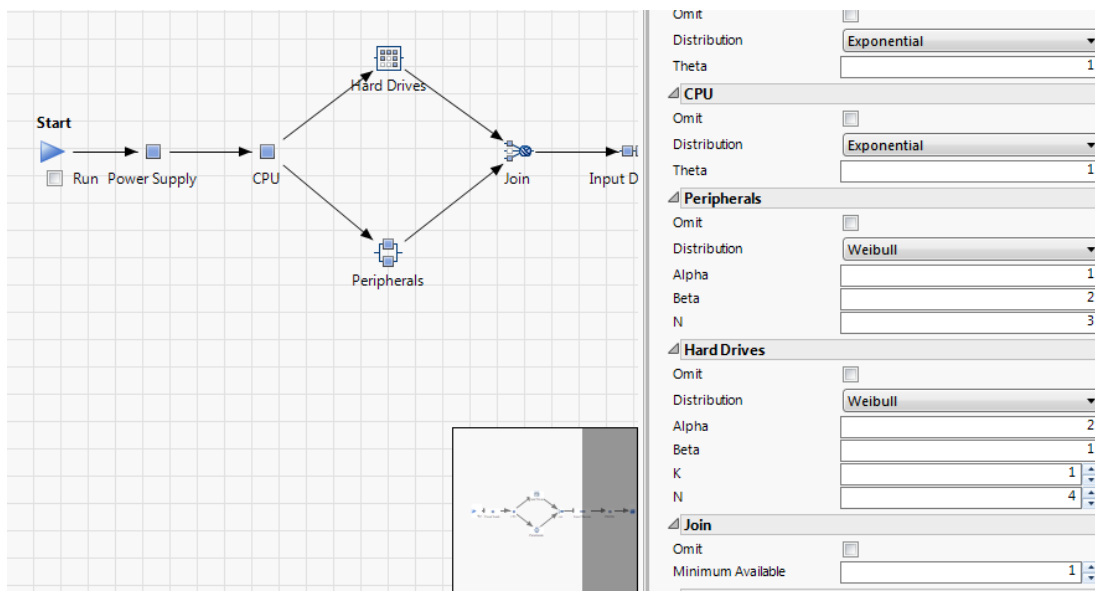
The Shape toolbar includes the following drawing tools:

-  **Basic** Adds a single block shape to the system diagram.
-  **Series** Adds a series block shape that represents a group of components connected in a series. All components must function for the system to function.
-  **Parallel** Adds a parallel block shape that represents a group of components that are in parallel. At least one of the components must function for the system to function.
-  **K out of N** Adds a k -out-of- n block shape where you specify k and n . At least k of the components must function for the system to function.
-  **Knot** Adds a knot shape to the system diagram that enables you to join Parallel or K out of N components that have different Distribution property settings.

JMP PRO Preview Window

When the system diagram is too large to appear in the workspace, the Preview window enables you to reposition the portion of the system diagram that appears in the workspace. The viewing area is indicated in the Preview window by a white background. Drag the viewing area to reposition the view of the Workspace to another part of the system diagram.

Figure 11.13 Preview Window with Visible Part of Diagram on White Background



JMP PRO Reliability Block Diagram Platform Options

The Reliability Block Diagram red triangle menu contains the following options for the Reliability Block Diagram window:

Save and Save As Enables you to save a Reliability Block Diagram, or save an existing Reliability Block Diagram with a new name, to a JMP Scripting Language (JSL) script that is automatically executed when it is opened in JMP. See the *Scripting Guide* for more information about Auto-Submit scripts.

Note: The Save and Save As red triangle options are equivalent to File > Save and File > Save As. They are available in the red triangle menu for convenience.

Show Design Diagram If **Show Design Comparisons** is selected, this option enables you to hide or show the system diagram in the Workspace.

Show Design Comparisons Displays a Distribution Overlay profiler and Remaining Life Distribution Overlay profiler for the selected system diagrams. See [“Show Design Comparisons”](#) on page 306.

Import Component Distribution Settings Enables you to import configuration settings for the system diagram from a data table. The table must contain columns for the diagram category, diagram name, component name, distribution, and one or more parameters. The number of parameters depends on the specified distribution.

Note: The strings in the imported table must be exact matches to the strings in the system diagram.

New Design Item Enables you to add a new design diagram to the Designs panel. See [“Add a New Design Item”](#) on page 307.

New Library Item Enables you to add a new item to the Library panel. The Library panel lets you create subsystem diagrams for use in multiple system designs. See [“Add a New Library Item”](#) on page 308.

Note: Library items are available only for use in diagrams that are contained in the Designs panel of a JSL file. They are not available to other diagrams in other JSL files.

Show Design Comparisons

The **Show Design Comparisons** option displays a Distribution Overlay plot and a Remaining Life Distribution Overlay plot. These plots appear in the Workspace below the System Diagram. See Figure 11.14 and Figure 11.15.

- The Distribution Overlay plot shows the probability of system failure for each of the designs in the Designs panel.
- The Remaining Life Distribution Overlay plot shows the conditional probability of failure for each design in the Designs panel, given that the systems have survived to a specified survival time. Enter the Survival Time in the box to the right of the plot. Alternatively, select the small rectangle at the origin and drag it to the right to dynamically set the Survival Time.

Check boxes enable you to select which designs are represented in the plots. This enables you to compare a subset of designs with respect to failure probabilities.

Tip: To hide the system diagram in the Workspace, click the Reliability Block Diagram red triangle and deselect **Show Design Diagram**. Alternatively, to show more of the Design Comparisons, reposition the horizontal splitter upward.

Figure 11.14 Distribution Overlay Example

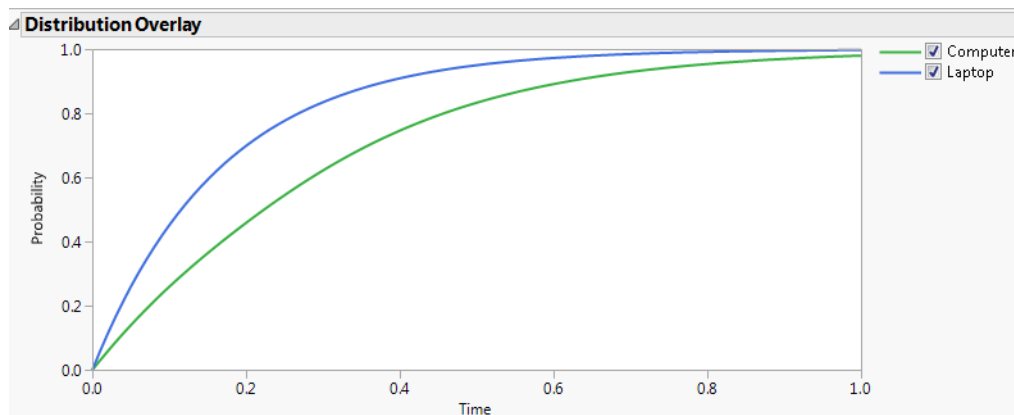
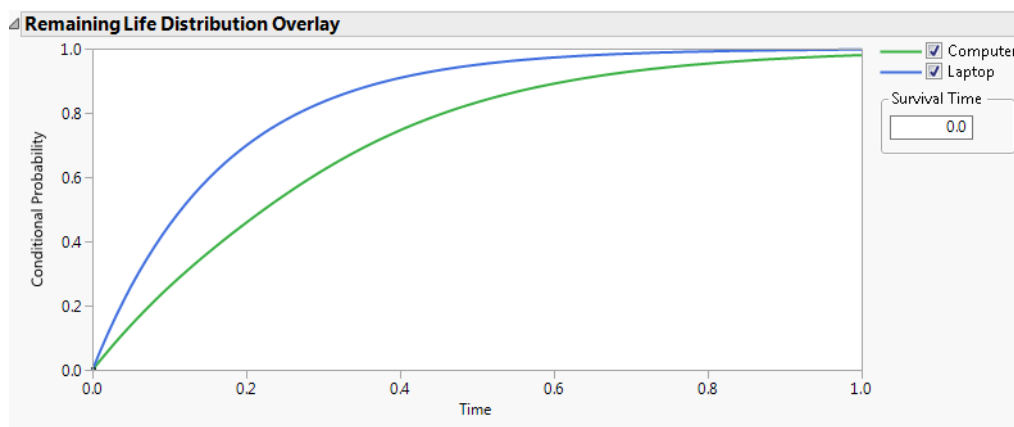


Figure 11.15 Remaining Life Distribution Overlay Example



JMP PRO Add a New Design Item

The **New Design Item** option adds a new system design to the Designs panel list.

- Click the Reliability Block Diagram red triangle and select **New Design Item**.
- A design named New Diagram X, where X is a number that identifies the diagram name, is added to the Designs panel list.
- To name the design, select the label and enter a name.
- Use the Shape toolbar to draw the system diagram. (See [“Example Using the Reliability Block Diagram Platform”](#) on page 295.)

Tip: You can also copy and paste the body of a diagram from another block diagram template. The Start and End shapes are not copied. After you paste the diagram, you must reconnect the Start and End shapes.

- Configure the components.
- Save the file.

To display a design in the System Diagram window and to view its Profiler panel, double-click its icon in the Designs panel.

JMP PRO Add a New Library Item

The **New Library Item** option adds a new sub-system to the Library panel list.

Tip: Drag a system design from the Designs panel to the Library panel to add it to the Library as a sub-system.

- Click the Reliability Block Diagram red triangle and select **New Library Item**.
- A sub-system called New Diagram X, where X is a number that identifies the diagram name, is added to the Library panel list.
- To name the sub-system, select the label and enter a name.
- Use the Shape toolbar to draw the sub-system diagram.
- Configure the components.
- Create connections from the Start block to the End block.
- Save the file.

JMP PRO Workspace Options

The Workspace panel supports several right-click commands for adjusting the view of the panel. Here are some of the available commands:

Zoom In Enables you to zoom in to the system diagram.

Zoom Out Enables you to zoom out from the system diagram.

Tip: On Windows, you can press Ctrl and use the mouse scroll wheel to zoom in to and out from the diagram.

Zoom Scale Opens the Set Zoom Scale window enabling you to zoom in or out using a scaling factor. The scaling factor ranges from 0.75 (75%) to 5.0 (500%).

Preview Enables you to hide or show the Preview window in the Workspace panel.

Align Selected Vertices Vertically Enables you to align shapes on a vertical line.

Align Selected Vertices Horizontally Enables you to align shapes on a horizontal line.

Create Series Diagram Enables you to create a series diagram using the nodes in the diagram. The order of the series is determined by the order of node creation. This option is available only when the diagram does not contain any arrows.

Create Parallel Diagram Enables you to create a parallel diagram using the nodes in the diagram. This option is available only when the diagram does not contain any arrows.

JMP PRO Configuration Settings

Each component in a Reliability Block Diagram can be assigned a failure distribution. The available failure distributions are listed in [Table 11.2](#) on page 309. To see the formulas and parameterization for these failure distributions, see “Distributions” on page 82 in the “Life Distribution” chapter.

JMP PRO Distribution Configurations

When a component is added to the diagram, an outline for that component appears beneath the Configuration outline.

Note: You can omit a component from the analysis by checking the box next to Omit.

Select a Distribution and enter the required parameter values.

Table 11.2 Distributions and Parameters

Property Type	Required Inputs
Exponential	Theta
Weibull	Alpha, Beta
Lognormal	location, scale
Loglogistic	location, scale
Fréchet	location, scale
GenGamma	mu, sigma, lambda
DS Weibull	Alpha, Beta, Defective Probability

Table 11.2 Distributions and Parameters (Continued)

Property Type	Required Inputs
DS Lognormal	location, scale, Defective Probability
DS Loglogistic	location, scale, Defective Probability
DS Fréchet	location, scale, Defective Probability
Nonparametric	data or data file

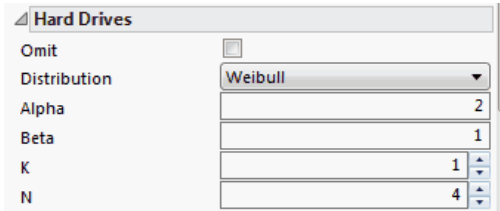
To view Configuration settings for the components in a selected design or subsystem, do the following:

- To view the Configuration settings for a specific component, select the component’s shape.
- To view Configuration settings for more than one component, select multiple components’ shapes using the Arrow tool or pressing Control and clicking.
- To view Configuration settings for all components, deselect any shapes by clicking in a blank portion of the Workspace.

For each component in the diagram, you can either omit the component from calculations by checking the box next to Omit or:

- From the Distribution list, select the appropriate distribution. For all selections other than Nonparametric, enter parameter values for the distribution. For Nonparametric, see [“Specify a Nonparametric Distribution”](#) on page 311.
- For Series and Parallel components, you must also enter a value for N, the total number of components contained in the series or parallel shape.
- For K out of N components, you must also enter K, the minimum number of components that must function for the system to function, and N, the total number of components in the shape.
- For Knot components, enter the Minimum Available Dependencies. The Knot shape enables you to configure a *k-out-of-n* shape where the shapes that are joined have different distributions. The Minimum Number of Dependencies is *k*, the minimum number of paths leading to the Knot that need to function in order for the system to function.

Figure 11.16 Example of a Weibull Configuration for a K out of N Shape

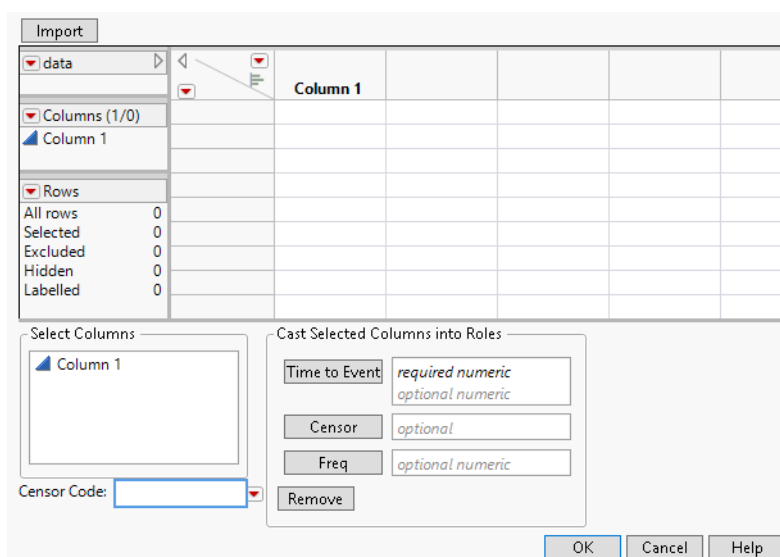


JMP PRO Specify a Nonparametric Distribution

The Nonparametric option under Distribution enables you to approximate an arbitrary distribution. Enter data or import a file containing a large set of data. This data is used to approximate the distribution.

After selecting Nonparametric, click the  icon next to Data. The Provide Data window appears, enabling you to either enter data or import a data file. Once you have imported or entered your data, the data is used to calculate a nonparametric distribution for the component.

Figure 11.17 Provide Data Window



To import data from a file, do the following.

1. Open the JMP data table that contains the data to import.
2. In the Provide Data window, click **Import**.
The Select Data Table window appears.
3. From the Data Table list, select the data table.
4. Click **OK**.
5. In the panel beneath the data grid, specify which columns represent Time to Event data, as well as Censor and Freq data if appropriate.
6. In the Provide Data window, click **OK**.

To enter data manually, do the following.

1. Create columns for Time to Event data, as well as Censor and Freq data, if appropriate.

2. Enter the data into the columns.
3. In the panel beneath the data grid, specify which columns represent Time to Event data, as well as Censor and Freq data if appropriate.
4. In the Provide Data window, click **OK**.



Options for Design and Library Items

The available options for Design and Library items are listed below:

- Show Configuration shows or hides the Configuration outline. See [“Configuration Settings”](#) on page 309.
- Options to show various profilers. See [“Profilers”](#) on page 313 and the following:
 - [“Distribution Profiler”](#) on page 313
 - [“Remaining Life Distribution Profiler”](#) on page 313
 - [“Reliability Profiler”](#) on page 314
 - [“Quantile Profiler”](#) on page 314
 - [“Density Profiler”](#) on page 314
 - [“Hazard Profiler”](#) on page 314
 - [“Cumulative Hazard Profiler”](#) on page 314
- Options to show plots for component importance measures and mean time to failure. See [“Component Importance and Time to Failure”](#) on page 315 and the following:
 - [“Birnbaum’s Component Importance”](#) on page 315
 - [“Remaining Life BCI”](#) on page 316
 - [“Component Integrated Importance”](#) on page 316
 - [“Mean Time to Failure”](#) on page 317
- Options to show overlay plots for the components in a system diagram. See [“Component Plots”](#) on page 317 and the following:
 - [“Component Distribution Functions”](#) on page 317
 - [“Component Reliability Functions”](#) on page 318
 - [“Component Density Functions”](#) on page 318
 - [“Component Hazard Functions”](#) on page 319
 - [“Component Cumulative Hazard Functions”](#) on page 319
- [“Print Algebraic Reliability Formula”](#) on page 320
- [“Generate Algebraic Expression Data Table”](#) on page 321

- “Clone and Delete” on page 321

JMP PRO Profilers

Various profilers are provided to help you analyze the reliability properties of a system. You can view profilers for each system diagram that is listed in the Designs and the Library panels. This section describes the profilers. The profilers appear in the Profilers panel of the report.

Red Triangle Options for Profilers

Each profiler has a red triangle menu that contains the following commands:

Reset Factor Grid Displays a window where you can enter a current setting and values that control the display. See *Profilers* for more information about setting the factor grid.

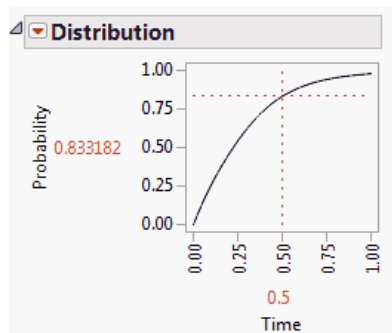
Factor Settings Select this option to configure the profiler’s settings and to link the profilers. See *Profilers* for more information about Factor Settings.

Note: Adjust the X and Y axes of a profiler to view the desired portion of the graph.

JMP PRO Distribution Profiler

The Distribution Profiler displays the probability that the system fails as a function of time. The Distribution profiler for the system appears by default.

Figure 11.18 Distribution

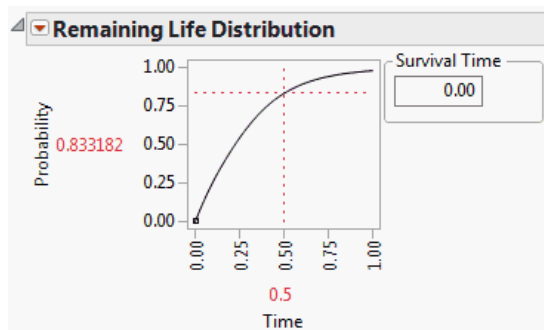


JMP PRO Remaining Life Distribution Profiler

The Remaining Life Distribution profiler for the system shows the probability that the system fails given that it has survived a specified amount of time, designated as Survival Time. By default, Survival Time is set to zero, and the remaining life distribution function is equivalent to the distribution function.

Enter a value for **Survival Time** to indicate the time to which the system has survived without failing. As an alternative to entering a value, select the small rectangle at the origin of the graph and drag it to the right to dynamically set the Survival Time.

Figure 11.19 Remaining Life Distribution



JMP^{PRO} Reliability Profiler

The Reliability profiler for the system shows the probability that the system will function as a function of time. The reliability function is also known as the survival function.

JMP^{PRO} Quantile Profiler

The Quantile profiler for the system shows time as a function of the failure probability. Note that the Quantile function is the inverse of the Distribution function.

JMP^{PRO} Density Profiler

The Density profiler for the system shows the probability density function associated with the system's failure distribution function.

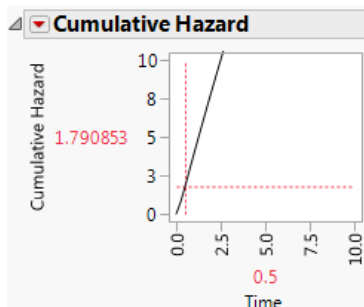
JMP^{PRO} Hazard Profiler

The Hazard profiler for the system shows the instantaneous failure rate at a given time.

JMP^{PRO} Cumulative Hazard Profiler

The Cumulative Hazard profiler for the system shows the cumulative hazard function as a function of time.

Figure 11.20 Cumulative Hazard



JMP PRO Component Importance and Time to Failure

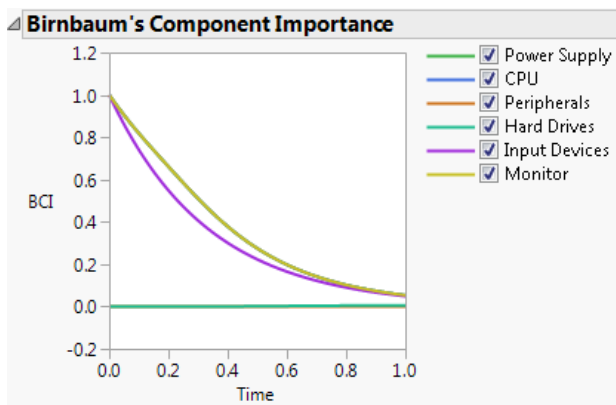
Plots that analyze component importance and mean time to failure are provided for each system diagram that is listed in the Designs and the Library panels. This section describes the component importance measures and mean time to failure. These plots and statistics appear in the profilers panel of the report.

Note: Check boxes in the legend of each component importance plot enable you to select which components to view in that plot.

JMP PRO Birnbaum's Component Importance

Select **Show BCI** to show an overlay plot of the Birnbaum's Component Importance measures for each component of the selected system diagram. A component's BCI at a given time is the probability that the system fails if the component fails. A component with a large BCI is critical to system reliability.

Figure 11.21 Birnbaum's Component Importance

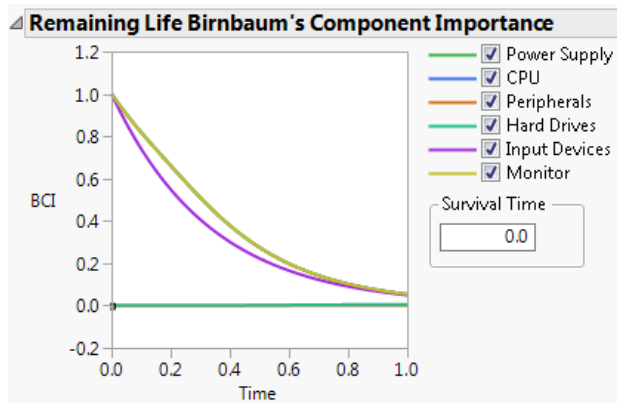


JMP PRO Remaining Life BCI

Select **Show Remaining Life BCI** to show an overlay plot of the Birnbaum's Component Importance for Remaining Life. The BCI for Remaining Life is the probability that the system fails if the component fails, given that the system has survived a specified amount of time, designated as Survival Time. By default, Survival Time is set to zero, and the BCI for Remaining Life is equivalent to the Birnbaum's Component Importance.

Enter a value for **Survival Time** to indicate the time to which the system has survived without failing. As an alternative to entering a value, select the small rectangle at the origin of the graph and drag it to the right to dynamically set the Survival Time.

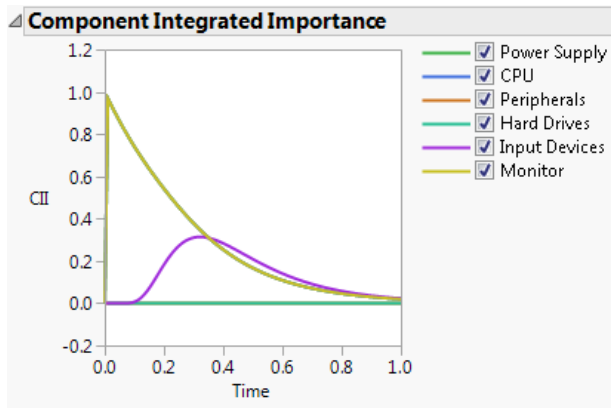
Figure 11.22 Birnbaum's Component Importance for Remaining Life



JMP PRO Component Integrated Importance

Select **Show Component Integrated Importance** to show an overlay plot of the integrated importance measures for the components of the Reliability Block Diagram. The integrated importance measure for each component takes into account the failure rate of the component as well as the likelihood of failing instantaneously. See Si et al. (2012).

Figure 11.23 Component Integrated Importance



JMP PRO Mean Time to Failure

Select **Show MTTF** to view the Mean Time to Failure (MTTF) for the system.

Note: The formula that is used to calculate the Mean Time to Failure depends on the specified failure distributions and Configuration settings for each component in the system.

JMP PRO Component Plots

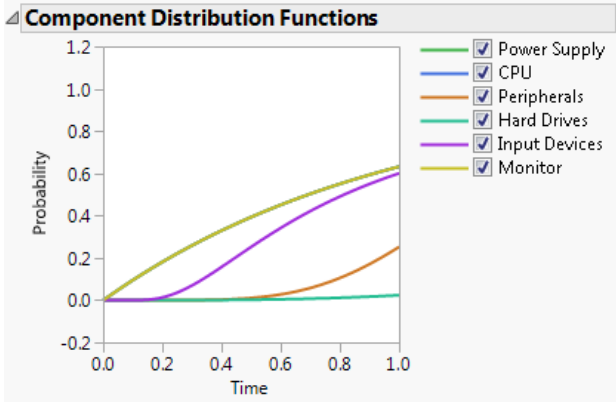
A system diagram usually contains many individual components. The reliability functions of each of these components can be examined in overlay plots. You can view component overlay plots for each system diagram that is listed in the Designs and the Library panels. This section describes the component overlay plots. These plots appear in the profilers panel of the report.

Note: Check boxes in the legend of each component plot enable you to select which components to view in that plot.

JMP PRO Component Distribution Functions

Select **Show Component Distribution Functions** to show an overlay plot of the distribution functions for the components of the Reliability Block Diagram.

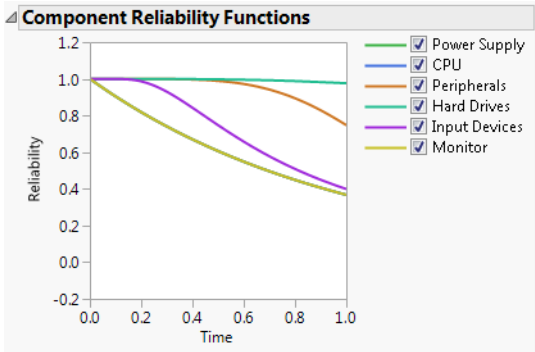
Figure 11.24 Component Distribution Functions



JMP PRO Component Reliability Functions

Select **Show Component Reliability Functions** to show an overlay plot of the reliability functions for the components of the Reliability Block Diagram.

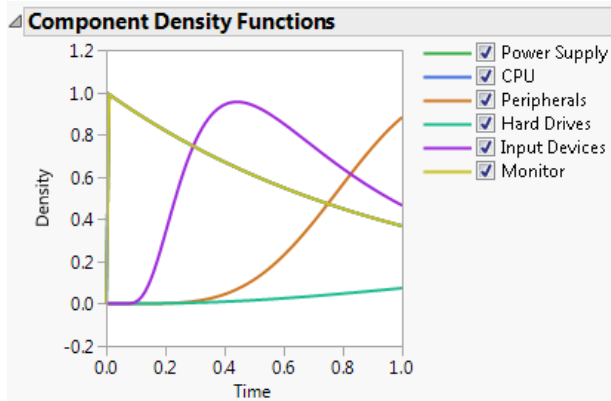
Figure 11.25 Component Reliability Functions



JMP PRO Component Density Functions

Select **Show Component Density Functions** to show an overlay plot of the density functions for the components of the Reliability Block Diagram.

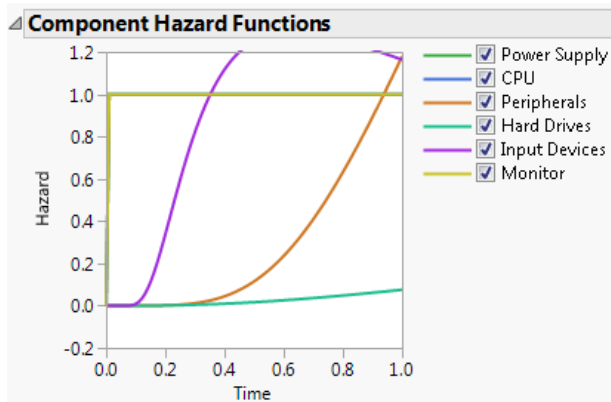
Figure 11.26 Component Density Functions



JMP PRO Component Hazard Functions

Select **Show Component Hazard Functions** to show an overlay plot of the hazard functions for the components of the Reliability Block Diagram.

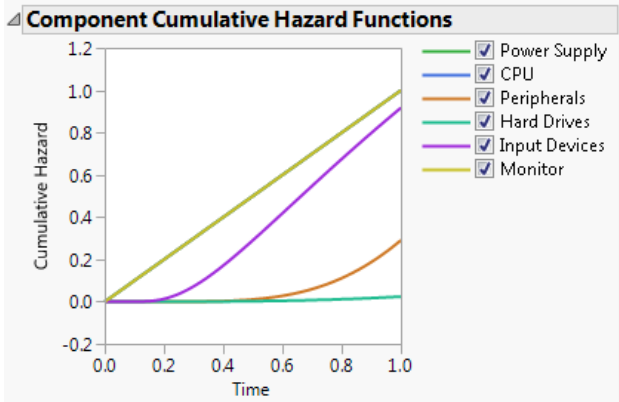
Figure 11.27 Component Hazard Functions



JMP PRO Component Cumulative Hazard Functions

Select **Show Component Cumulative Hazard Functions** to show an overlay plot of the cumulative hazard functions for the components of the Reliability Block Diagram.

Figure 11.28 Component Cumulative Hazard Functions



JMP PRO Print Algebraic Reliability Formula

The Print Algebraic Reliability to the Log Window option prints the reliability formula for the selected system diagram to the Log window.

1. (Windows only) Select **View > Log**.
Or
(macOS only) Select **Window > Log**.
2. Open the exampleRBDcomplete.jsl file that you created.
3. In the Designs panel, select **Computer**.
4. Click the Computer red triangle and select **Print Algebraic Reliability to the Log Window**.

The Log window displays the formula for the Algebraic Reliability of the selected block diagram.

Figure 11.29 Algebraic Reliability in Log Window

/ * :

```
Reliability Block Diagram[]Algebraic Reliability:  
R["Power Supply"] * R["CPU"] * R["Peripherals"] * R["Input Devices"] * R["Monitor"]  
+ R["Power Supply"] * R["CPU"] * F["Peripherals"] * R["Hard Drives"] * R[  
"Input Devices"] * R["Monitor"]
```

Note: Reliability formulas use "R" to represent a component's reliability and "F" to represent probability of a component's failure.

Generate Algebraic Expression Data Table

The Generate Algebraic Expression Data Table option creates a new data table that contains the algebraic reliability expression as a column formula. The table contains a column for each component in the system. The final column contains the system reliability algebraic expression as a JSL column formula that use the preceding columns as arguments.

Clone and Delete

Each system diagram that is listed in the Designs and the Library panels has options that aid in managing the design and library items. This section describes the Clone and Delete operations.

Clone a Design or Library Item

The Clone option creates a new system design or library sub-system identical to the selected system design or library sub-system.

- Select the **Clone** option from the red triangle menu for the design or library item to be copied.
A new design or library item is added to the Designs or Library list.
- Save the file.

Delete a Design or Library Item

The Delete option removes the selected system design or a library sub-system from the panel.

- Select the **Delete** option from the red triangle menu for the design or library item to be deleted.
The specified design or library item is deleted from the Designs or Library list.
- Save the file.

Chapter 12

JMP[®] PRO Repairable Systems Simulation

Estimate Outage Time for a Complex System

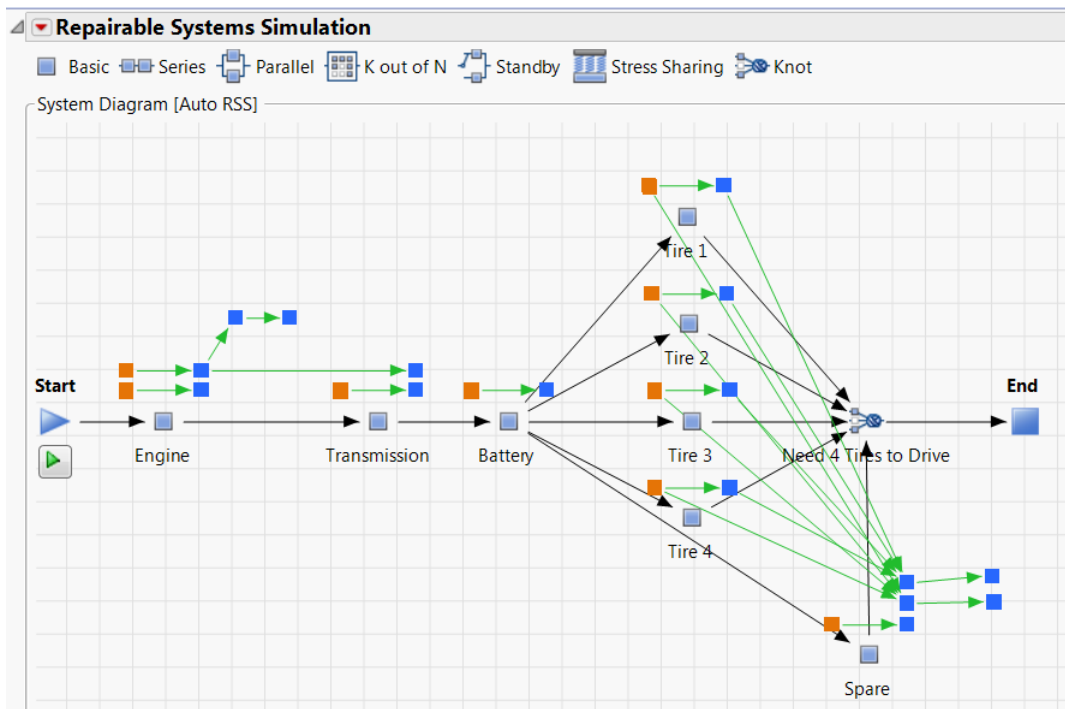
The Repairable Systems Simulation platform is available only in JMP Pro.

The Repairable Systems Simulation (RSS) platform enables you to explore the reliability within complex repairable systems. A *repairable system* consists of individual components that age and require maintenance. The state of a repairable system changes as its components fail and are subsequently repaired.

Repairable systems are involved in many aspects of modern life. They can range in size from refrigerators and houses to power plants and telephone communication networks.

The RSS platform enables you to simulate different system configurations in order to optimize your repairable system. For example, you can extend the useful life of costly components or maximize production outputs while maintaining a safe system operation.

Figure 12.1 Example of a Repairable Systems Simulation Diagram



Contents

Overview of the Repairable Systems Simulation Platform	325
Example Using the Repairable Systems Simulation Platform	325
The Repairable Systems Simulation Window	327
System Diagram	328
Shapes Toolbar	329
Configuration Panel	330
Add Event Panel	333
Add Action Panel	334
Repairable Systems Simulation Platform Options	337
Options for Block Items	338
Distribution Options	338
Specify a Nonparametric or Estimated Distribution	339
Event Settings	340
Action Settings	341
Repairable Systems Simulation Results	342
Results Table	342
Results Explorer	343
Repairable Systems Simulation Results Options	345
RSS Results Explorer Point Estimation Profilers	346

JMP[®] PRO Overview of the Repairable Systems Simulation Platform

The Repairable Systems Simulation (RSS) platform enables you to simulate a repairable system and to analyze its *outage time*. Outage time is the total time a system is not in the On state as a result of either planned maintenance or unintentional failure.

A repairable system is represented by a system diagram. Block shapes in the system diagram represent one or more components that connect to other components. A component that is performing work in the block is said to be *functional*. Components can be connected in series or in parallel. When components are connected in series, the system fails if any of the components in the series fails. When components are connected in parallel, the system fails if all of the parallel components fail.

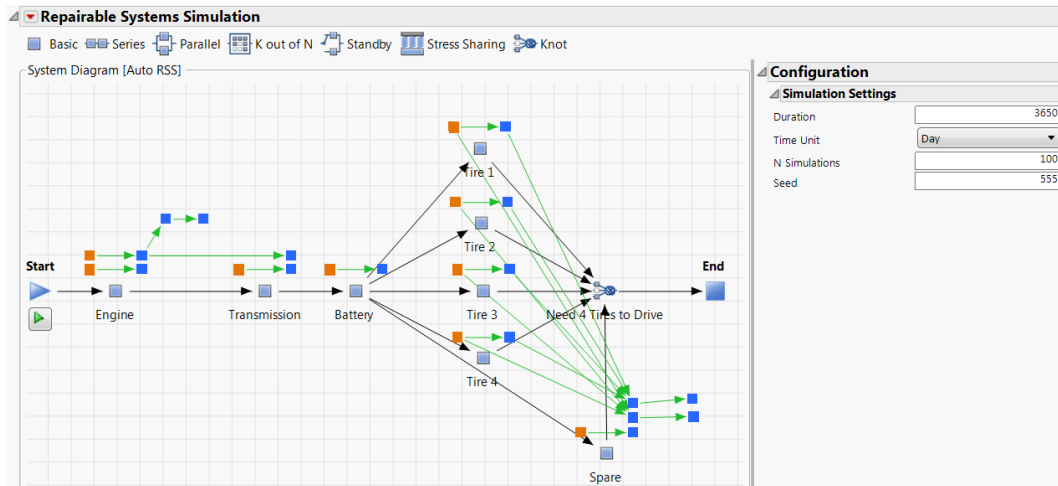
A component pathway is said to be *uninterrupted* when at least one path between the Start and End blocks does not pass through a block shape failure. During a simulation, the system remains in the On state as long as there is an uninterrupted component pathway. If a block shape failure interrupts the component pathway, the system is set to the Down state. If a block shape fails but the component pathway is not interrupted, the system remains in the On state.

You can add unique events and actions to characterize the repairable nature of individual components. An *event* is a specific occurrence within the system, such as component failure, scheduled maintenance, or the start of a simulation iteration. Events trigger the execution of one or more *actions*, which express how components behave. Actions can alter the state of a component or the entire system.

JMP[®] PRO Example Using the Repairable Systems Simulation Platform

In this example, you run 100 iterations of a simulation of a car system and analyze the system's estimated outage time over the course of 10 years.

1. Select **Help > Sample Data**, click **Open the Sample Scripts Directory**, and open Car Repair Simulation.jsl.

Figure 12.2 Car Repair System Diagram

An RSS window that diagrams a car system appears. In the diagram, the first block on the left is the Start block. Notice in the Configuration panel, on the right side of the RSS window, that the simulation is set to run for 3650 days.

2. (Optional) Enter 555 next to Seed.

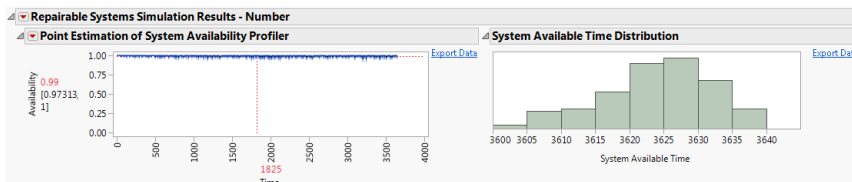
Because the simulation involves random component failures, this action ensures that you obtain the exact results shown below.

3. Click the green triangle below the Start block to simulate the car system.

A data table that contains the simulation results appears.

4. Click the green arrow next to the Launch Repairable Systems Simulation Results Explorer script.

A window appears containing the Repairable Systems Simulation Results report. For more information about how to interpret these results, see [“Repairable Systems Simulation Results”](#) on page 342.

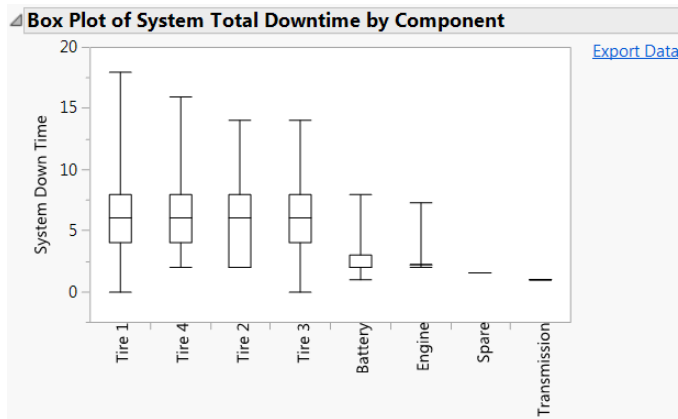
Figure 12.3 Partial RSS Report

You predict that the car will be available between 3,600 and 3,640 days over the next 10 years. Because the values shown in the Point Estimation of System Availability graph are

close to one, you conclude that the car will be mostly available to drive over the next 10 years. You are interested in which components cause the most downtime for the system.

5. Click the red triangle next to Repairable Systems Simulation Results - Number and select **Box Plot of Total System Downtime by Component**.

Figure 12.4 Partial Box Plot of Total System Downtime by Component Report

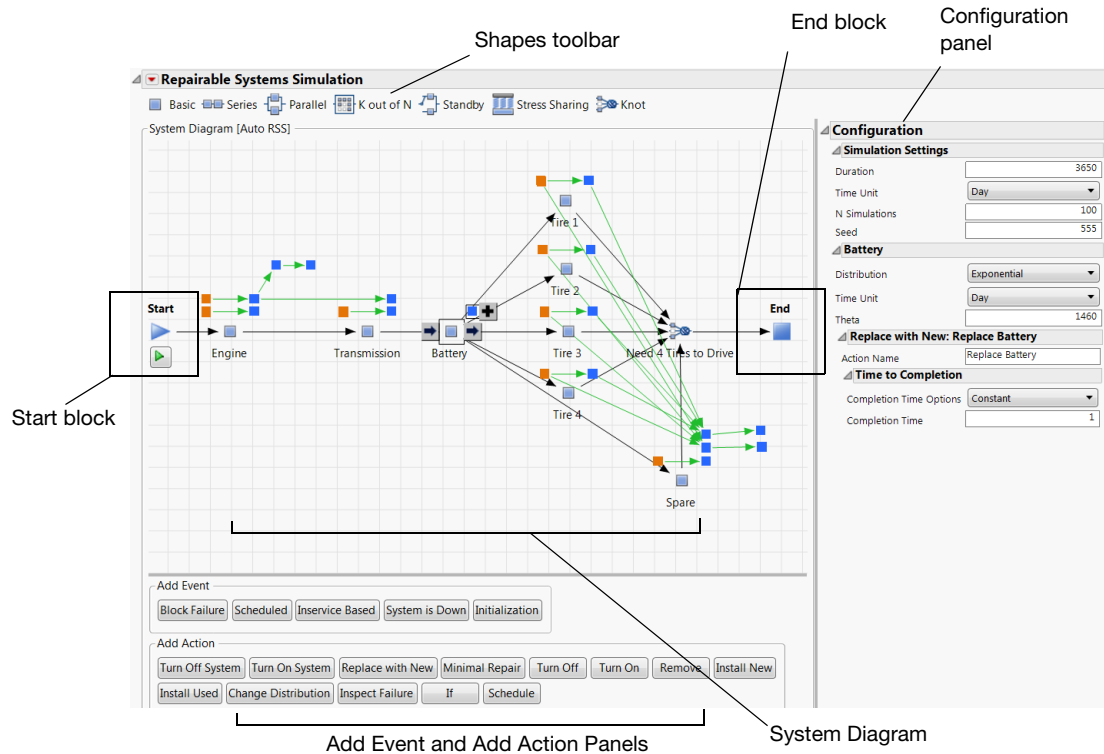


You conclude that the tires cause the most downtime for the car system. To increase the average total system availability, you might consider using more durable tires or carrying another spare tire.

JMP PRO The Repairable Systems Simulation Window

Launch the Repairable Systems Simulation platform by selecting **Analyze > Reliability and Survival > Repairable Systems Simulation**. The Repairable Systems Simulation window is divided into the following panels:

- “System Diagram” on page 328
- “Shapes Toolbar” on page 329
- “Configuration Panel” on page 330
- “Add Event Panel” on page 333
- “Add Action Panel” on page 334

Figure 12.5 The Repairable Systems Simulation Window


JMP PRO System Diagram

The System Diagram is the space where you can diagram a repairable system. A new System Diagram contains a Start block and an End block.

The following options are available in a pop-up menu when you right-click in the System Diagram:

Run Simulation Runs the simulation using the current Simulation Settings in the Configuration panel. The results of the simulation are reproducible if a nonzero value is specified for Seed.

Run Multithreaded Simulation Runs a multithreaded simulation of the system. The results of the simulation are not reproducible when you run a multithreaded simulation.

Diagram Operation Contains options to control the appearance of the diagram:

Delete (Available only if a block, event, or action is selected.) Removes the selected items from the diagram.

Show Block Names Shows or hides the names that appear below the blocks in the diagram.

Zoom In Increases the Zoom Scale value by a factor of 0.9.

Zoom Out Decreases the Zoom Scale by a factor of 0.9.

Tip: On Windows, you can press Ctrl and use the mouse scroll wheel to zoom in to and out from the diagram.

Zoom Scale Enables the user to set the scale of the zoom on a scale of 0 to 10000. The original Zoom Scale value is 1. The Zoom Scale value increases as you zoom into the diagram.

Preview Shows or hides the Preview window in the lower right corner of the System Diagram.

Show Event Action Link Across Blocks Shows or hides the green arrows that represent action links from one block to another. If this option is turned off, action links for selected events and actions still appear.

Align Selected Vertices Vertically (Available only if more than one block is selected.) Updates the horizontal position of the selected blocks so that they are aligned vertically.

Align Selected Vertices Horizontally (Available only if more than one block is selected.) Updates the vertical position of the selected blocks so that they are aligned horizontally.

When you right-click an item in the System Diagram, additional submenu items appear in the pop-up menu. The submenu items that appear depend on the type of item. These options enable you to add, remove, or change the events and actions attached to a block item.

Shapes Toolbar

The Shapes toolbar contains the block shapes used to represent components in the System Diagram. Add block shapes to the System Diagram by clicking the icon for the block shape and dragging it onto the System Diagram. The Shapes toolbar includes the following block shape icons:

 **Basic** Adds a single block shape that represents a single component.

 **Series** Adds a series block shape that represents a group of identical components that are connected in series. All of the components must be functional for the block to remain functional.

 **Parallel** Adds a parallel block shape that represents a group of identical components that are connected in parallel. At least one of the components must be functional for the block to remain functional.

 **K out of N** Adds a k -out-of- n block shape that represents a group of n identical components that are connected in parallel. At least k of the n components must be functional for the block to remain functional.

 **Standby** Adds a standby block shape that represents a group of n identical components that are connected in parallel. Only k components are *active*, or performing work in the system. The remaining *inactive* components act as standby components that are activated when any of the k active components fail. The block must have at least k functional and activated components for the block to remain functional.

 **Stress Sharing** Adds a stress sharing block shape that represents n identical components that are connected in parallel. Components fail one at a time, and remaining components fail at a quicker rate. The block must have at least one functional component and the block must successfully reallocate stress across remaining components for the block to remain functional.

 **Knot** Adds a knot block shape that represents a combination of the block shapes that point to the knot. At least a specified number of the block shapes that point to the knot must be functional for the knot to remain functional.

For information about the settings available for block shapes, see “[Configuration Panel](#)” on page 330.

Configuration Panel

The Simulation Settings report appears at the top of the Configuration panel. The component settings for every selected component appear below the Simulation Settings outline.

Configuration settings for selected events and actions appear below the block shape to which the action or event belongs.

Simulation Settings

Alter the settings for the simulation in the Simulation Settings report, which is available in the Configuration panel. The following settings are available:

Duration The length of time that is simulated in each iteration.

Time Unit The unit of time to use for the simulation.

Note: Recurring events and non-immediate actions use the Time Unit specified in Simulation Settings.

N Simulations The number of iterations in the simulation.

Seed (Optional) A random seed that ensures the reproducibility of simulation results. By default, the Seed is set to zero, which does not produce reproducible results. When you save the analysis to a script, the random seed that you enter is saved to the script.

Caution: If you run a simulation by right-clicking and selecting **Run Multithreaded Simulation**, the results are not reproducible, even if you specify a random seed.

Block Settings

Select a block shape in the System Diagram to see its settings in the Configuration panel.

Each block shape, except the knot block, has a Turn On System Exemption option. If this option is selected for a block, the block is not turned on when the Turn On System action is invoked.

Each block shape, except the knot block, has a failure distribution that determines the rate at which the block shape's individual components randomly fail. The failure distribution for a basic block determines the rate at which the block fails, because basic blocks represent only one component. For more information about the available failure distribution options, see ["Distribution Options"](#) on page 338. Under the Distribution option, you can specify the Time Unit and distribution-dependent parameters for the block.

Series and Parallel

A series block fails when one of its components fails. A parallel block fails when all of its components fail. The following option is available for series and parallel blocks:

N Specifies the number of identical components contained in the block.

K-out-of-N

K-out-of-N blocks contain n identical components. The block fails when fewer than k of the components are functional. The following options are available for K-out-of-N blocks:

K Specifies the minimum number of functional components required for the block to remain functional.

N Specifies the number of identical components contained in the block.

Standby

Standby blocks have secondary components, called standby components, that are inactive. Active components perform work within a standby block. Inactive components do not perform work within a standby block, and are activated one at a time as active components fail. Occasionally, the activation process is not successful. A component switch might fail when activating a standby component. A standby block fails when less than k of its n identical components are active. The following options are available for standby blocks:

K Specifies the number of components that are initially active. This is also the minimum number of active components that is required for the block to remain functional.

N Specifies the total number of identical components within the block. The difference between k and n is equal to the number of standby components.

Switch Type Specifies the mechanism that activates a single standby component if any active component fails.

Single Switch A single switch exists in the block. If the activation of a standby component fails, then the block also fails.

Individual Switches A switch exists for each standby component. If the activation of a standby component fails, that standby component cannot be activated. The standby block attempts to activate the next standby component until a standby component is activated. If no remaining switches are functional and fewer than k of the components are active, then the block fails.

Switch Reliability Specifies the probability of success of activating a standby component when any active component fails.

Standby Type Specifies the state and failure distribution of the standby components.

Cold Standby components do not age until they are activated.

Warm Standby components age according to a secondary failure distribution while they are inactive. When standby components are activated, they age according to the primary failure distribution. Use the secondary failure distribution to mimic reduced stress on standby components that are not performing work in the standby block.

Stress Sharing

A stress sharing block distributes stress equally among its components. As components fail, the components that remain functional experience increased stress and subsequently fail at an increased rate.

N Specifies the total number of identical components contained in the block.

Switch Reliability Specifies the probability of successfully reallocating stress among the remaining functional components. The block fails if the reallocation of stress fails.

Stress Sharing Type Specifies how stress is shared among functional components.

Basic (default) Specifies that stress is shared equally among the remaining functional components. This type of stress sharing is referred to as *Load Sharing*. The characteristic life of individual components is proportional to the number of components that share the work load.

Custom Specifies that components share stress according to the JSL code that appears in the Sharing Formula option. The Sharing Formula defines how stress changes when components fail.

Knot

Knot blocks do not have a failure distribution. Knot blocks fail only if the number of connected blocks shapes that are functional falls below the specified minimum number. The following option is available for knot blocks:

Minimum Available Specifies the minimum number of functional blocks that must point to the knot block for the knot block to remain functional.



Add Event Panel

Events represent discrete occurrences in a simulation. You can use events to trigger actions. There is no limit to how many actions a single event can trigger. The following options are available in the Add Event Panel:

Note: Some block shapes do not have all of the available event types.

Block Failure Occurs when the block fails. If the component pathway is interrupted, the system is set to the Down state.

Scheduled Occurs at a specified recurring interval. The Max Occurrence option sets a limit on the number of occurrences of this event in a simulation. By default, Max Occurrence is set to missing. In this case, the event continues to occur until the end of the simulation. See [“Event Settings”](#) on page 340.

Inservice Based Occurs when the component reaches a specified age. A component's age is the cumulative time that it has been functional and active. Age stops increasing when a component fails or its block is turned off. The Max Occurrence sets a limit on how many times this event can take place in a simulation. By default, the event continues to occur until the end of the simulation. See [“Event Settings”](#) on page 340.

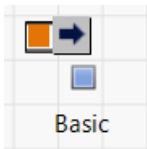
System is Down Occurs when the system is set to either the Down or the Off state. The system state can be changed intentionally or unintentionally.

Initialization Occurs when each simulation iteration begins. Use this event to arrange actions that have to complete before the system runs.

Create an Event

- 1. To create an event, select a block shape in the System Diagram to display the Add Event and Add Action panels at the bottom of the System Diagram.
- 2. Select one of the options in the Add Event panel to define an event for the selected block shape.

Figure 12.6 Block Failure Event



An orange square that represents the new event appears above the component. Notice that a connection arrow appears on the right side of the selected event.

JMP[®] PRO Add Action Panel

An action defines component and system behavior that is triggered by either a connected event or a connected action. Connected actions are triggered upon completion of the prior action. The following options are available in the Add Action Panel:

Note: Some block shapes do not have all of the available action types.

Table 12.1 Action Options

Action Name	Blocks that can use this Action	Starting behavior	Completion Behavior
Turn off System	All		All blocks are turned off. The system is set to the Off state if it was not already in the Down state.

Table 12.1 Action Options *(Continued)*

Action Name	Blocks that can use this Action	Starting behavior	Completion Behavior
Turn on System	All		All blocks that were not already failed are turned on. The system is set to the On state if there is an uninterrupted component pathway. Blocks with the Turn On System Exemption option selected are not turned on by this action.
Replace with New	All	Turns off the original block, and sets the system to the Down state if the component pathway is interrupted.	The block becomes new. Its age is reset and the block is turned on. The system is set to the On state if there is an uninterrupted component pathway.
Minimal Repair	Basic and Series	Equivalent to starting behavior for the Replace with New action.	Turns the block on. The system is set to the On state if there is an uninterrupted component pathway. The age of the block is not reset.
Turn On Block	All	Cancels the action if the block is in the Down state or is currently removed from the system.	Turns the block on increments Turn On Count by one.
Turn Off Block	All		Turns the block off. The system is set to the Down state if the component pathway is interrupted.
Remove Block	All		The block is removed from the system. The system is set to the Down state if the component pathway is interrupted.

Table 12.1 Action Options (*Continued*)

Action Name	Blocks that can use this Action	Starting behavior	Completion Behavior
Install New	All	Equivalent to starting behavior for Replace with New action.	Turns the block off. The age of the block is reset.
Install Used	Basic	Equivalent to starting behavior for Replace with New action.	Turns the block off. You specify a new age and failure distribution for the block.
Change Distribution	Basic	Turns the block off. The system is set to the Down state if the component pathway is interrupted.	Changes the failure distribution of the block. Use this to mimic operation changes over time or cumulative damage to the block.
Inspect Failure	All		Triggers connected actions if the block has failed or is removed from the system.
If	All		Triggers connected actions if the specified condition script is true.
Schedule	All		Triggers connected actions at a specified interval. You can limit the number of scheduled intervals or allow the action to continue through the end of the simulation iteration.

Create an Action

1. To create an action, select a block shape in the System Diagram to display the Add Event and Add Action panels at the bottom of the System Diagram.
2. Select one of the options from the Add Action panel to define an action for the selected block shape.

A blue action square is created above the selected block shape. Actions are triggered when a connected event occurs. Create an event that triggers the action that you defined.

3. Select one of the options from the Add Event panel to define an event for the selected block shape.

An orange event square is created above the block shape. Notice the connection arrow on the right side of the event.

4. Click the connection arrow and drag it to the blue action square that you created in step 2.

A green arrow connects the event and action squares. When the event occurs in a simulation iteration, it triggers the connected action.

5. Select the blue action square that you created in step 2.

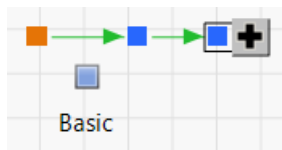
You can connect additional actions that are triggered upon completion of the previous action by using the addition sign that appears on the right side of the selected action.

6. Click the addition sign and drag it to an empty area in the System Diagram.

A list of the available actions appears.

7. Select one of the options from the action list to create an action that is connected to the first action.

Figure 12.7 Create an Action



A green arrow connects the two actions. The second action is triggered upon completion of the first action.



Repairable Systems Simulation Platform Options

The Repairable Systems Simulation red triangle menu contains the following options:

Save and Save As Enables you to save a Repairable Systems Simulation to a JMP Scripting Language (JSL) script that is automatically executed when it is opened in JMP. See the *Scripting Guide* for more information about Auto-Submit scripts.

Note: The Save and Save As red triangle options are equivalent to selecting **File > Save** and **File > Save As**, respectively. They are available in the red triangle menu for convenience.

Import Component Distribution Settings Enables you to import configuration settings for the system diagram from a data table. The table must contain columns for the component name, distribution, and one or more parameters. The number of parameters depends on the specified distribution.

Note: The strings in the imported table must be exact matches to the strings in the system diagram.

JMP PRO Options for Block Items

- To view settings for the components in a selected design or subsystem, do the following:
- To view the Configuration settings for a block shape, select the block shape in the System Diagram.
 - To view Configuration settings for more than one block shape, select multiple block shapes using the Arrow tool, press Ctrl, and click multiple block shapes.

JMP PRO Distribution Options

Each block shape randomly fails according to its specified failure distribution. In addition, you specify a time unit for the distribution. The block shape’s time unit can be different from the time unit option that you specify in the simulation settings. The Turn On Count option is based on the number of times the individual block shape has been set to the On state.

The available failure distributions are listed in Table 12.2.

Table 12.2 Distributions and Additional Parameters

Property Type	Required Inputs
Exponential	Theta
Weibull	Alpha, Beta
Lognormal	location, scale
Loglogistic	location, scale
Fréchet	location, scale
GenGamma	mu, sigma, lambda
DS Weibull	Alpha, Beta, Defective Probability

Table 12.2 Distributions and Additional Parameters *(Continued)*

Property Type	Required Inputs
DS Lognormal	location, scale, Defective Probability
DS Loglogistic	location, scale, Defective Probability
DS Fréchet	location, scale, Defective Probability
Nonparametric	data or data file
Estimated	estimated distribution, data, or data file

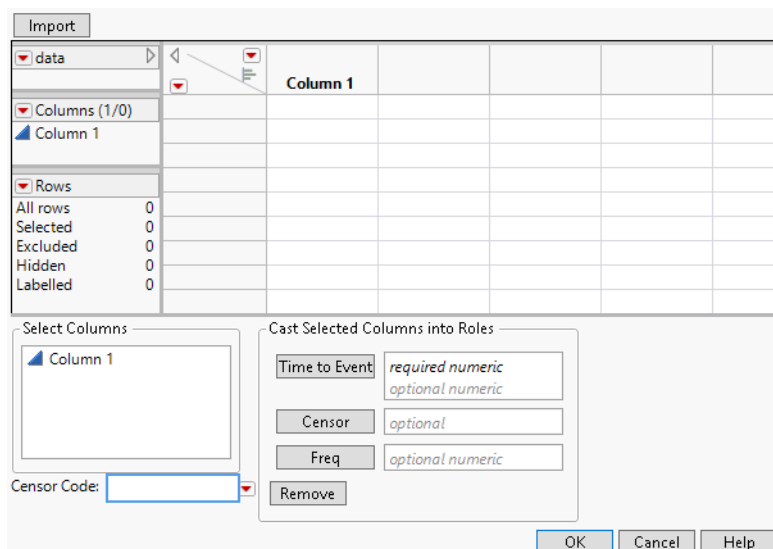
To see the formulas and parameterization for these failure distributions, see [“Distributions”](#) on page 82 in the “Life Distribution” chapter.

JMP PRO Specify a Nonparametric or Estimated Distribution

The Nonparametric and Estimated options under Distribution enable you to approximate an arbitrary distribution. You can enter data manually or import a file that contains data. These data are used to approximate the distribution.

After selecting Nonparametric or Estimated, click the  icon next to Data. The Provide Data window appears, which enables you to either enter data or import a data file. After you have imported or entered your data, the data are used to calculate a distribution for the component.

Figure 12.8 Provide Data Window



The screenshot shows the 'Provide Data' window in JMP PRO. The window has a title bar 'Import'. Below the title bar is a table with columns and rows. The 'Columns' section shows 'Column 1' selected. The 'Rows' section shows 'All rows' selected. The 'Select Columns' section shows 'Column 1' selected. The 'Cast Selected Columns into Roles' section shows 'Time to Event' as 'required numeric', 'Censor' as 'optional', and 'Freq' as 'optional numeric'. There is a 'Censor Code' field and a 'Remove' button. At the bottom are 'OK', 'Cancel', and 'Help' buttons.

To import data from a file, do the following:

1. Before clicking the  icon, open the JMP data table that contains the data to import.
2. Click the  icon next to Data.
3. In the Provide Data window, click **Import**.
The Select Data Table window appears.
4. From the Data Table list, select a data table.
5. Click **OK**.
6. In the panel beneath the data grid, specify the columns that represent Time to Event data.
7. Specify Censor and Freq columns, if appropriate.
8. In the Provide Data window, click **OK**.

To enter data manually, do the following:

1. Create columns for Time to Event data.
2. Create columns for Censor and Freq data, if appropriate.
3. Enter the data in the columns.
4. In the panel beneath the data grid, specify which columns represent Time to Event data, as well as Censor and Freq data if appropriate. See step 5 and step 6 above.
5. In the Provide Data window, click **OK**.

Event Settings

Select an event in the System Diagram to see its settings in the Configuration panel. You can change the Event Name setting to distinguish the event from others. Only the Scheduled and Inservice events have additional settings, which are listed here:

Scheduled

Scheduled events have the following options:

Recurring Interval Specifies the length of the time interval between event occurrences.

Max Occurrences Specifies the maximum number of times this event can occur.

Inservice

Inservice events have the following options:

Inservice Specifies the interval of component run time to wait before this event occurs.

Note: Component run time accumulates only when the system is in the On state and the component is functional.

Max Occurrences Specifies the maximum number of times this event can occur.

JMP PRO Action Settings

Select an action in the System Diagram to see its settings in the Configuration panel. You can change the Action Name setting to distinguish the action from others. By default, actions are completed immediately. Non-immediate completion time options are available in the Completion Time Options menu.

Figure 12.9 Action Settings

The screenshot shows the 'Replace with New: Replace Engine' configuration panel. It has two main sections: 'Action Name' with a text field containing 'Replace Engine', and 'Time to Completion'. The 'Time to Completion' section includes a 'Completion Time Options' dropdown menu set to 'Constant' and a 'Completion Time' text field containing the value '0'.

The following options appear in the Completion Time Options list:

Immediate (default) Specifies that no time passes between starting and completion behavior.

Constant Specifies that the time lapse is always the specified Completion Time.

Choice Specifies that the time lapse is randomly chosen from the specified list of comma-separated values.

Uniform Specifies that the time lapse is randomly chosen from a uniform distribution with the specified Minimum and Maximum.

Triangle Specifies that the time lapse is randomly chosen from a triangular distribution with the specified Minimum, Mode, and Maximum.

Normal Specifies that the time lapse is randomly chosen from a normal distribution with the specified Mean and Standard Deviation.



Repairable Systems Simulation Results

- “Results Table”
- “Results Explorer”
- “Repairable Systems Simulation Results Options”
- “RSS Results Explorer Point Estimation Profilers”



Results Table

When you click the green arrow under the Start block to run the simulation, the results of the simulation appear in a data table. This table describes the events and subsequent actions that occurred in each simulation iteration.

The results table contains the following columns:

Sim ID Identifies the simulation iteration to which the event or action belongs.

Time Gives the exact time that the event or action took place in the simulation.

Subject Gives the name of the block shape to which the event or action is connected or System in the case of system events and actions.

Predicate Gives the name of the event or action that took place.

State Gives the state of the Subject at that exact Time. The initialization and termination of each action are denoted with Start and Finish, respectively.

Note Gives an additional description of an action. The end of each simulation iteration is denoted with End.

Figure 12.10 RSS Results Table

	Sim ID	Time	Subject	Predicate	State	Note
1	1	0	Spare	Initialization Spar...	Start	
2	1	0	Spare	Initialization Spar...	Off	
3	1	0	Spare	Initialization Spar...	Finish	
4	1	65.024584796	Tire 2	Tire 2 is unrepair...	Down	
5	1	65.024584796	System	Turn Down System	Down	Unintended
6	1	65.024584796	Engine	Turn Off Block by...	Off	
7	1	65.024584796	Transmission	Turn Off Block by...	Off	
8	1	65.024584796	Battery	Turn Off Block by...	Off	
9	1	65.024584796	Tire 1	Turn Off Block by...	Off	
10	1	65.024584796	Tire 3	Turn Off Block by...	Off	
11	1	65.024584796	Tire 4	Turn Off Block by...	Off	
12	1	65.024584796	Tire 2	Replace Tire 2	Removed	
13	1	65.024584796	Tire 2	Replace Tire 2	Start	

Notice that the first entry in Figure 12.10 is the start of the Initialization Spare in Trunk action, which has an immediate completion time. The action turns the Spare component off and then finishes while Time is still 0.

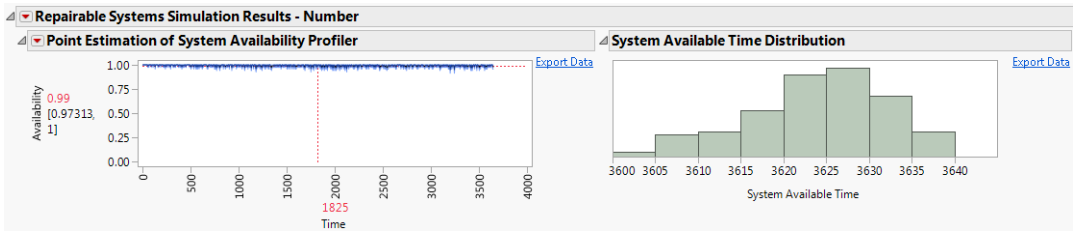
The next event that occurs is Tire 2 is Unrepairable when approximately 65 days have passed in the first simulation iteration. In row five, the system is unintentionally set to the Down state because only three tires are functional. The Need 4 Tires to Drive knot requires at least four tires, including the Spare, to be functional. The Tire 2 is Unrepairable event simultaneously triggers the Replace Tire 2 action and the Use Spare action. The Drive with Spare action is triggered in row eleven, and the system is set to the On state.

Notice that Tire 2 failed and was replaced by the Spare at the same time. No time elapsed between the time that the system was set to the Down state and the time that the system was returned to the On state. Because the subsequent actions had an immediate completion time, the Tire 2 is Unrepairable event did not cause any system outage time. The results explorer shows the system outage time that accumulates when actions are non-immediate.

JMP[®] PRO Results Explorer

When you click the green arrow next to the Launch Repairable Systems Simulation Results Explorer script in the results data table, a report appears that contains an analysis of the simulation results.

Figure 12.11 Partial RSS Explorer Report



By default, the report contains two types of graphs. The first type of graph, shown on the left side of Figure 12.11, is a point estimate graph. The point estimate graph displays aggregated simulation results on the vertical axis plotted against simulation iteration time on the horizontal axis. The range of the horizontal axis is the duration that you specify in the simulation settings. The estimated probability of the system being available at a given time is shown in red next to the vertical axis. Below the point estimate is a 95% confidence interval for the point estimate.

Tip: To see the point estimation of system availability at a specific time in the simulation iteration, click the red number below the horizontal axis. Specify a time within the range of the simulation duration and press Enter.

The second type of graph, shown on the right side of Figure 12.11, is a histogram of the system available time from the simulation iterations. The bins in the histograms are representative of one of the following:

- Total simulation time.
- Proportion of the simulation duration.

Tip: Select the **Export Data** option beside a graph to create a data table with the data that the graph uses.

By default, the report contains the following graphs:

- Point Estimation of System Availability Profiler** Shows a profiler graph over time of the estimated probability that the system is in the On or the Off state during a simulation iteration.
- System Available Time Distribution** Shows a histogram of the total time that the system was available for each simulation iteration.
- System Availability Distribution** Shows a histogram of the total time that the system was available for each simulation iteration divided by the duration of the simulation. This distribution is the same as the System Available Time Distribution, except that it is shown as a proportion of the duration of the simulation.

Point Estimation of System In Service Probability Profiler Shows a profiler graph over time of the estimated probability that the system is in the On state during a simulation iteration.

System In Service Time Distribution Shows a histogram of the total time that the system was in the On state for each simulation iteration.

System In Service Probability Distribution Shows a histogram of the total time that the system was in the On state for each simulation iteration divided by the duration of the simulation. This distribution is the same as System In Service Time Distribution, except that the distribution is shown as a proportion of the duration of the simulation.

Point Estimation of System Unplanned Outage Profiler Shows a profiler graph over time of the estimated probability that the system is in the Down state during a simulation iteration.

System Unplanned Outage Distribution Shows a histogram of the total time that the system was in the Down state for each simulation iteration.

System Unplanned Outage Percentage Distribution Shows a histogram of the total time that the system was in the Down state for each simulation iteration divided by the duration of the simulation. This distribution is the same as System Unplanned Outage Distribution, except that the distribution is shown as a proportion of the duration of the simulation.

Repairable Systems Simulation Results Options

The Repairable Systems Simulation Results red triangle menu contains the following options:

Point Estimation of Component Availability For each block shape in the system diagram, this option shows the following reports:

- Point Estimation of Availability Profiler
- Available Time Distribution
- Availability Distribution
- Point Estimation of Unplanned Outage Profiler
- Unplanned Outage Time Distribution
- Unplanned Outage Distribution

System Availability by Period Shows a graph that enables you to examine the average system availability during specified time periods. Use the By Time panel to define the time period as an interval of time. Use the By Event panel to define the time period as a specified number of occurrences of one event. When you click **Go**, the report shows a bar chart of the average system availability during each period.

Box Plot of System Total Downtime by Component Shows a box plot of the total time the system was in the Down or Off states caused by individual components. The components are ordered by contribution from the most outage time to the least outage time, similar to a Pareto chart.

Box Plot of System Unplanned Total Outage Time by Component Shows a box plot of total time the system was in the Down state caused by individual components. The components are ordered by contribution from the most outage time to the least outage time, similar to a Pareto chart.

See *Using JMP* for more information about the following options:

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.



RSS Results Explorer Point Estimation Profilers

Each point estimation profiler within the Repairable Systems Simulation Results Explorer report has a red triangle menu that contains the following options:

Confidence Intervals Shows or hides the 95% confidence intervals on the point estimation graphs.

Reset Factor Grid Opens the Factor Settings window, which enables you to modify the parameters of the point estimation graph. See *Profilers* for more information about setting the factor grid.

Factor Settings Provides additional options affecting the Factor Grid. See *Profilers* for more information about Factor Settings.

Chapter 13

Survival Analysis

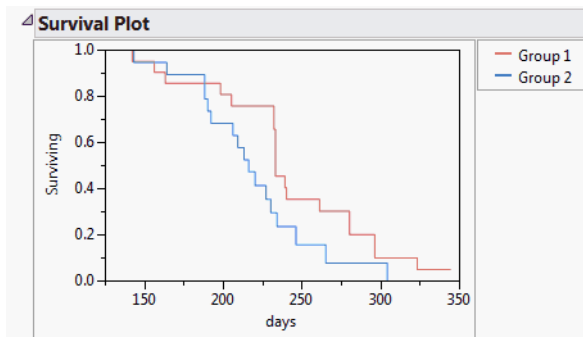
Analyze Survival Time Data

Survival data contain duration times until the occurrence of a specific event and are sometimes referred to as *event-time* response data. The event is usually failure, such as the failure of an engine or death of a patient. If the event does not occur before the end of a study for an observation, the observation is said to be *censored*.

The Survival platform fits a single Y that represents time to event (or time to failure). Use the Survival platform to examine the distribution of the failure times.

Tip: To fit explanatory models, use the Fit Parametric Survival platform or the Fit Proportional Hazards platform.

Figure 13.1 Example of a Survival Plot



Contents

Overview of the Survival Analysis Platform.....	349
Example of Survival Analysis	350
Launch the Survival Platform	352
The Survival Plot	353
Survival Platform Options	354
Exponential, Weibull, and Lognormal Plots and Fits	356
Fitted Distribution Plots.....	360
Competing Causes	361
Additional Examples of the Survival Platform.....	362
Example of Survival Analysis	362
Example of Competing Causes	364
Example of Interval Censoring	367
Statistical Reports for Survival Analysis	368

Overview of the Survival Analysis Platform

Survival data need to be analyzed with specialized methods for two reasons:

1. The survival times usually have specialized non-normal distributions, like the exponential, Weibull, and lognormal.
2. Some of the data could be censored.

Survival functions are calculated using the nonparametric Kaplan-Meier method for one or more groups of either *complete* or *right-censored* data. Complete data have no censored values. Right-censoring is when you do not know the exact survival time, but you know that it is greater than the specified value. Right-censoring occurs when the study ends without all the units failing, or when a patient has to leave the study before it is finished. The censored observations cannot be ignored without biasing the analysis. The elements of a survival model are:

- A time indicating how long until the unit (or patient) either experienced the event or was censored. Time is the model response (Y).
- A censoring indicator that denotes whether an observation experienced the event or was censored. JMP uses the convention that the code for a censored unit is 1 and the code for a non-censored event is zero.
- Explanatory variables (if a regression model is used.)
- Interval censoring is when a data point is somewhere on an interval between two values. If interval censoring is needed, then two Y variables hold the lower and upper limits bounding the event time.

Common terms used for reliability and survival data include lifetime, life, survival, failure-time, time-to-event, and duration.

The Survival platform computes product-limit (Kaplan-Meier) survival estimates for one or more groups. It can be used as a complete analysis or is useful as an exploratory analysis to gain information for more complex model fitting. The Kaplan-Meier Survival platform does the following:

- Shows a plot of the estimated survival function for each group. A plot for the whole sample is optional.
- Calculates and lists survival function estimates for each group and for the combined sample.
- Shows exponential, Weibull, and lognormal diagnostic failure plots to graphically check the appropriateness of using these distributions for further regression modeling. Parameter estimates are available on request.
- Computes the Log Rank and generalized Wilcoxon Chi-square statistics to test homogeneity of the estimated survival function across groups.

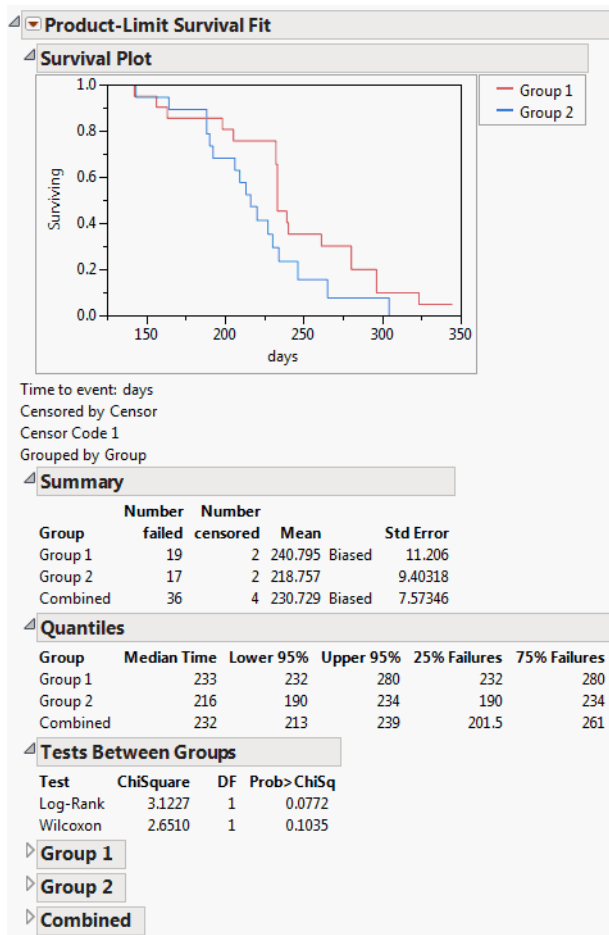
- Analyzes competing causes, prompting for a cause of failure variable, and estimating a Weibull failure time distribution for censoring patterns corresponding to each cause.

Example of Survival Analysis

An experiment was undertaken to characterize the survival time of rats exposed to a carcinogen in two treatment groups. The data are in the `Rats.jmp` sample data table. The event in this example is death. The objective is to see whether rats in one treatment group live longer (more days) than rats in the other treatment group.

1. Select **Help > Sample Data Library** and open `Rats.jmp`.
The data in the `days` column is the survival time. Notice that some observations are censored.
2. Select **Analyze > Reliability and Survival > Survival**.
3. Select `days` and click **Y, Time to Event**.
4. Select `Group` and click **Grouping**.
5. Select `Censor` and click **Censor**.
6. Click **OK**.

Figure 13.2 Survival Plot for Rats.jmp Data

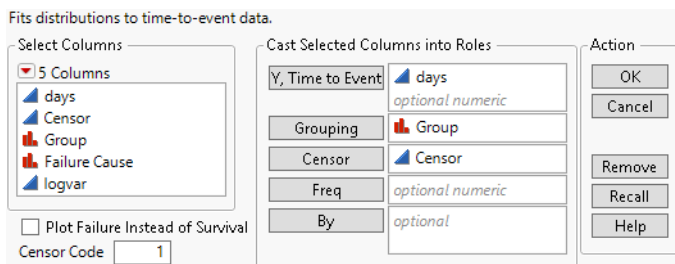


It appears that the rats in treatment group 1 are living longer than the rats in treatment group 2.

Launch the Survival Platform

Launch the Survival platform by selecting **Analyze > Reliability and Survival > Survival**.

Figure 13.3 The Survival Launch Window



For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The Survival launch window contains the following options:

Y, Time to Event Identifies the time to event or time to censoring. If you have interval censoring, specify two Y variables, representing the lower and upper limits.

Grouping Classifies the data into groups that are fit separately.

Censor Identifies censored values. Enter the value that identifies censoring in the Censor Code box. This column can contain more than two distinct values under the following conditions:

- All censored rows contain the value that is entered in the Censor Code box.
- Non-censored rows have a value other than what is in the Censor Code box.

Freq Indicates the column whose values are the frequencies of observations for each row when there are multiple units recorded. If the value is 0 or a positive integer, then the value represents the frequencies or counts of observations for each row.

By Performs a separate analysis for each level of a classification or grouping variable.

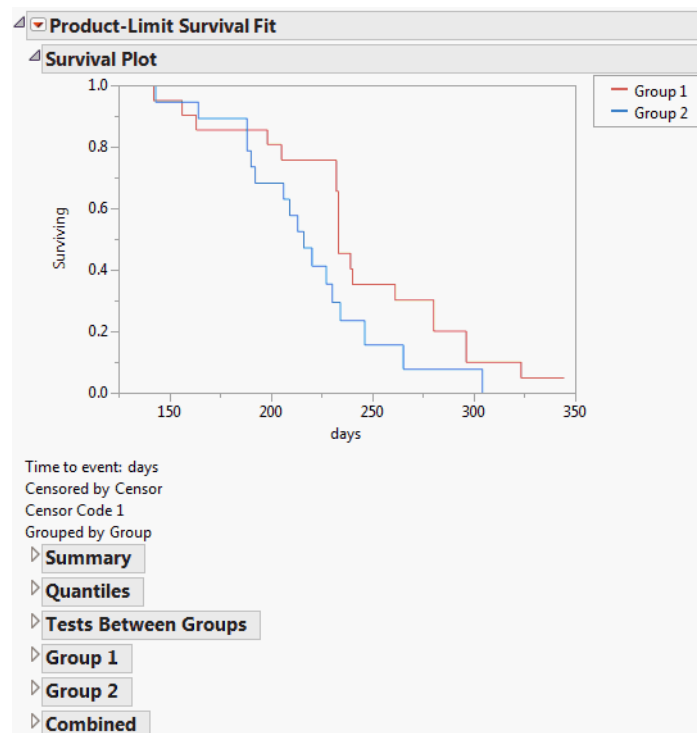
Plot Failure Instead of Survival Shows a failure probability plot instead of its reverse (a survival probability plot).

Censor Code Identifies the value in the Censor column that designates right-censored observations. The default value is 1.

The Survival Plot

The Survival platform shows overlay step plots of estimated survival functions for each group. A legend identifies groups by color and line type.

Figure 13.4 The Survival Plot



Reports beneath the plot show summary statistics and quantiles for survival times. Estimated survival times for each observation are computed within groups. Survival times are computed from the combined sample. When there is more than one group, statistical tests compare the survival curves.

If there are any failures that occur at time zero, the Failures at Time Zero report appears when you request a distribution fit from the Product-Limit Survival Fit red triangle menu. This report contains a table of counts of zero-time failures for each level of the Grouping variable. If there is no Grouping variable specified, there is one row in the table labeled Combined. The table also contains a column labeled Prob Time>0. This column is the proportion of observations in each group that has a nonzero failure time.

Survival Platform Options

The Product-Limit Survival Fit red triangle menu contains the following options:

Survival Plot Shows the overlaid survival plots for each group.

Failure Plot Shows the overlaid failure plots (proportion failing over time) for each group (in the tradition of the reliability literature.) A failure plot reverses the vertical axis to show the number of failures rather than the number of survivors.

Note: The Failure Plot option replaces the Reverse Y Axis option found in older versions of JMP (which is still available in scripts).

Plot Options Contains the following options:

Note: The first seven options (Show Points, Show Kaplan Meier, Show Combined, Show Confid Interval, Show Simultaneous CI, Show Shaded Pointwise CI, and Show Shaded Simultaneous CI) and the last two options (Fitted Survival CI, Fitted Failure CI) pertain to the initial survival plot and failure plot. The other five (Midstep Quantile Points, Connect Quantile Points, Fitted Quantile, Fitted Quantile CI Lines, Fitted Quantile CI Shaded) pertain only to the distributional plots.

Show Points Shows the sample points at each step of the survival plot. Failures appear at the bottom of the steps, and censorings are indicated by points above the steps.

Show Kaplan Meier Shows the Kaplan-Meier curves. This option is on by default.

Show Combined Shows the survival curve for the combined groups in the Survival Plot.

Show Confid Interval Shows the pointwise 95% confidence bands on the survival plot for groups and for the combined plot when it appears with the Show Combined option.

Show Points, Show Combined When you select the Show Points and Show Combined options, the survival plot for the total or combined sample appears as a gray line. The points also appear at the plot steps of each group.

Show Simultaneous CI Shows the simultaneous confidence bands for all groups on the plot. Meeker and Escobar (1998, ch. 3) discuss pointwise and simultaneous confidence intervals and the motivation for simultaneous confidence intervals in survival analysis.

Midstep Quantile Points Changes the plotting positions to use the *modified Kaplan-Meier* plotting positions. These plotting positions are equivalent to taking mid-step positions of the Kaplan-Meier curve, rather than the bottom-of-step positions. This option is recommended, so it is on by default.

Connect Quantile Points Shows the lines in the plot. This option is on by default.

Fitted Quantile Shows the straight-line fit on the fitted Weibull, lognormal, or exponential plot. This option is on by default.

Fitted Quantile CI Lines Shows the 95% confidence bands for the fitted Weibull, lognormal, or exponential plot.

Fitted Quantile CI Shaded Shows the display of the 95% confidence bands for a fit as a shaded area or dashed lines.

Fitted Survival CI Shows the confidence intervals (on the survival plot) of the fitted distribution.

Fitted Failure CI Shows the confidence intervals (on the failure plot) of the fitted distribution.

Exponential Plot Plots the cumulative exponential failure probability by time for each group. Lines that are approximately linear empirically indicate the appropriateness of using an exponential model for further analysis. For example, in Figure 13.5, the lines for Group 1 and Group 2 in the Exponential Plot are curved rather than straight. This indicates that the exponential distribution is not appropriate for this data. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

Exponential Fit Produces the Exponential Parameters table and the linear fit to the exponential cumulative distribution function in the Exponential Plot (Figure 13.5). The parameter Theta corresponds to the mean failure time. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

Weibull Plot Plots the cumulative Weibull failure probability by log(time) for each group. A Weibull plot that has approximately parallel and straight lines indicates a Weibull survival distribution model might be appropriate to use for further analysis. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

Weibull Fit Produces the linear fit to the Weibull cumulative distribution function in the Weibull plot and two popular forms of Weibull estimates. These estimates are shown in the Extreme-value Parameter Estimates table and the Weibull Parameter Estimates tables (Figure 13.5). The Alpha parameter is the 0.632 quantile of the failure-time distribution. The Extreme-value table shows a different parameterization of the same fit, where $\Lambda = \ln(\text{Alpha})$ and $\Delta = 1/\text{Beta}$. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

LogNormal Plot Plots the cumulative lognormal failure probability by log(time) for each group. A lognormal plot that has approximately parallel and straight lines indicates a lognormal distribution is appropriate to use for further analysis. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

LogNormal Fit Produces the linear fit to the lognormal cumulative distribution function in the lognormal plot and the LogNormal Parameter Estimates table shown in Figure 13.5. Mu and Sigma correspond to the mean and standard deviation of a normally distributed natural logarithm of the time variable. See [“Exponential, Weibull, and Lognormal Plots and Fits”](#) on page 356.

Fitted Distribution Plots Use in conjunction with the fit options to show three plots corresponding to the fitted distributions: Survival, Density, and Hazard. If you have not performed a fit, no plot appears. See [“Fitted Distribution Plots”](#) on page 360.

Competing Causes Performs an estimation of the Weibull model using the specified causes to indicate a failure event and other causes to indicate censoring. The fitted distribution appears as a dashed line in the Survival Plot. See [“Competing Causes”](#) on page 361.

Estimate Survival Probability Estimates survival probabilities for the time values that you specify.

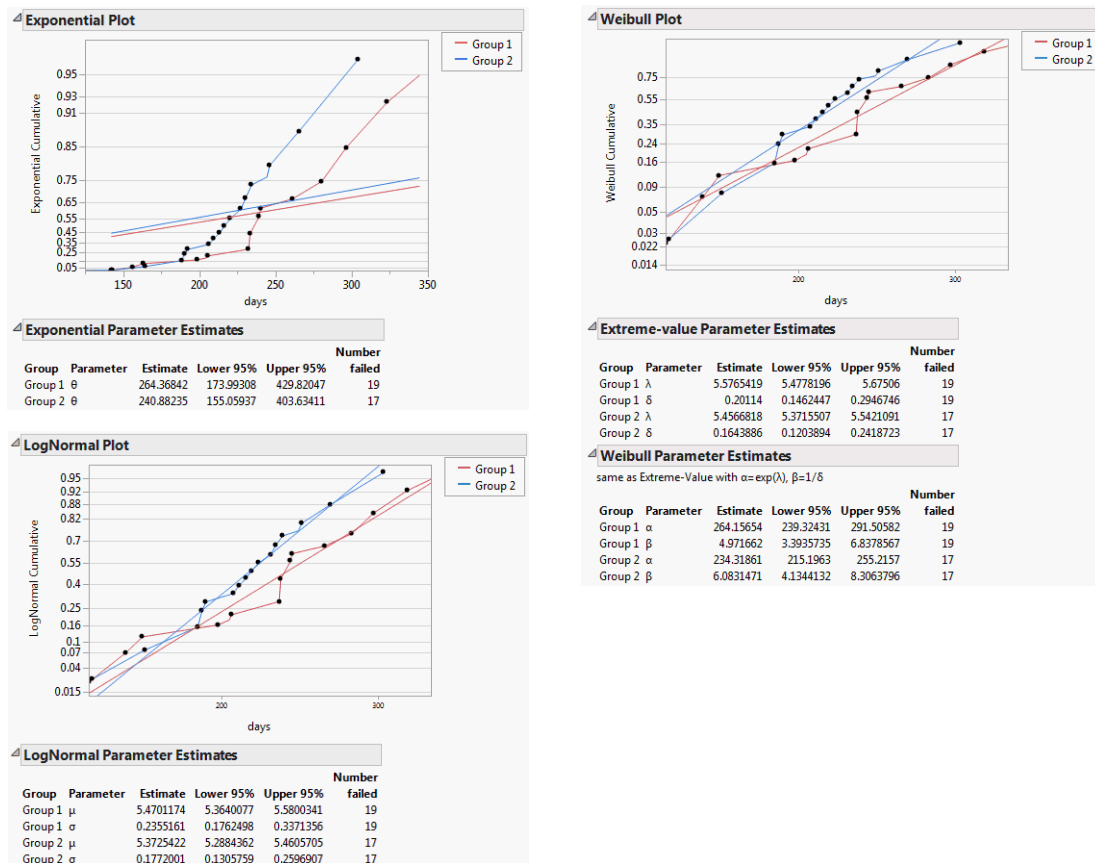
Estimate Time Quantile Estimates a time quantile for each survival probability that you specify.

Save Estimates Creates a data table containing survival and failure estimates, along with confidence intervals, and other distribution statistics.

Exponential, Weibull, and Lognormal Plots and Fits

For each of the three supported distributions in the Survival platform, there is a plot command and a fit command. Use the plot command to see whether the event markers seem to follow a straight line. The markers tend to follow a straight line when the distributional fit is suitable for the data. Then, use the fit commands to estimate the parameters.

Figure 13.5 Exponential, Weibull, and Lognormal Plots and Reports



The following table shows what to plot to make a straight line fit for that distribution:

Table 13.1 Straight Line Fits for Distribution

Distribution Plot	Horizontal Axis	Vertical Axis	Interpretation
Exponential	time	$-\log(S)$	slope is $1/\theta$
Weibull	$\log(\text{time})$	$\log(-\log(S))$	slope is β
Lognormal	$\log(\text{time})$	$\text{Probit}(1-S)$	slope is $1/\sigma$

Note: S = product-limit estimate of the survival distribution.

Exponential

The exponential distribution is the simplest distribution for modeling time-to-event data. The exponential distribution has only one parameter, θ . It is a constant-hazard distribution, with no memory of how long it has survived to affect how likely an event is. The parameter θ is the expected lifetime.

Weibull

The Weibull distribution is the most popular distribution for modeling time-to-event data. The Weibull distribution can have two or three parameters. The Survival platform fits the two-parameter Weibull distribution. Authors parameterize this distribution in many different ways (Table 13.2). JMP reports two of these parameterizations: the Weibull alpha-beta parameterization and a parameterization based on the smallest extreme value distribution.

The alpha-beta parameterization, shown in the Weibull Parameter Estimates report, is widely used in the reliability literature (Nelson 1990). The α parameter is interpreted as the quantile at which 63.2% of the units fail. The β parameter determines how the hazard rate changes over time. If $\beta > 1$, the hazard rate increases over time; if $\beta < 1$, the hazard rate decreases over time; and if $\beta = 1$, the hazard rate is constant over time. A Weibull distribution with a constant hazard function is equivalent to an exponential distribution.

The lambda-delta extreme value parameterization is shown in the Extreme-Value Parameter Estimates report. This parameterization is sometimes desirable in a statistical sense because it places the Weibull distribution in a location-scale setting (Meeker and Escobar 1998, p. 86). The location parameter is λ , and the scale parameter is δ . In relation to the alpha-beta parameterization, λ is equal to the natural log of α , and δ is equal to the reciprocal of β . Therefore, the δ parameter determines how the hazard rate changes over time. If $\delta > 1$, the hazard rate decreases over time; if $\delta < 1$, the hazard rate increases over time; and if $\delta = 1$, the hazard rate is constant over time. A Weibull distribution with a constant hazard function is equivalent to an exponential distribution.

Table 13.2 Various Weibull Parameters in Terms of α and β in JMP

JMP Weibull	α	β
Wayne Nelson	$\alpha = \alpha$	$\beta = \beta$
Meeker and Escobar	$\eta = \alpha$	$\beta = \beta$
Tobias and Trindade	$c = \alpha$	$m = \beta$
Kececioglu	$\eta = \alpha$	$\beta = \beta$
Hosmer and Lemeshow	$\exp(X \beta) = \alpha$	$\lambda = \beta$
Blishke and Murthy	$\beta = \alpha$	$\alpha = \beta$

Table 13.2 Various Weibull Parameters in Terms of alpha and beta in JMP (Continued)

Kalbfleisch and Prentice	$\lambda = 1/\alpha$	$p = \beta$
JMP Extreme Value	$\lambda = \log(\alpha)$	$\delta = 1/\beta$
Meeker and Escobar s.e.v.	$\mu = \log(\alpha)$	$\sigma = 1/\beta$

Lognormal

The lognormal distribution is also very popular for modeling time-to-event data. The lognormal distribution is equivalent to the distribution where if you take the log of the values, the distribution is normal. If you want to fit a normal distribution to your data, you can take the $\exp()$ of it and model your data with a lognormal distribution. See [“Additional Examples of Fitting Parametric Survival”](#) on page 383 in the “Fit Parametric Survival” chapter.

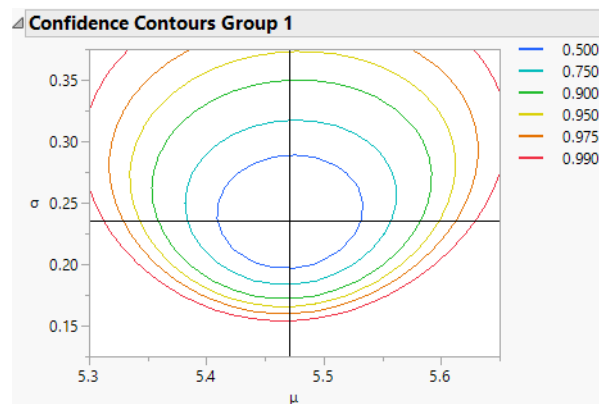
Additional Options

To see additional options for the exponential, Weibull, and lognormal fits, press Shift, click the red triangle next to Product-Limit Survival Fit, and select the desired fit.

Use these options to do the following tasks:

- Set the confidence level for the limits.
- Set the constrained value for theta (in the case of an exponential fit), sigma (in the case of a lognormal fit) or beta (in the case of a Weibull fit). See [“WeiBayes Analysis”](#) on page 360.
- Obtain a Confidence Contour Plot for the Weibull and lognormal fits (when there are no constrained values).

Figure 13.6 Confidence Contour Plot



WeiBayes Analysis

JMP can constrain the values of the Theta (Exponential), Beta (Weibull), and Sigma (LogNormal) parameters when fitting these distributions. This feature is needed in *WeiBayes* situations, for example:

- Where there are few or no failures
- There are existing historical values for beta
- There is still a need to estimate alpha

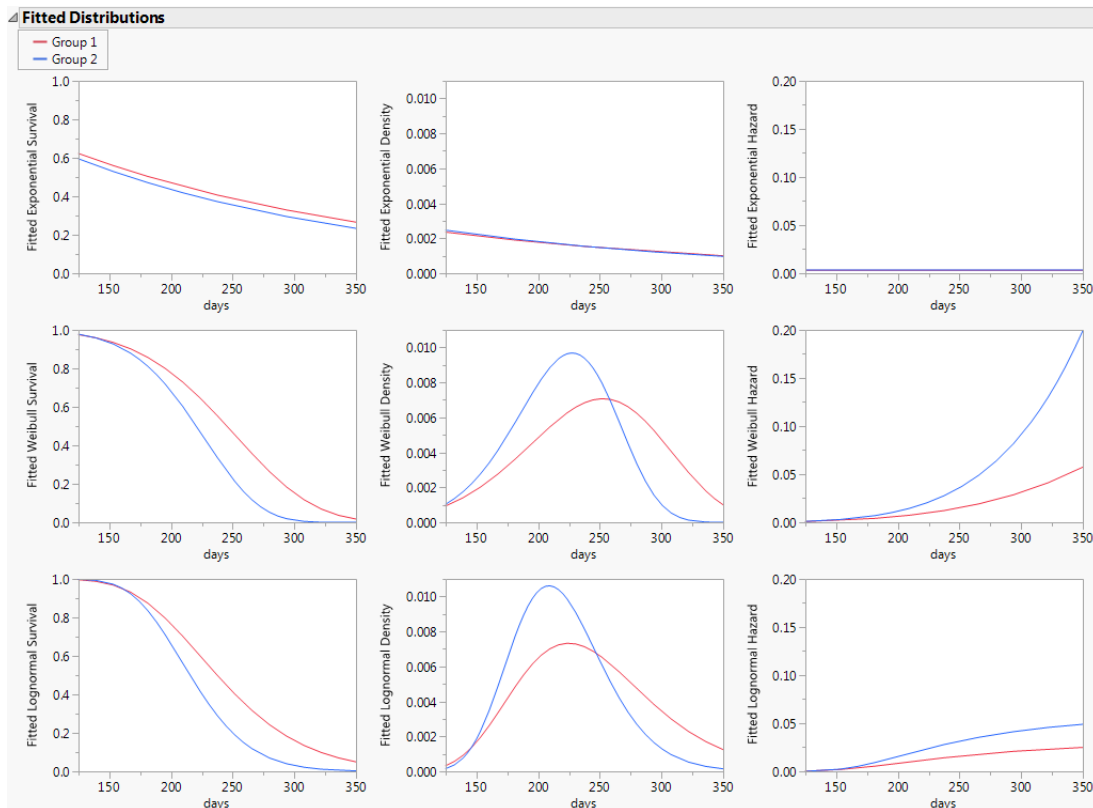
For more information about WeiBayes situations, see Abernethy ([1996](#)).

With no failures, the standard technique is to add a failure at the end. Then, the estimates reflect a type of lower bound on the alpha value, rather than a real estimate. However, the WeiBayes feature allows for a true estimation.

Fitted Distribution Plots

Use the Fitted Distribution Plots option to see Survival, Density, and Hazard plots for the exponential, Weibull, and lognormal distributions. The plots share the same axis scaling so that the distributions can be easily compared.

Figure 13.7 Fitted Distribution Plots for Three Distributions



These plots can be transferred to other graphs through the use of graphic scripts. To copy the graph, right-click in the plot to be copied and select **Edit > Copy Frame Contents**. Right-click in the destination plot and select **Edit > Paste Frame Contents**.

Competing Causes

Sometimes there are multiple causes of failure in a system. For example, suppose that a manufacturing process has several stages and the failure of any stage causes a failure of the whole system. If the different causes are independent, the failure times can be modeled by an estimation of the survival distribution for each cause. A censored estimation is undertaken for a given cause by treating all the event times that are not from that cause as censored observations.

The Competing Causes red triangle menu contains the following options:

Omit Causes Removes the specified cause value and recalculates the survival estimates.

Save Cause Coordinates Adds a new column to the current table called $\log(-\log(\text{Surv}))$. This information is often used to plot against the time variable with a grouping variable, such as the code for type of failure.

Weibull Lines Adds Weibull lines to the plot.

Hazard Plot Adds a Hazard Plot.

Simulate Creates a new data table containing time and cause information from the Weibull distribution, as estimated by the data.

Additional Examples of the Survival Platform

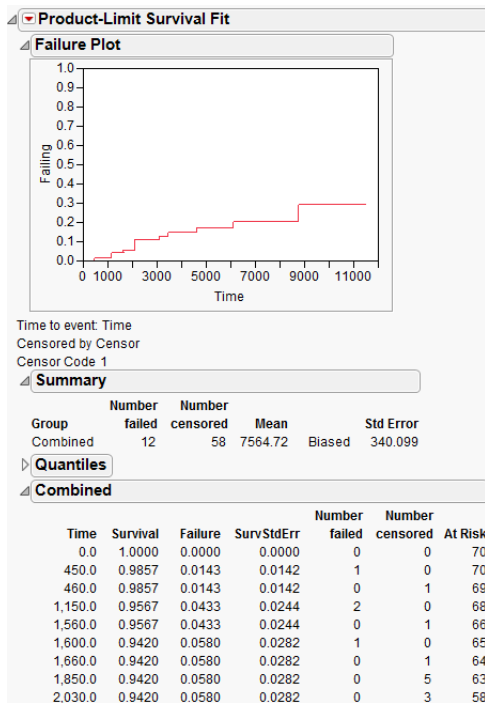
- [“Example of Survival Analysis”](#)
- [“Example of Competing Causes”](#)
- [“Example of Interval Censoring”](#)

Example of Survival Analysis

The failure of diesel generator fans was studied by Nelson ([1982](#), p. 133) and Meeker and Escobar ([1998](#), app. C1).

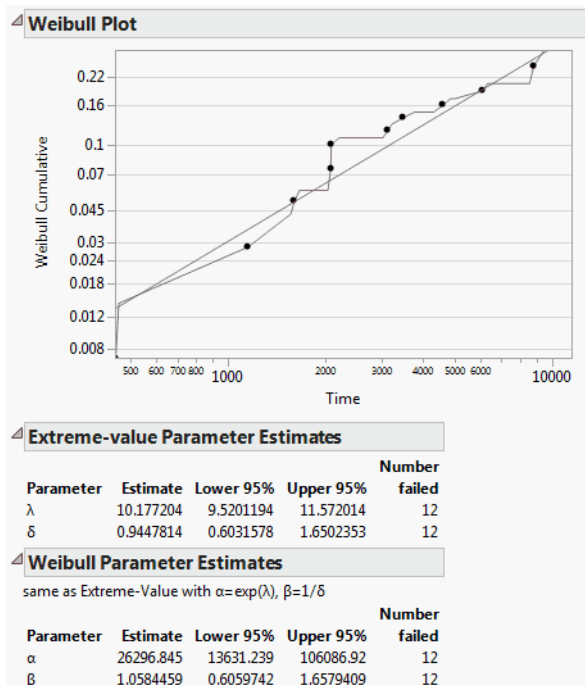
1. Select **Help > Sample Data Library** and open Reliability/Fan.jmp.
2. Select **Analyze > Reliability and Survival > Survival**.
3. Select Time and click **Y, Time to Event**.
4. Select Censor and click **Censor**.
5. Select the check box for **Plot Failure instead of Survival**.
6. Click **OK**.

Figure 13.8 Fan Initial Output



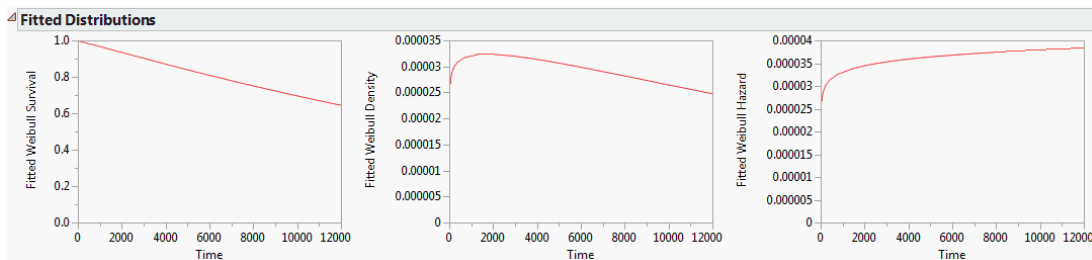
Notice that the probability of failure increases over time. Often the next step is to explore distributional fits, such as a Weibull model. Click the red triangle next to Product-Limit Survival Fit and select **Weibull Plot** and **Weibull Fit**.

Figure 13.9 Weibull Output for Fan Data



Because the fit is reasonable and the Beta estimate is near 1, you can conclude that this looks like an exponential distribution, which has a constant hazard rate. Click the Weibull Plot red triangle and select **Fitted Distribution Plots**. Three views of the Weibull fit appear.

Figure 13.10 Fitted Distribution Plots



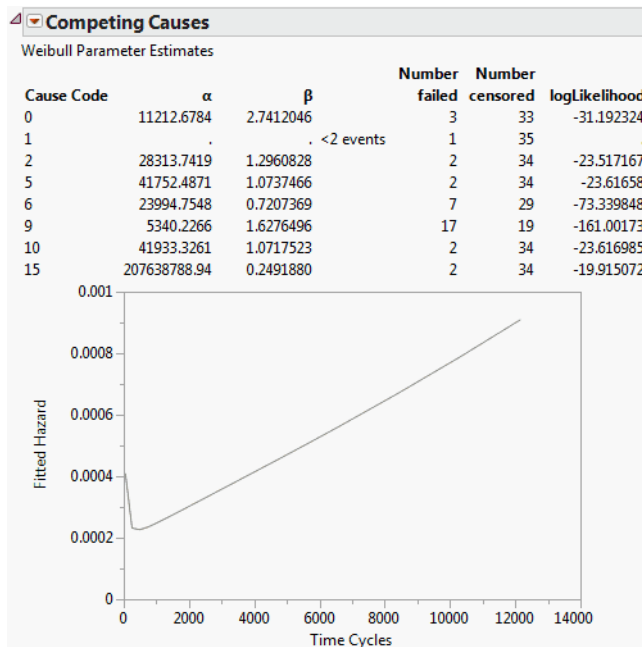
Example of Competing Causes

Nelson (1982) discusses the failure times of a small electrical appliance that has a number of causes of failure. One group (Group 2) of the data is represented in the JMP sample data table Appliance.jmp.

1. Select **Help > Sample Data Library** and open Reliability/Appliance.jmp.

2. Select **Analyze > Reliability and Survival > Survival**.
3. Select Time Cycles and click **Y, Time to Event**.
4. Click **OK**.
5. Click the red triangle next to Product-Limit Survival Fit and select **Competing Causes**.
6. Click Cause Code, and click **OK**.
7. Click the Competing Causes red triangle and select **Hazard Plot**.

Figure 13.11 Competing Causes Report and Hazard Plot

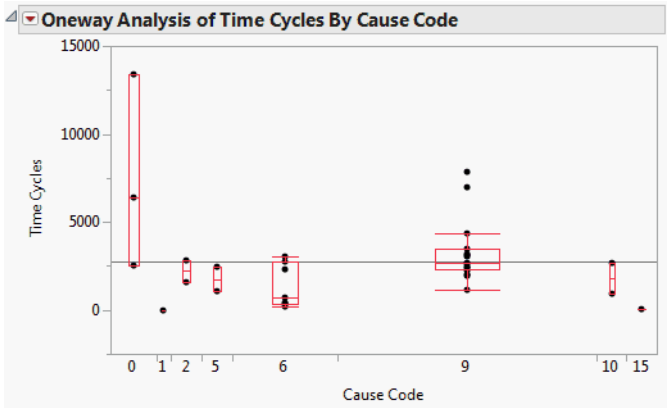


The survival distribution for the whole system is simply the product of the survival probabilities. The Competing Causes table shows the Weibull estimates of Alpha and Beta for each failure cause.

In this example, most of the failures were due to cause 9. Cause 1 occurred only once and could not produce good Weibull estimates. Cause 15 happened for very short times and resulted in a small beta and large alpha. Recall that alpha is the estimate of the 63.2% quantile of failure time, which means that causes with early failures often have very large alphas. If these causes do not result in early failures, then these causes do not usually cause later failures.

Figure 13.12 shows the Fit Y by X plot of Time Cycles by Cause Code with the Quantiles option in effect. This plot further illustrates how the alphas and betas relate to the failure distribution.

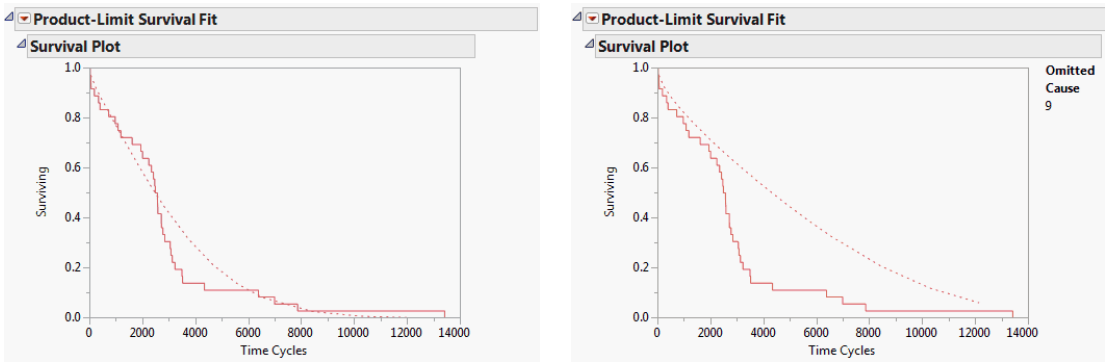
Figure 13.12 Fit Y by X Plot of Time Cycles by Cause Code



In this example, recall that cause 9 was the source of most of the failures. If cause 9 was corrected, how would that affect the survival due to the remaining causes? Select the **Omit Causes** option to remove a cause value and recalculate the survival estimates.

Figure 13.13 shows the survival plots with all competing causes and without cause 9. You can see that the survival rate (represented by the dashed line) without cause 9 does not improve much until 2,000 cycles. It then becomes much better and remains improved, even after 10,000 cycles.

Figure 13.13 Survival Plots with Omitted Causes



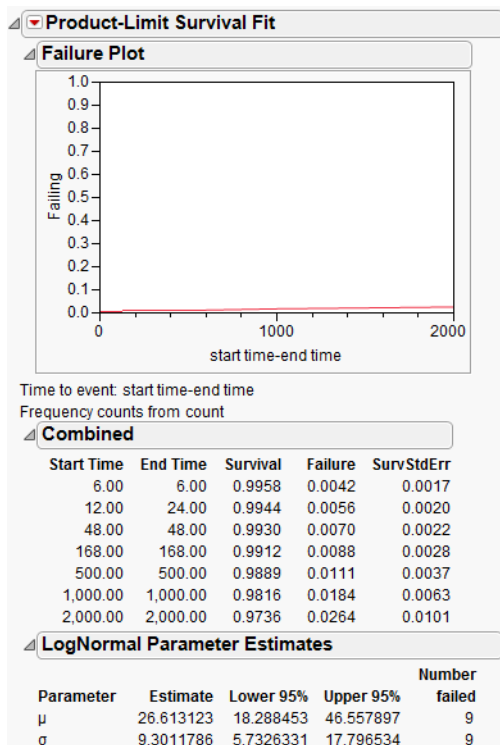
Example of Interval Censoring

With interval censored data, you know only that the events occurred in some time interval. The Turnbull method is used to obtain nonparametric estimates of the survival function.

In this example from Nelson (1990, p. 147), microprocessor units are tested and inspected at various times and the failed units are counted. Missing values in one of the columns indicate that you do not know the lower or upper limit, and therefore the event is left or right censored, respectively.

1. Select **Help > Sample Data Library** and open Reliability/Microprocessor Data.jmp.
2. Select **Analyze > Reliability and Survival > Survival**.
3. Select start time and end time and click **Y, Time to Event**.
4. Select count and click **Freq**.
5. Select the check box next to **Plot Failure instead of Survival**.
6. Click **OK**.
7. Click the red triangle next to Product-Limit Survival Fit and select **LogNormal Fit**.

Figure 13.14 Interval Censoring Output



The resulting Turnbull estimates are shown. Turnbull estimates might have gaps in time where the survival probability is not estimable. In this example, such gaps occur between 6 and 12, 24 and 48, 48 and 168, and so on.

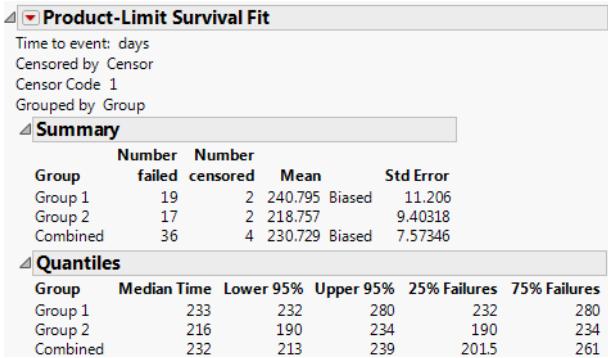
At this point, select a distribution to see its fitted estimates —in this case, a Lognormal distribution is fit. Notice that the failure plot shows very small failure rates for these data.

Statistical Reports for Survival Analysis

For data that is not interval censored, the initial reports show Summary and Quantiles data (Figure 13.15). The Summary data shows the number of failed and number of censored observations for each group (when there are groups) and for the whole study. The mean and standard deviations are also adjusted for censoring. For computational details about these statistics, see the LIFETEST Procedure chapter in SAS Institute Inc. (2020).

The Quantiles data shows time to failure statistics for individual and combined groups. These include the median survival time, with upper and lower 95% confidence limits. The median survival time is the time (number of days) at which half the subjects have failed. The quartile survival times (25% and 75%) are also included.

Figure 13.15 Summary Statistics for the Univariate Survival Analysis



The Summary report gives estimates for the mean survival time, as well as the standard error of the mean. The estimated mean survival time is defined as follows:

$$\hat{\mu} = \sum_{i=1}^D \hat{S}(t_{i-1})(t_i - t_{i-1}) \text{ with a standard error of } \hat{\sigma}(\hat{\mu}) = \sqrt{\frac{m}{m-1} \sum_{i=1}^{D-1} \frac{A_i^2}{n_i(n_i - d_i)}}$$

where

$$\hat{S}(t_i) = \prod_{j=1}^i \left(1 - \frac{d_j}{n_j}\right)$$

$$A_i = \sum_{j=i}^{D-1} \hat{S}(t_j)(t_{j+1} - t_j)$$

$$m = \sum_{j=1}^D d_j$$

$\hat{S}(t_i)$ is the survival distribution at time t_i

D is the number of distinct event times

n_i is the number of surviving units just prior to t_i

d_i is the number of units that fail at t_i

t_0 is defined to be 0

When there are multiple groups, the Tests Between Groups table provides statistical tests for homogeneity among the groups. Kalbfleisch and Prentice (1980, ch. 1), Hosmer and Lemeshow (1999, ch. 2), and Klein and Moeschberger (1997, ch. 7) discuss statistics and comparisons of survival curves.

Figure 13.16 Tests between Groups

Tests Between Groups			
Test	ChiSquare	DF	Prob>ChiSq
Log-Rank	3.1227	1	0.0772
Wilcoxon	2.6510	1	0.1035

Test Names two statistical tests of the hypothesis that the survival functions are the same across groups.

Chi-Square Provides the Chi-square approximations for the statistical tests.

The **Log-Rank** test places more weight on larger survival times and is more useful when the ratio of hazard functions in the groups being compared is approximately constant. The hazard function is the instantaneous failure rate at a given time. It is also called the *mortality rate* or *force of mortality*.

The **Wilcoxon** test places more weight on early survival times and is the optimum rank test if the error distribution is logistic. See Kalbfleisch and Prentice (1980).

DF Provides the degrees of freedom for the statistical tests.

Prob>ChiSq Lists the probability of obtaining a Chi-square value greater than the one computed if the survival functions are the same for all groups.

Figure 13.17 shows an example of the product-limit survival function estimates for one group.

Figure 13.17 Example of Survival Estimates Table

Group 1						
days	Survival	Failure	SurvStdErr	Number failed	Number censored	At Risk
0.000	1.0000	0.0000	0.0000	0	0	21
142.000	0.9524	0.0476	0.0465	1	0	21
156.000	0.9048	0.0952	0.0641	1	0	20
163.000	0.8571	0.1429	0.0764	1	0	19
198.000	0.8095	0.1905	0.0857	1	0	18
204.000	0.8095	0.1905	0.0857	0	1	17
205.000	0.7589	0.2411	0.0941	1	0	16
232.000	0.6577	0.3423	0.1053	2	0	15
233.000	0.4554	0.5446	0.1114	4	0	13
239.000	0.4048	0.5952	0.1099	1	0	9
240.000	0.3542	0.6458	0.1072	1	0	8
261.000	0.3036	0.6964	0.1031	1	0	7
280.000	0.2024	0.7976	0.0902	2	0	6
296.000	0.1012	0.8988	0.0678	2	0	4
323.000	0.0506	0.9494	0.0493	1	0	2
344.000	0.0506	0.9494	0.0493	0	1	1

Note: When the final time recorded is a censored observation, the report indicates a *biased* mean estimate. The biased mean estimate is a lower bound for the true mean.

Chapter 14

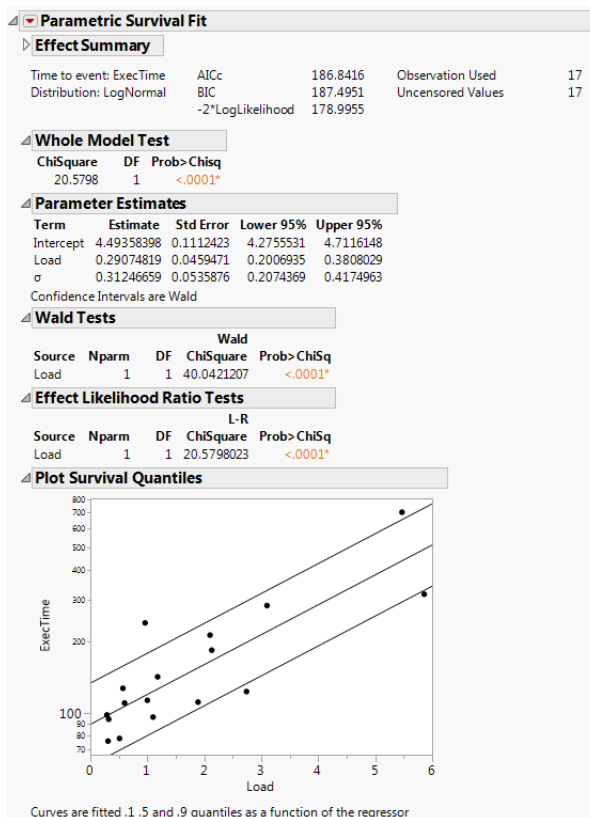
Fit Parametric Survival

Fit Survival Data Using Regression Models

Survival times can be expressed as a function of one or more variables. When this is the case, use a regression platform that fits a linear regression model while taking into account the survival distribution and censoring. The Fit Parametric Survival platform fits the time to event Y (with censoring) using linear regression models that can involve both location and scale effects. The fit is performed using the Weibull, lognormal, exponential, Fréchet, loglogistic, smallest extreme value (SEV), normal, largest extreme value (LEV), and logistic distributions.

Note: The Fit Parametric Survival platform is a slightly customized version of the Fit Model platform. You can also fit parametric survival models using the Nonlinear platform.

Figure 14.1 Example of a Parametric Survival Fit



Contents

Overview of the Fit Parametric Survival Platform.....	373
Example of Parametric Regression Survival Fitting.....	373
Launch the Fit Parametric Survival Platform	375
The Parametric Survival Fit Report	377
The Parametric Survival - All Distributions Report.....	378
Parametric Competing Cause Report.....	380
Fit Parametric Survival Options	380
Nonlinear Parametric Survival Models	382
Additional Examples of Fitting Parametric Survival.....	383
Arrhenius Accelerated Failure LogNormal Model	383
Interval-Censored Accelerated Failure Time Model	387
Analyze Censored Data Using the Nonlinear Platform	388
Left-Censored Data.....	390
Weibull Loss Function Using the Nonlinear Platform.....	391
Fitting Simple Survival Distributions Using the Nonlinear Platform.....	394
Statistical Details for the Fit Parametric Survival Platform.....	397
Loss Formulas for Survival Distributions	397

Overview of the Fit Parametric Survival Platform

Survival times can be expressed as a function of one or more variables. When this is the case, use a regression platform that fits a linear regression model while taking into account the survival distribution and censoring. The Fit Parametric Survival platform fits the time to event Y (with censoring) using linear regression models that can involve both location and scale effects. The fit is performed using the Weibull, lognormal, exponential, Fréchet, loglogistic, SEV, normal, LEV, and logistic distributions.

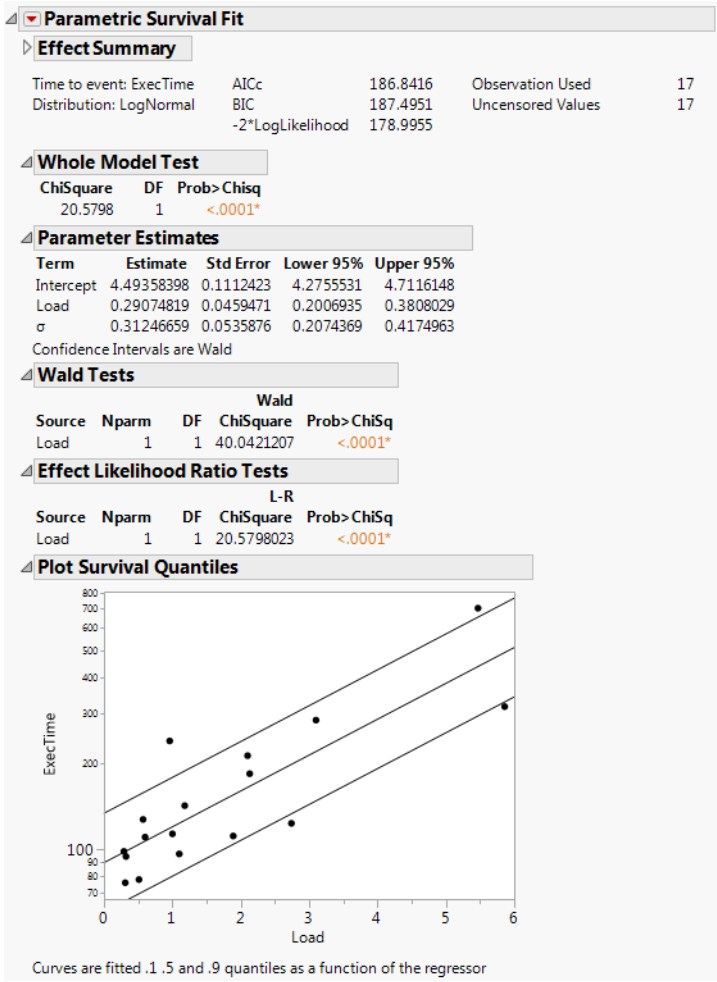
Example of Parametric Regression Survival Fitting

The data table Comptime.jmp contains data on the analysis of computer program execution time whose lognormal distribution depends on the effect Load.

Note: The data in Comptime.jmp comes from Meeker and Escobar ([1998](#), p. 434).

1. Select **Help > Sample Data Library** and open Reliability/Comptime.jmp.
2. Select **Analyze > Reliability and Survival > Fit Parametric Survival**.
3. Select ExecTime and click **Time to Event**.
4. Select Load and click **Add**.
5. Change the **Distribution** from **Weibull** to **Lognormal**.
6. Click **Run**.

Figure 14.2 Computing Time Output



When there is only one effect, a plot of the survival quantiles for three survival probabilities are shown as a function of the effect.

Time quantiles are desired for when 90% of jobs are finished under a system load of 5. See Meeker and Escobar (1998, p. 438).

- Click the Parametric Survival Fit red triangle and select **Estimate Quantile**.
- Type 5 in the first row beneath **Load**.
- Type 0.9 in the first row beneath **p**.
- Click **Go**.

Figure 14.3 Estimates of Time Quantile

Estimates of Quantile				
Load	p	Quantile	Lower 95%	Upper 95%
5	0.9	571.21575	401.29076	813.09482

The report estimates that 90% of the jobs will be done by 571 seconds of execution time under a system load of 5.

Launch the Fit Parametric Survival Platform

Launch the Fit Parametric Survival platform by selecting **Analyze > Reliability and Survival > Fit Parametric Survival**.

Figure 14.4 The Fit Parametric Survival Launch Window

Tip: To change the alpha level, click the Model Specification red triangle and select **Set Alpha Level**.

For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The Fit Parametric Survival launch window contains the following options:

Time to Event Contains the time to event or time to censoring. With interval censoring, specify two Y variables, where one Y variable gives the lower limit and the other Y variable gives the upper limit for each unit.

Censor Specifies a column with indicators to identify right-censored observations. Select the value that identifies right-censored observations from the Censor Code menu. The Censor column is used only when one Time to Event column is entered.

Freq Specifies a column that contains the frequencies or counts of observations for each row when there are multiple units recorded.

Cause Specifies a column that contains multiple failure causes. This column is particularly useful for estimating competing causes. A separate parametric fit is performed for each cause value. Failure events can be coded with either numeric or categorical (labels) values.

By Performs a separate analysis for each level of a classification or grouping variable.

Location and Scale Effects Specifies location and scale effects. For more information about the Construct Model Effects options, see *Fitting Linear Models*.

Personality Indicates the fitting method. Parametric Survival should always be selected.

Distribution Choose the desired response distribution that is appropriate for your data. Choose the All Distributions option to fit all the distributions and compare the fits. If you choose All Distributions, the report shows a comparison of the distribution fits. See [“The Parametric Survival - All Distributions Report”](#) on page 378.

Note: By default, the All Distributions option fits a model for the log-location-scale distributions that appear above All Distributions in the Distribution menu. Select **Preferences > Platforms > Fit Parametric Survival > Include location-scale distributions in All Distributions** to change the behavior of the All Distributions option to also include the location-scale distributions.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

The Parametric Survival Fit Report

If you select All Distributions in the launch window, a Parametric Survival Fit report appears for each distribution. If you specify a Cause column in the launch window, a Parametric Survival Fit report appears for each cause. Otherwise, only one Parametric Survival Fit report appears. Each Parametric Survival Fit report contains the following:

Effect Summary Shows an interactive report that enables you to add or remove effects from the model. See *Fitting Linear Models*.

Model Fit Details The Time to event shows which Y column is specified, and the Distribution shows which distribution is fit. AICc, BIC, and -2Loglikelihood are all measures of the model fit. These measures allow for comparisons to other model fits. Observation Used and Uncensored Values are summary statistics for the data. See *Fitting Linear Models*.

Whole Model Test Compares the complete fit with an intercept-only fit. If there is only an intercept term, the fit is the same as that from the Life Distribution platform.

Parameter Estimates Shows the estimates of the regression parameters.

 A link to launch the Generalized Regression platform appears below the Parameter Estimates table. The link enables you to perform variable selection using the Generalized Regression platform and appears under the following circumstances:

- The model has no scale effects.
- No Cause column is specified in the launch window.
- The Distribution specified in the launch window is Normal, Lognormal, or Weibull.

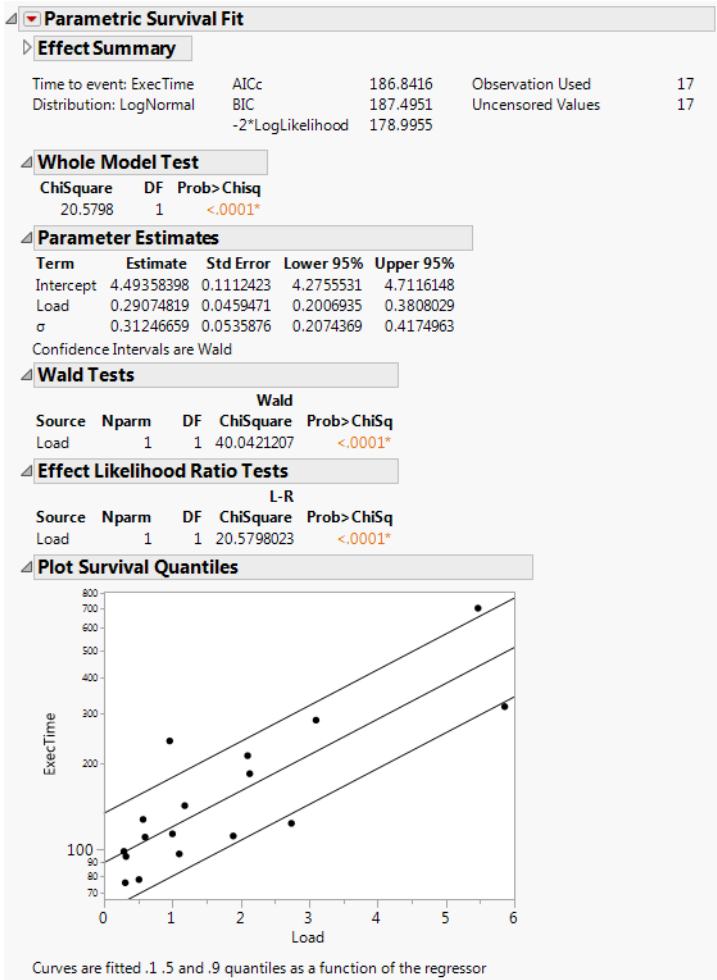
Alternate Parameterization (Available only for the Weibull distribution.) Shows the parameter estimates for the α and β parameterization of the Weibull distribution. For more information about this parameterization, see “Weibull” on page 85 in the “Life Distribution” chapter.

Wald Tests Shows a Wald Chi-square test for each term in the model.

Effect Likelihood Ratio Tests Compare the log-likelihood from the fitted model to one that removes each term from the model individually.

Plot Survival Quantiles Shows the data points plotted with the 0.1, 0.5, and 0.9 quantiles.

Figure 14.5 The Parametric Survival Fit Report



The Parametric Survival - All Distributions Report

The Parametric Survival - All Distributions report appears only when you select All Distributions in the launch window. By default, this report contains a Model Comparison report and Distribution Overlay plot. The Quantile Function Overlay plot is available in the red triangle menu next to Parametric Survival - All Distributions.

Model Comparison Table that lists fit statistics (AICc and BIC) for the fitted distributions. The distributions with the smallest AICc and BIC values are labeled in the right-most column. If one distribution has the smallest value for both AICc and BIC, that distribution is labeled “Best”. The Parametric Survival Fit report corresponding to the distribution with

the smallest AICc is open by default. For more information about these statistics, see *Fitting Linear Models*.

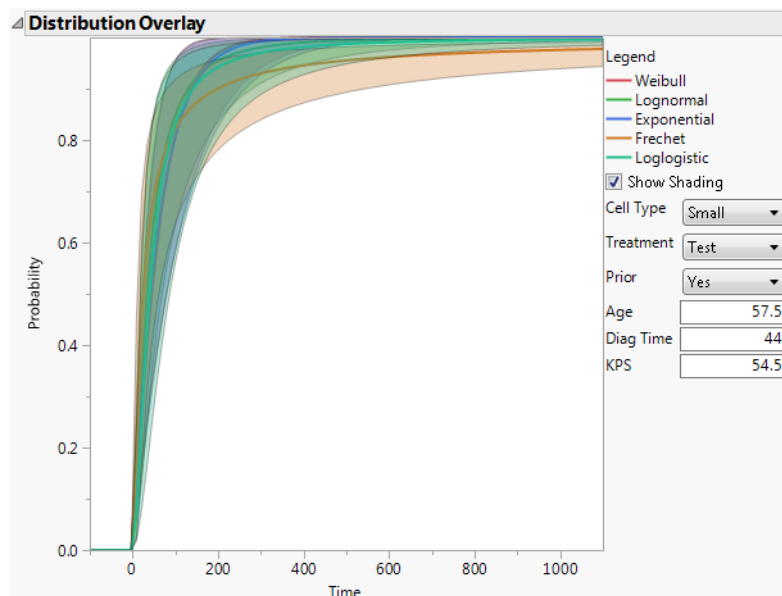
Distribution Overlay Plot of overlaid distribution functions for the fitted distributions at specified values of the effects.

Quantile Function Overlay Plot of overlaid quantile functions for the fitted distributions at specified values of the effects.

Both the Distribution Overlay and Quantile Function Overlay plots show curves for each of the fitted distributions overlaid on the same graph. By default, each curve has a shaded Wald-based confidence interval. To the right of each plot, there is legend, an option to show shading of confidence intervals, and controls that enable you to specify different values of the effects. Figure 14.6 shows an example of a Distribution Overlay plot.

Note: You can change the α level for the shaded confidence intervals by selecting Set Alpha Level from the red triangle menu in the Fit Model launch window. The default α level is 0.05.

Figure 14.6 Distribution Overlay Plot



Parametric Competing Cause Report

If you specify a Cause column in the launch window, the Parametric Competing Cause report appears. If you also specify All Distributions in the launch window, this report is labeled Parametric Competing Cause - All Distributions. The Parametric Competing Cause report contains the following:

Summary by Cause Table that lists fit statistics (AICc and BIC) for each cause. If All Distributions is selected for the Distribution option in the launch window, this table includes fit statistics for each cause within each distribution fit.

Model Comparison (Available only when All Distributions is selected for the Distribution option in the launch window.) Shows a table that lists fit statistics (AICc and BIC) for the fitted distributions. The distributions with the smallest AICc and BIC values are labeled in the right-most column. If one distribution has the smallest value for both AICc and BIC, that distribution is labeled “Best”. The Parametric Survival Fit report corresponding to the distribution with the smallest AICc is open by default.

For more information about the AICc and BIC statistics, see *Fitting Linear Models*.

Fit Parametric Survival Options

The Parametric Survival Fit red triangle menu contains the following options:

Likelihood Ratio Tests Produces tests that compare the log-likelihood from the fitted model to one that removes each term from the model individually.

Wald Tests Produces chi-square test statistics and p -values for Wald tests of whether each parameter is zero.

Likelihood Confidence Intervals Specifies the type of confidence intervals shown in the Parameter Estimates table for each parameter. When this option is selected, a profile likelihood confidence interval appears. Otherwise, a Wald interval is shown. In the report, the interval type is noted below the Parameter Estimates table. This option is on by default when the computational time for the profile likelihood confidence intervals is not large.

Note: You can change the α level for the confidence intervals by selecting Set Alpha Level from the red triangle menu in the Fit Model launch window. The default α level is 0.05.

Correlation of Estimates Produces a correlation matrix for the model effects with each other and with the parameter of the fitting distribution.

Covariance of Estimates Produces a covariance matrix for the model effects with each other and with the parameter of the fitting distribution.

Estimate Survival Probability Estimates the failure and survival probabilities for the given time values. Specify effect values and one or more time values. JMP then calculates the survival and failure probabilities with 95% confidence limits for all possible combinations of the entries.

Estimate Quantile Estimates the quantiles for the given probabilities. Specify effect values and one or more quantile probabilities. JMP then calculates the time quantiles and 95% confidence limits for all possible combinations of the entries.

Note: For the Estimate Survival Probability and Estimate Quantile options, you can change the alpha level from the default of 0.05.

Residual Probability Plot Shows a probability plot of the standardized residuals.

Save Residuals Saves the residuals to a new column in the data table. For interval-censored observations, two columns of residuals are saved to the data table.

Distribution Profiler Shows the response surfaces of the failure probability versus individual explanatory and response variables.

Quantile Profiler Shows the response surfaces of the response variable versus the explanatory variable and the failure probability.

Distribution Plot by Level Combinations Shows three probability plots for assessing model fit. The plots show different lines for each combination of the X levels.

Separate Location A probability plot assuming equal scale parameters and separate location parameters. This is useful for assessing the parallelism assumption.

Separate Location and Scale A probability plot assuming different scale and location parameters. This is useful for assessing if the distribution is adequate for the data. This plot is not shown for the Exponential distribution.

Regression A probability plot for which the distribution parameters are functions of the X variables.

Save Probability Formula Saves the estimated probability formula to a new column in the data table.

Save Quantile Formula Saves the estimated quantile formula to a new column in the data table. Selecting this option displays a pop-up window, asking you to enter a probability value for the quantile of interest.

Publish Probability Formula Creates a probability formula and saves it as a formula column script in the Formula Depot platform. If a Formula Depot report is not open, this option creates a Formula Depot report. See *Predictive and Specialized Modeling*.

Publish Quantile Formula Creates a quantile formula and saves it as a formula column script in the Formula Depot platform. If a Formula Depot report is not open, this option creates a Formula Depot report. See *Predictive and Specialized Modeling*.

Model Dialog Relaunches the launch window.

Effect Summary Shows the interactive Effect Summary report that enables you to add or remove effects from the model. See *Fitting Linear Models*.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Nonlinear Parametric Survival Models

Use the Nonlinear platform for survival models in the following instances:

- The model is nonlinear.
- You need a distribution other than Weibull, lognormal, exponential, Fréchet, loglogistic, SEV, normal, LEV, or logistic.
- You have censoring that is not the usual right, left, or interval censoring.

With the ability to estimate parameters in specified loss functions, the Nonlinear platform becomes a powerful tool for fitting maximum likelihood models. For complete information about the Nonlinear platform, see *Predictive and Specialized Modeling*.

To fit a nonlinear model when data are censored, you must first use the formula editor to create a parametric equation that represents a loss function adjusted for censored observations. Then use the Nonlinear platform to estimate the parameters using maximum likelihood.

Additional Examples of Fitting Parametric Survival

- [“Arrhenius Accelerated Failure LogNormal Model”](#)
- [“Interval-Censored Accelerated Failure Time Model”](#)
- [“Analyze Censored Data Using the Nonlinear Platform”](#)
- [“Left-Censored Data”](#)
- [“Weibull Loss Function Using the Nonlinear Platform”](#)
- [“Fitting Simple Survival Distributions Using the Nonlinear Platform”](#)

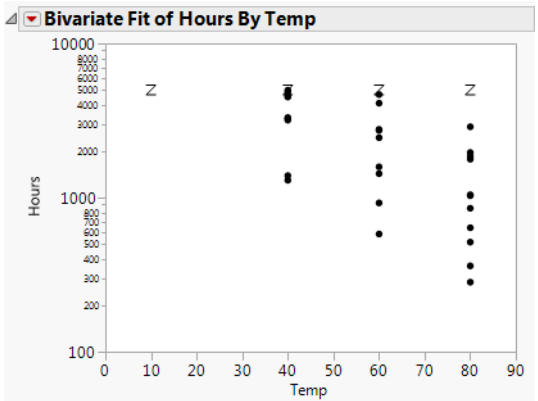
Arrhenius Accelerated Failure LogNormal Model

In the Devalt.jmp data, units are stressed by heating, in order to make them fail soon enough to obtain enough failures to fit the distribution.

Note: The data in Devalt.jmp comes from Meeker and Escobar ([1998](#), p. 493).

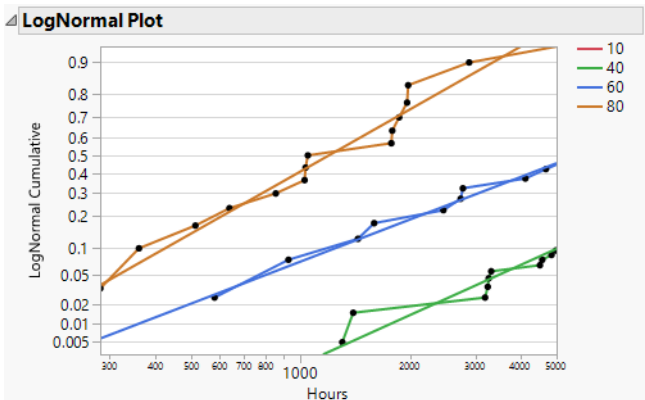
1. Select **Help > Sample Data Library** and open Reliability/Devalt.jmp.
First, use the Bivariate platform to see a plot of hours by temperature using the log scale for time.
2. Select **Analyze > Fit Y by X**.
3. Select Hours and click **Y, Response**.
4. Select Temp and click **X, Factor**.
5. Click **OK**.

Figure 14.7 Bivariate Plot of Hours by Log Temp



- Next, use the survival platform to produce a LogNormal plot of the data for each temperature.
6. Select **Analyze > Reliability and Survival > Survival**.
 7. Select Hours and click **Y, Time to Event**.
 8. Select Censor and click **Censor**.
 9. Select Temp and click **Grouping**.
 10. Select Weight and click **Freq**.
 11. Click **OK**.
 12. Click the red triangle next to Product-Limit Survival Fit and select **LogNormal Plot** and **LogNormal Fit**.
 13. Click **OK**.

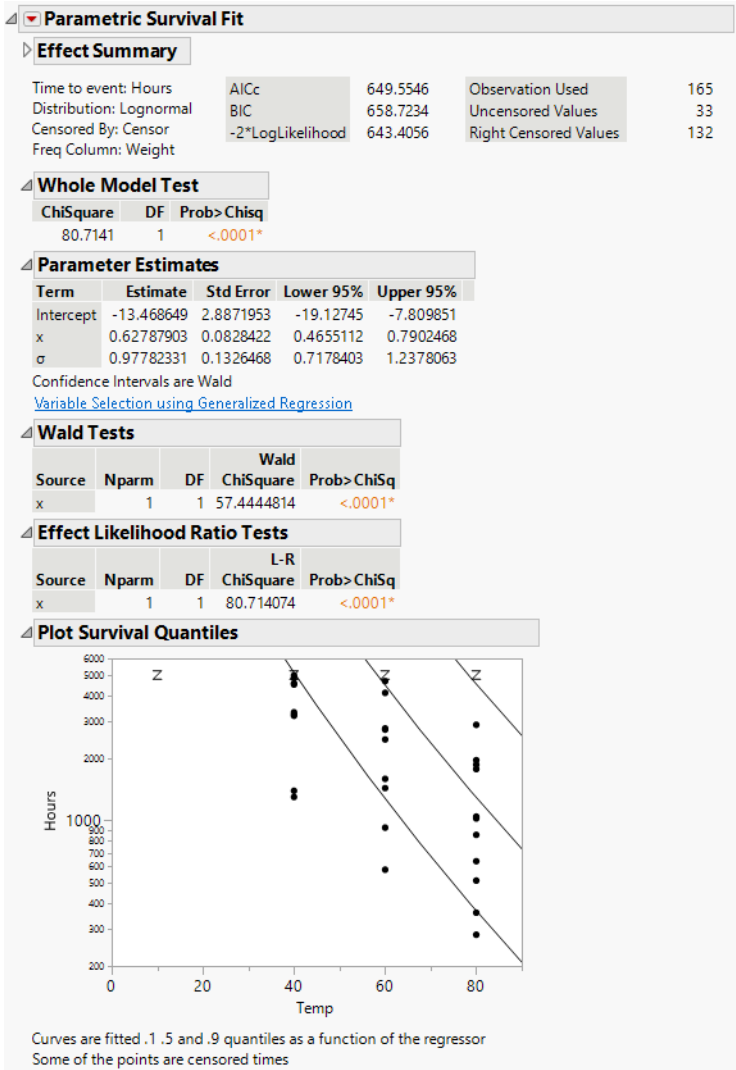
Figure 14.8 Lognormal Plot



Next, use the Fit Parametric Survival platform to fit one model using an effect for temperature.

14. Select **Analyze > Reliability and Survival > Fit Parametric Survival**.
15. Select Hours and click **Time to Event**.
16. Select x and click **Add**.
17. Select Censor and click **Censor**.
18. Select Weight and click **Freq**.
19. Change the **Distribution** type to **Lognormal**.
20. Click **Run**.

Figure 14.9 Devalt Parametric Output



The result shows the regression fit of the data:

- If there is only one effect and it is continuous, then a plot of the survival as a function of the effect is shown. Lines are at 0.1, 0.5, and 0.9 survival probabilities.
- If the effect column has a formula in terms of one other column, as in this case, the plot is done with respect to the inner column. In this case, the effect was the column x, but the plot is done with respect to Temp, of which x is a function.

Finally, get estimates of survival probabilities extrapolated to a temperature of 10 degrees Celsius for the times 30000 and 10000 hours.

21. Click the Parametric Survival Fit red triangle and select **Estimate Survival Probability**.
22. Enter the values shown in Figure 14.10 into the Dialog to Estimate Survival.

The Arrhenius transformation of 10 degrees is 40.9853, the effect value.

Figure 14.10 Estimating Survival Probabilities

Dialog to Estimate Survival

Enter term values and values on the right, and then click Go.

x	Time	Alpha
40.9853	30000	0.0500
.	10000	
.	.	.
.	.	.
.	.	.
.	.	.
.	.	.
.	.	.
.	.	.
.	.	.
.	.	.

Go

23. Click **Go**.

Figure 14.11 Survival Probabilities

Estimates of Survival

x	Time	Prob Failure	Lower 95%	Upper 95%	Prob Survival
40.9853	30000	0.0227781	0.002432	0.118393	0.9772219
40.9853	10000	0.0008951	4.353e-5	0.0101183	0.9991049

The Estimates of Survival report shows the estimates and a confidence interval.

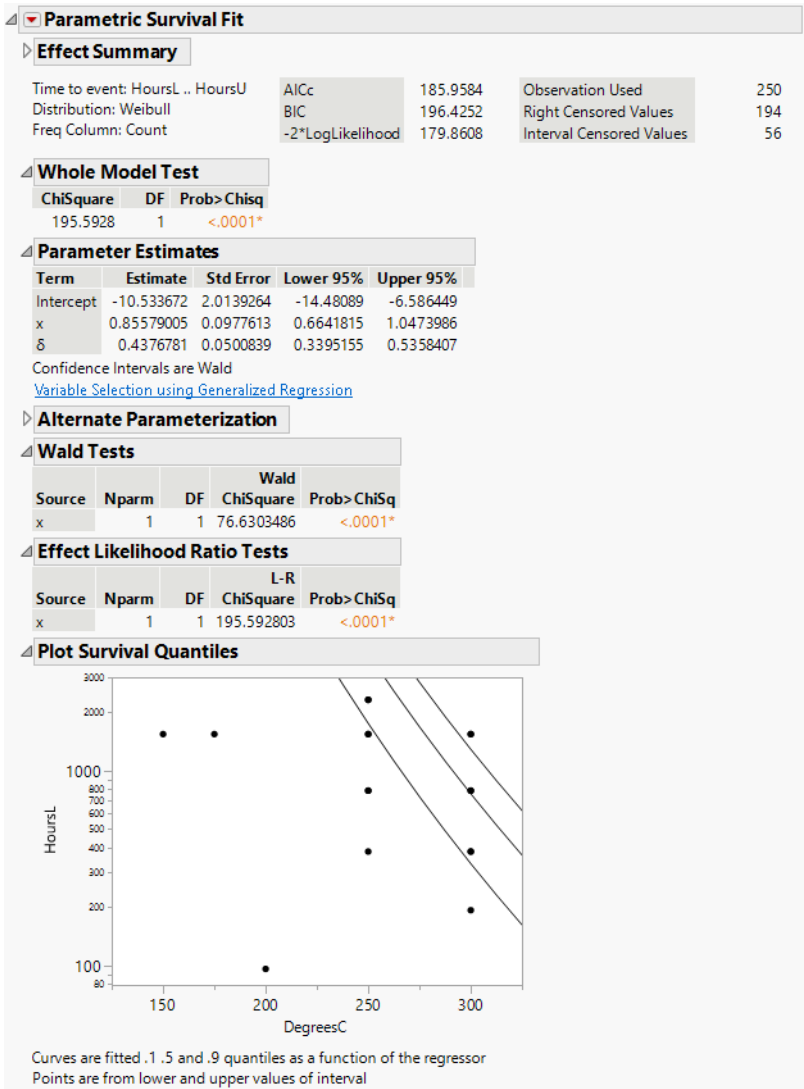
Interval-Censored Accelerated Failure Time Model

The ICdevice02.jmp data shows failures that were found to have happened between inspection intervals. The model uses two y -variables, containing the upper and lower bounds on the failure times. Right-censored times are shown with missing upper bounds.

Note: The data in ICdevice02.jmp comes from Meeker and Escobar (1998, p. 640).

1. Select **Help > Sample Data Library** and open Reliability/ICdevice02.jmp.
2. Select **Analyze > Reliability and Survival > Fit Parametric Survival**.
3. Select HoursL and HoursU and click **Time to Event**.
4. Select Count and click **Freq**.
5. Select x and click **Add**.
6. Click **Run**.

Figure 14.12 ICDevice Output



The resulting regression shows a plot of time by degrees.

Analyze Censored Data Using the Nonlinear Platform

You can analyze left-censored data using the Nonlinear platform. For the left-censored data, zero is the censored value because it also represents the smallest known time for an observation.

Note: The Tobit model is popular in economics for responses that must be positive or zero, with zero representing a censored point.

Note the following about the Tobit2.jmp data table:

- The response variable is a measure of the durability of a product and cannot be less than zero (Durable, is left-censored at zero).
- Age and Liquidity are independent variables.
- The table also includes the model and Tobit loss function. The model in residual form is $\text{durable} - (b_0 + b_1 * \text{age} + b_2 * \text{liquidity})$. To see the formula associated with Tobit Loss, right-click the column and select **Formula**.

Fit the Tobit model:

1. Select **Help > Sample Data Library** and open Reliability/Tobit2.jmp.
2. Select **Analyze > Specialized Modeling > Nonlinear**.
3. Select Model and click **X, Predictor Formula**.
4. Select Tobit Loss and click **Loss**.
5. Click **OK**.
6. Click **Go**.
7. Click **Confidence Limits**.

Figure 14.13 Solution Report

Solution				
	Loss	DFE	Avg Loss	Sqrt Avg Loss
	28.92596097	16	1.8078726	1.3445715
Parameter	Estimate	ApproxStdErr	Lower CL	Upper CL
b0	15.277120063	16.0327167	-22.059913	54.7012259
b1	-0.134007501	0.21893135	-0.7365041	0.32841861
b2	-0.045135584	0.05826851	-0.1858906	0.09330941
Sigma	5.5693492202	1.72814365	3.31613117	11.5291865
Solved By: Analytic NR				
Correlation of Estimates				

Left-Censored Data

The Tobit model from the previous section assumes a normal distribution that is censored at zero. An observation of zero is considered to be left-censored. In the Fit Parametric Survival platform, left-censored observations are specified using two response columns. In this example, you create a new column that indicates observations for which left censoring has occurred. For left-censored observations, the new column contains a missing value. Otherwise, it contains the observed value of *durable*. The new column is then used as the left side of interval-censored observations, and the existing column *durable* is used as the right side of the intervals.

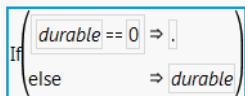
1. Select **Help > Sample Data Library** and open Reliability/Tobit2.jmp.

Create the Left Censoring Column:

2. Select **Cols > New Columns**.
3. Type *durable0* for **Column Name**.
4. Select **Column Properties** and click **Formula**.
5. Select **Conditional > If** and select *durable*.
6. Select **Comparison > a == b**, type 0, and press Enter.
7. Select the box labeled “else clause” and select *durable*.

Figure 14.14 shows the completed formula.

Figure 14.14 Column Formula for *durable0*



8. Click **OK**.
9. Click **OK**.

Fit the Tobit Model:

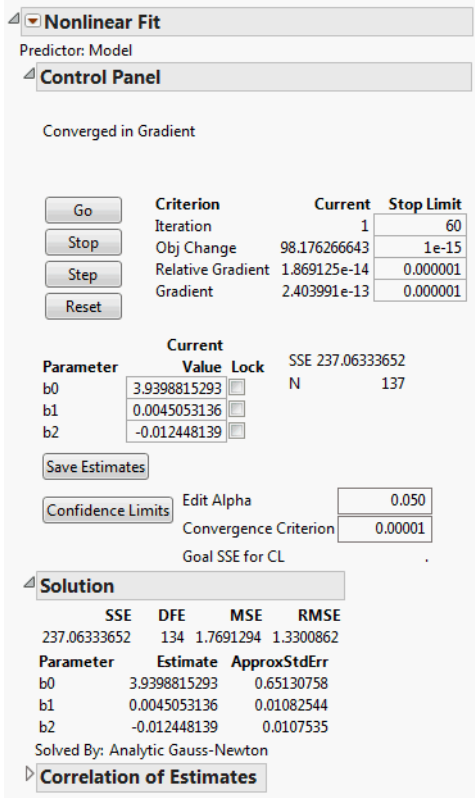
10. Select **Analyze > Reliability and Survival > Fit Parametric Survival**.
11. Select *durable0* and click **Time to Event**.
12. Select *durable* and click **Time to Event**.

You must use two response columns to specify left-censored observations. The direction of censoring is determined by the order of the columns in the Time to Event role.

13. Select *age* and *liquidity* and click **Add**.
14. Change the **Distribution** type to **Normal**.

- 3. Select Model and click **X, Predictor Formula**.
- 4. Click **OK**.
- 5. Click **Go**.

Figure 14.16 Initial Parameter Values in the Nonlinear Fit Control Panel



The report computes the least squares parameter estimates for this model.

- 6. Click **Save Estimates**.

The parameter estimates in the column formulas are set to those estimated by this initial nonlinear fitting process.

The Weibull column contains the Weibull formula, explained in [“Weibull Loss Function”](#) on page 397.

To continue with the fitting process:

- 7. Select **Analyze > Specialized Modeling > Nonlinear** again.
- 8. Select Model and click **X, Predictor Formula**.
- 9. Select Weibull loss and click **Loss**.

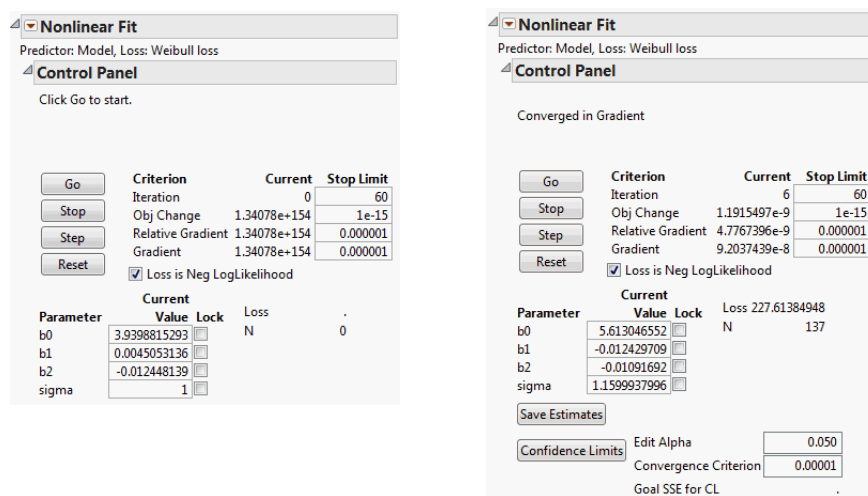
10. Click **OK**.

The Nonlinear Fit Control Panel on the left in Figure 14.17 appears. There is now the additional parameter called sigma in the loss function. Because it is in the denominator of a fraction, a starting value of 1 is reasonable for sigma. When using any loss function other than the default, the **Loss is Neg LogLikelihood** box on the Control Panel is checked by default.

11. Click **Go**.

The fitting process converges as shown on the right in Figure 14.17.

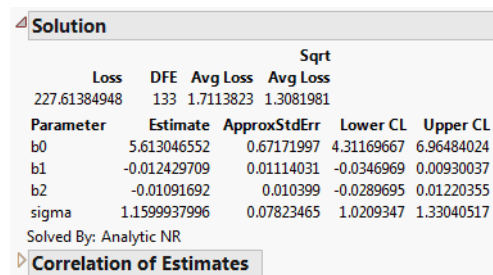
Figure 14.17 Nonlinear Model with Custom Loss Function



The fitting process estimates the parameters by maximizing the negative log of the Weibull likelihood function.

12. (Optional) Click **Confidence Limits** to show lower and upper 95% confidence limits for the parameters in the Solution table.

Figure 14.18 Solution Report



Note: Because the confidence limits are profile likelihood confidence intervals instead of the standard asymptotic confidence intervals, they can take time to compute.

You can also run the model with the predefined exponential and lognormal loss functions. Before you fit another model, reset the parameter estimates to the least squares estimates, as they might not converge otherwise. To reset the parameter estimates:

13. (Optional) Click the Nonlinear Fit red triangle and select **Revert to Original Parameters**.

Fitting Simple Survival Distributions Using the Nonlinear Platform

The following examples show how to use maximum likelihood methods to estimate distributions from time-censored data when there are no effects other than the censor status.

The Loss Function Templates folder has templates with formulas for exponential, extreme value, loglogistic, lognormal, normal, and one-and two-parameter Weibull loss functions. To use these loss functions, copy your time and censor values into the Time and censor columns of the loss function template. To run the model, select **Nonlinear** and assign the loss column as the **Loss** variable. Because both the response model and the censor status are included in the loss function and there are no other effects, you do not need a prediction column (model variable).

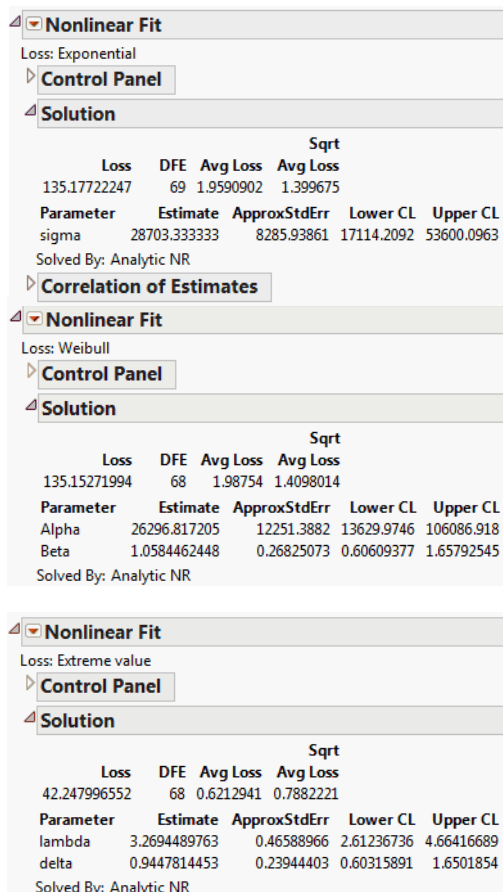
Exponential, Weibull, and Extreme-Value Loss Function

The Fan.jmp data table can be used to illustrate the Exponential, Weibull, and Extreme value loss functions discussed in Nelson (1982). The data are from a study of 70 diesel fans that accumulated a total of 344,440 hours in service. The fans were placed in service at different times. The response is failure time of the fans or run time, if censored.

Tip: To view the formulas for the loss functions, in the Fan.jmp data table, right-click the Exponential, Weibull, and Extreme value columns and select **Formula**.

1. Select **Help > Sample Data Library** and open Reliability/Fan.jmp.
2. Select **Analyze > Specialized Modeling > Nonlinear**.
3. Select Exponential and click **Loss**.
4. Click **OK**.
5. Make sure that the Loss is Neg LogLikelihood check box is selected.
6. Click **Go**.
7. Click **Confidence Limits**.
8. Repeat these steps, but select Weibull and Extreme value instead of Exponential.

Figure 14.19 Nonlinear Fit Results



Lognormal Loss Function

The Locomotive.jmp data can be used to illustrate a lognormal loss. The lognormal distribution is useful when the range of the data is several powers of e .

Tip: To view the formula for the loss function, in the Locomotive.jmp data table, right-click the logNormal column and select **Formula**.

The lognormal loss function can be very sensitive to starting values for its parameters. Because the lognormal distribution is similar to the normal distribution, you can create a new variable that is the natural log of Time and use **Distribution** to find the mean and standard deviation of this column. Then, use those values as starting values for the Nonlinear platform. In this example, the mean of the natural log of Time is 4.72 and the standard deviation is 0.35.

1. Select **Help > Sample Data Library** and open Reliability/Locomotive.jmp.

- 2. Select **Analyze > Specialized Modeling > Nonlinear.**
- 3. Select logNormal and click **Loss.**
- 4. Click **OK.**
- 5. Type 4.72 in the box next to Mu.
- 6. Type 0.35 in the box next to sigma.
- 7. Click **Go.**
- 8. Click **Confidence Limits.**

Figure 14.20 Solution Report

Solution				
		Sqrt		
Loss	DFE	Avg Loss	Avg Loss	
237.09354534	94	2.5222718	1.5881662	
Parameter	Estimate	ApproxStdErr	Lower CL	Upper CL
Mu	5.1169247193	0.10416697	4.93735853	5.35966436
sigma	0.7054940302	0.0932136	0.55446171	0.93294904
Solved By: Numerical NR				
Correlation of Estimates				

The maximum likelihood estimates of the lognormal parameters are 5.11692 for Mu and 0.7055 for Sigma. The corresponding estimate of the median of the lognormal distribution is the antilog of 5.11692, ($e^{5.11692}$), which is approximately 167. This represents the typical life for a locomotive engine.

Statistical Details for the Fit Parametric Survival Platform

This section contains statistical details for the Fit Parametric Survival platform.

Loss Formulas for Survival Distributions

The following formulas are for the negative log-likelihoods to fit common parametric models. Each formula uses the calculator `if` conditional function with the uncensored case of the conditional first and the right-censored case as the `Else` clause. You can copy these formulas from tables in the Loss Function Templates folder in Sample Data and paste them into your data table.

Exponential Loss Function

In the exponential loss function shown here, `sigma` represents the mean of the exponential distribution and `Time` is the age at failure.

$$\text{- IfMZ} \left(\begin{array}{l} \text{Censor} == 0 \Rightarrow -\text{Log}\left(\frac{\text{sigma}}{\text{Time}}\right) - \frac{\text{Time}}{\text{sigma}} \\ \text{else} \quad \quad \quad \Rightarrow -\left(\frac{\text{Time}}{\text{sigma}}\right) \end{array} \right)$$

A characteristic of the exponential distribution is that the instantaneous failure rate remains constant over time. This means that the chance of failure for any subject during a given length of time is the same regardless of how long a subject has been in the study.

Weibull Loss Function

The Weibull density function often provides a good model for the lifetime distributions. You can use the Survival platform for an initial investigation of data to determine whether the Weibull loss function is appropriate for your data.

$$\text{- IfMZ} \left(\begin{array}{l} \text{Censor} == 0 \Rightarrow \frac{\text{Model}}{\text{sigma}} - \text{Exp}\left(\frac{\text{Model}}{\text{sigma}}\right) - \text{Log}\left(\frac{\text{sigma}}{\text{Model}}\right) \\ \text{else} \quad \quad \quad \Rightarrow -\text{Exp}\left(\frac{\text{Model}}{\text{sigma}}\right) \end{array} \right)$$

There are examples of one-parameter, two-parameter, and extreme-value functions in the Loss Function Templates folder.

Lognormal Loss Function

The formula shown below is the lognormal loss function where $\text{Normal Distribution}(\text{model}/\text{sigma})$ is the standard normal distribution function. The hazard function has value 0 at $t = 0$, increases to a maximum, and then decreases. The hazard function approaches zero as t becomes large.

$$- \text{If} \left(\begin{array}{l} \text{censor} == 0 \Rightarrow -0.5 * \left(\frac{\text{Model}}{\text{sigma}} \right)^2 - 0.5 * \text{Log}(2 * \text{Pi}) - \text{Log}(\text{sigma}) \\ \text{else} \quad \Rightarrow \text{Log} \left(1 - \text{Normal Distribution} \left(\frac{\text{Model}}{\text{sigma}} \right) \right) \end{array} \right)$$

Loglogistic Loss Function

If Y is distributed as the logistic distribution, $\text{Exp}(Y)$ is distributed as the loglogistic distribution.

$$- \text{IfMZ} \left(\begin{array}{l} \text{censor} == 0 \Rightarrow \frac{\text{Model}}{\text{sigma}} - 2 * \text{Log} \left(1 + \text{Exp} \left(\frac{\text{Model}}{\text{sigma}} \right) \right) - \text{Log}(\text{sigma}) \\ \text{else} \quad \Rightarrow - \text{Log} \left(1 + \text{Exp} \left(\frac{\text{Model}}{\text{sigma}} \right) \right) \end{array} \right)$$

Chapter 15

Fit Proportional Hazards

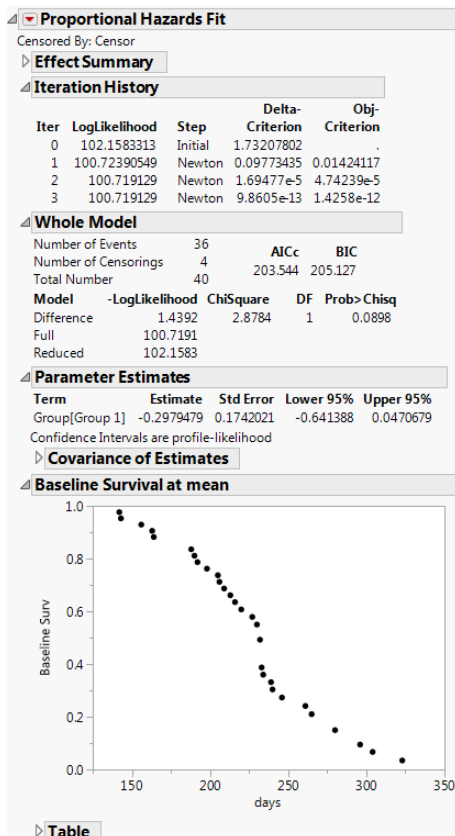
Fit Survival Data Using Semiparametric Regression Models

The Fit Proportional Hazards platform fits the Cox proportional hazards model, which assumes a multiplying relationship between covariates (predictors) and the hazard function.

Proportional hazards models are popular regression models for survival data with covariates. This model is semiparametric. The linear model is estimated, but the form of the hazard function is not. Time-varying covariates are not supported.

Note: The Fit Proportional Hazards platform is a slightly customized version of the Fit Model platform.

Figure 15.1 Example of a Proportional Hazards Fit



Contents

Overview of the Fit Proportional Hazards Platforms 401

Example of the Fit Proportional Hazards Platform 401

 Risk Ratios for One Nominal Effect with Two Levels 403

Launch the Fit Proportional Hazards Platform 404

The Fit Proportional Hazards Report. 406

Fit Proportional Hazards Options 407

Example Using Multiple Effects and Multiple Levels 408

 Risk Ratios for Multiple Effects and Multiple Levels 410

Overview of the Fit Proportional Hazards Platforms

The proportional hazards model is a special semiparametric regression model proposed by D. R. Cox (1972) to examine the effect of explanatory variables on survival times. The survival time of each member of a population is assumed to follow its own hazard function.

The proportional hazards model is nonparametric in the sense that it involves an unspecified arbitrary baseline hazard function. It is parametric because it assumes a parametric form for the covariates. The baseline hazard function is scaled by a function of the model's (time-independent) covariates to give a general hazard function. Unlike the Kaplan-Meier analysis, proportional hazards computes parameter estimates and standard errors for each covariate. The regression parameters (β) associated with the explanatory variables and their standard errors are estimated using the maximum likelihood method. A conditional risk ratio (or hazard ratio) is also computed from the parameter estimates.

The survival estimates in proportional hazards are generated using an empirical method. See Lawless (1982). They represent the empirical cumulative hazard function estimates, $H(t)$, of the survivor function, $S(t)$, and can be written as $S_0 = \exp(-H(t))$. The hazard function is defined as follows:

$$H(t) = \sum_{j:t_j < t} \frac{d_j}{\sum_{l \in R_j} e^{x_l \beta}}$$

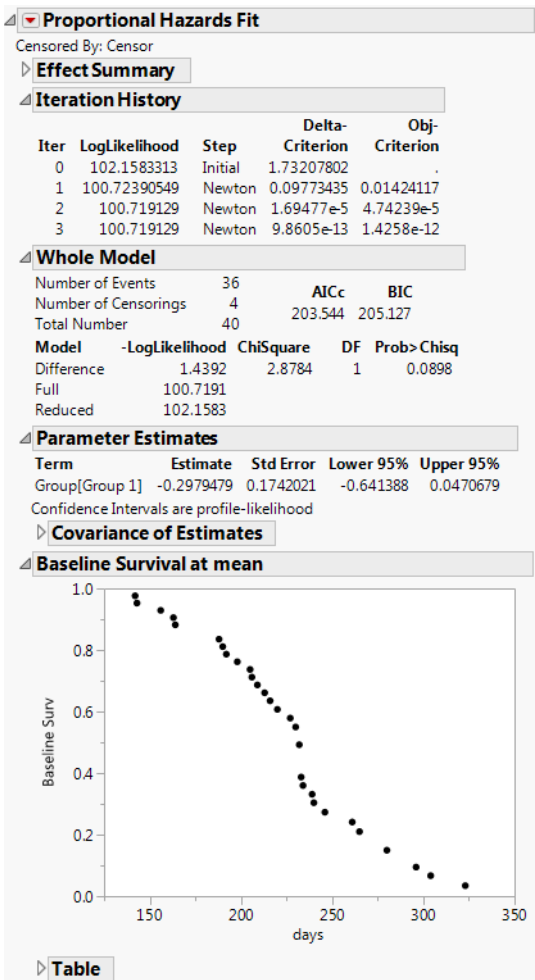
When there are ties in the response, meaning there is more than one failure at a given time event, the Breslow likelihood is used.

Example of the Fit Proportional Hazards Platform

This example illustrates one nominal effect with two levels. For an example with multiple effects and multiple levels, see [“Example Using Multiple Effects and Multiple Levels”](#) on page 408.

1. Select **Help > Sample Data Library** and open Rats.jmp.
2. Select **Analyze > Reliability and Survival > Fit Proportional Hazards**.
3. Select days and click **Time to Event**.
4. Select Censor and click **Censor**.
5. Select Group and click **Add**.
6. Click **Run**.

Figure 15.2 Proportional Hazards Fit Report for Rats.jmp Data



In the Rats.jmp data, there are only two groups. Therefore, in the Parameter Estimates report, a confidence interval that does not include zero indicates an alpha-level significant difference between groups. Also, in the Effect Likelihood Ratio Tests report, the test of the null hypothesis for no difference between the groups shown in the Whole Model Test table is the same as the null hypothesis that the regression coefficient for Group is zero.

Risk Ratios for One Nominal Effect with Two Levels

To show risk ratios for effects, select the **Risk Ratios** option from the red triangle menu. In this example, there is only one effect, and there are only two levels for that effect. The risk ratio for Group 2 is compared with Group 1 and appears in the Risk Ratios for Group report. See Figure 15.3. The risk ratio in this table is determined by computing the exponential of the parameter estimate for Group 2 and dividing it by the exponential of the parameter estimate for Group 1.

Note the following:

- The Group 1 parameter estimate appears in the Parameter Estimates table (Figure 15.2).
- The Group 2 parameter estimate is calculated by taking the negative value for the parameter estimate of Group 1.
- Reciprocal shows the value for 1/Risk Ratio.

Tip: To see the Reciprocal values, right-click in the Risk Ratios report and select **Columns > Reciprocal**.

For this example, the risk ratio for Group2/Group1 is calculated as follows:

$$\exp[-(-0.2979479)] / \exp(-0.2979479) = 1.8146558$$

This risk ratio value suggests that the risk of death for Group 2 is 1.81 times that for Group 1.

Figure 15.3 Risk Ratios for Group Table

Risk Ratios					
Risk Ratios for Group					
Level1	/Level2	Risk Ratio	Prob>Chisq	Lower 95%	Upper 95%
Group 2	Group 1	1.8146558	0.0872	0.9167104	3.5921661
Group 1	Group 2	0.5510687	0.0872	0.2783836	1.0908571

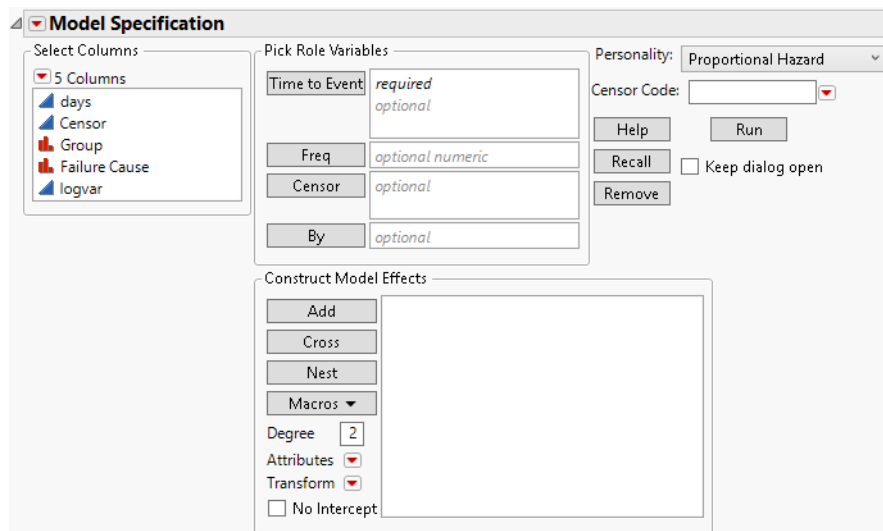
Normal approximations used for ratio confidence limits effects:
Group

For information about calculating risk ratios when there are multiple effects, or categorical effects with more than two levels, see [“Risk Ratios for Multiple Effects and Multiple Levels”](#) on page 410.

Launch the Fit Proportional Hazards Platform

Launch the Fit Proportional Hazards platform by selecting **Analyze > Reliability and Survival > Fit Proportional Hazards**.

Figure 15.4 The Fit Proportional Hazards Launch Window



Tip: To change the alpha level, click the Model Specification red triangle and select **Set Alpha Level**.

For more information about the options in the Select Columns red triangle menu, see *Using JMP*.

The Fit Proportional Hazards launch window contains the following options:

Time to Event Contains the time to event or time to censoring.

Censor Nonzero values in the censor column indicate censored observations. (The coding of censored values is usually equal to 1.) Uncensored values must be coded as 0.

Caution: The Censor column role does not honor the Missing Value Code column property.

Freq Column whose values are the frequencies or counts of observations for each row when there are multiple units recorded.

By Performs a separate analysis for each level of a classification or grouping variable.

Construct Model Effects Enters effects into your model. For more information about the Construct Model Effects options, see *Fitting Linear Models*.

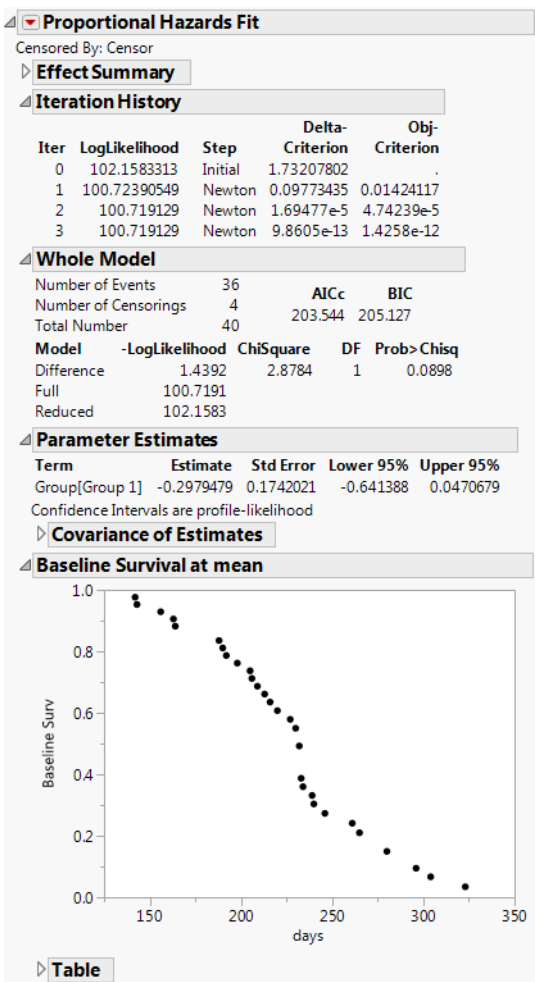
Personality Indicates the fitting method. Proportional Hazard should always be selected.

Censor Code Identifies the value in the Censor column that designates right-censored observations. After a Censor column is selected, JMP attempts to automatically detect the censor code and display it in the box. To change this, you can click the red triangle and select from a list of values. You can also enter a different value in the box. If the Censor column contains a Value Labels column property, the value labels appear in the list of values. Missing values are excluded from the analysis.

The Fit Proportional Hazards Report

Finding parameter estimates for a proportional hazards model is an iterative procedure. When the fitting is complete, the report in Figure 15.5 appears.

Figure 15.5 The Proportional Hazards Fit Report



Iteration History Lists iteration results occurring during the model calculations.

Whole Model Shows the negative of the log-likelihood function ($-\text{LogLikelihood}$) for the model with and without the covariates. Twice the positive difference between them gives a chi-square test of the hypothesis that there is no difference in survival time among the effects. The degrees of freedom (DF) are equal to the change in the number of parameters between the full and reduced models. See *Fitting Linear Models*.

Parameter Estimates Shows the parameter estimates for the covariates, their standard errors, and 95% upper and lower confidence limits. A confidence interval for a continuous column that does not include zero indicates that the effect is significant. A confidence interval for a level in a categorical column that does not include zero indicates that the difference between the level and the average of all levels is significant.

Effect Likelihood Ratio Tests Shows the likelihood ratio chi-square test of the null hypothesis that the parameter estimates for the effects of the covariates is zero.

Baseline Survival at mean Plots the baseline function estimates at each event time in the data. The values in the Table report are plotted here.

Fit Proportional Hazards Options

The Proportional Hazards Fit red triangle menu contains the following options:

Likelihood Ratio Tests Produces tests that compare the log-likelihood from the fitted model to one that removes each term from the model individually.

Wald Tests Produces chi-square test statistics and p -values for Wald tests of whether each parameter is zero.

Likelihood Confidence Intervals Specifies the type of confidence intervals shown in the Parameter Estimates table for each parameter. When this option is selected, a profile likelihood confidence interval appears. Otherwise, a Wald interval is shown. In the report, the interval type is noted below the Parameter Estimates table. This option is on by default when the computational time for the profile likelihood confidence intervals is not large.

Note: You can change the α level for the confidence intervals by selecting Set Alpha Level from the red triangle menu in the Fit Model launch window. The default α level is 0.05.

Risk Ratios Shows the risk ratios for the effects. For continuous columns, unit risk ratios and range risk ratios are calculated. The Unit Risk Ratio is $\text{Exp}(\text{estimate})$ and the Range Risk Ratio is $\text{Exp}[\text{estimate} \cdot (x_{\text{Max}} - x_{\text{Min}})]$. The Unit Risk Ratio shows the risk change over one unit of the regressor, and the Range Risk Ratio shows the change over the whole range of the regressor. For categorical columns, risk ratios are shown in separate reports for each effect. Note that for a categorical variable with k levels, only $k - 1$ design variables, or levels, are used.

Tip: To see Reciprocal values in the Risk Ratio report, right-click in the report and select **Columns > Reciprocal**.

Model Dialog Shows the completed launch window for the current analysis.

Effect Summary Shows the interactive Effect Summary report that enables you to add or remove effects from the model. See *Fitting Linear Models*.

See *Using JMP* for more information about the following options:

Local Data Filter Shows or hides the local data filter that enables you to filter the data used in a specific report.

Redo Contains options that enable you to repeat or relaunch the analysis. In platforms that support the feature, the Automatic Recalc option immediately reflects the changes that you make to the data table in the corresponding report window.

Save Script Contains options that enable you to save a script that reproduces the report to several destinations.

Save By-Group Script Contains options that enable you to save a script that reproduces the platform report for all levels of a By variable to several destinations. Available only when a By variable is specified in the launch window.

Example Using Multiple Effects and Multiple Levels

This example uses a proportional hazards model for the sample data, VA Lung Cancer.jmp. The data were collected from a randomized clinical trial, where males with inoperable lung cancer were placed on either a standard or a novel (test) chemotherapy (Treatment). The primary interest of this trial was to assess if the treatment type has an effect on survival time; special interest is given to the type of tumor (Cell Type). See Prentice (1973) and Kalbfleisch and Prentice (2002) for more information about this data set.

For the proportional hazards model, covariates include the following:

- Whether the patient had undergone previous therapy (Prior)
- Age of the patient (Age)
- Time from lung cancer diagnosis to beginning the study (Diag Time)
- A general medical status measure (KPS)

Age, Diag Time, and KPS are continuous measures and Cell Type, Treatment, and Prior are categorical (nominal) variables. The four nominal levels of Cell Type include Adeno, Large, Small, and Squamous.

This example illustrates the results for a model with more than one effect and a nominal effect with more than two levels. Risk ratios are also demonstrated, with example calculations for risk ratios for a continuous effect and risk ratios for an effect that has more than two levels.

1. Select **Help > Sample Data Library** and open VA Lung Cancer.jmp.
2. Select **Analyze > Reliability and Survival > Fit Proportional Hazards**.

Figure 15.7 Risk Ratios Report

Risk Ratios

Unit Risk Ratios

Per unit change in regressor

Term	Risk Ratio	Lower 95%	Upper 95%	Reciprocal
Age	0.991487	0.97357	1.009733	1.0085861
Diag Time	0.999908	0.982184	1.017952	1.000092
KPS	0.967905	0.957517	0.978405	1.0331596

Range Risk Ratios

Per change in regressor over entire range

Term	Risk Ratio	Lower 95%	Upper 95%	Reciprocal
Age	0.669099	0.283964	1.576584	1.4945466
Diag Time	0.992119	0.213096	4.619042	1.0079435
KPS	0.05484	0.020991	0.143271	18.23482

Risk Ratios for Cell Type

Level1	/Level2 ^	Risk Ratio	Prob>Chisq	Lower 95%	Upper 95%
Large	Adeno	0.4544481	0.0092*	0.2511038	0.8224612
Small	Adeno	0.7176217	0.2286	0.4181316	1.2316242
Squamous	Adeno	0.3047391	<.0001*	0.1690124	0.5494621
Small	Large	1.579106	0.0862	0.9370433	2.6611104
Squamous	Large	0.6705696	0.1574	0.3853373	1.166935
Adeno	Large	2.2004712	0.0092*	1.2158629	3.9824176
Squamous	Small	0.4246514	0.0019*	0.2476225	0.7282408
Adeno	Small	1.3934918	0.2286	0.811936	2.3915917
Large	Small	0.6332697	0.0862	0.375783	1.0671865
Adeno	Squamous	3.2814957	<.0001*	1.8199618	5.916725
Large	Squamous	1.4912695	0.1574	0.8569458	2.5951289
Small	Squamous	2.3548727	0.0019*	1.3731721	4.0384051

Risk Ratios for Treatment

Level1	/Level2 ^	Risk Ratio	Prob>Chisq	Lower 95%	Upper 95%
Test	Standard	1.3363418	0.1617	0.8903074	2.0058347
Standard	Test	0.7483115	0.1617	0.4985456	1.1232076

Risk Ratios for Prior

Level1	/Level2 ^	Risk Ratio	Prob>Chisq	Lower 95%	Upper 95%
Yes	No	1.0750063	0.7554	0.6820551	1.6943478
No	Yes	0.9302271	0.7554	0.5901976	1.4661572

Normal approximations used for ratio confidence limits effects: Cell Type Treatment Prior

Risk Ratios for Multiple Effects and Multiple Levels

Figure 15.7 shows the Risk Ratios for the continuous effects (Age, Diag Time, KPS) and the nominal effects (Cell Type, Treatment, Prior). For illustration, focus on the continuous effect, Age, and the nominal effect with four levels (Cell Type) for the VA Lung Cancer.jmp sample data.

For the continuous effect, Age, in the VA Lung Cancer.jmp sample data, the risk ratios are calculated as follows:

Unit Risk Ratios

$$\exp(\beta) = \exp(-0.0085494) = 0.991487$$

Range Risk Ratios

$$\exp[\beta(x_{\max} - x_{\min})] = \exp(-0.0085494 * 47) = 0.669099$$

For the nominal effect, Cell Type, all pairs of levels are calculated and are shown in the Risk Ratios for Cell Type table. Note that for a categorical variable with k levels, only $k - 1$ design variables, or levels, are used. In the Parameter Estimates table, parameter estimates are shown for only three of the four levels for Cell Type (Adeno, Large, and Small). The Squamous level is not shown, but it is calculated as the negative sum of the other estimates. Here are two example Risk Ratios for Cell Type calculations:

$$\text{Large/Adeno} = \exp(\beta_{\text{Large}})/\exp(\beta_{\text{Adeno}}) = \exp(-0.2114757)/\exp(0.57719588) = 0.4544481$$

$$\begin{aligned} \text{Squamous/Adeno} &= \exp[-(\beta_{\text{Adeno}} + \beta_{\text{Large}} + \beta_{\text{Small}})]/\exp(\beta_{\text{Adeno}}) \\ &= \exp[-(0.57719588 + (-0.2114757) + 0.24538322)]/\exp(0.57719588) = 0.3047391 \end{aligned}$$

Reciprocal shows the value for 1/Risk Ratio.

Appendix **A**

References

- Abernethy, R. B. (1996). *The New Weibull Handbook*. 2nd ed. North Palm Beach, FL: Robert B. Abernethy.
- Akaike, H. (1974). "A New Look at the Statistical Model Identification." *IEEE Transactions on Automatic Control* AC-19:716–723.
- Andrews, D. F., and Herzberg, A. M. (1985). *A Collection of Problems from Many Fields for the Student and Research Worker*. New York: Springer-Verlag.
- Burnham, K. P., and Anderson, D. R. (2002). *Model Selection and Multimodel Inference: A Practical Information-Theoretic Approach*. 2nd ed. New York: Springer-Verlag.
- Chow, S.-C. (2007). *Statistical Design and Analysis of Stability Studies*. Boca Raton, FL: Chapman & Hall/CRC.
- Cox, D. R. (1972). "Regression Models and Life-Tables." *Journal of the Royal Statistical Society, Series B* 34:187–220.
- Crow, L. H. (1975). *Reliability Analysis for Complex, Repairable Systems*. Technical Report No. 138, December 1975, US Army Materiel Systems Analysis Activity, Aberdeen Proving Ground, MD.
- Crow, L. H. (1982). "Confidence Interval Procedures for the Weibull Process with Applications to Reliability Growth." *Technometrics* 24:67–72.
- Escobar, L. A., Meeker, W. Q., Kugler, D. L., and Kramer, L. L. (2003). "Accelerated Destructive Degradation Tests: Data, Models, and Analysis." In *Mathematical and Statistical Methods in Reliability*, edited by B. H. Lindqvist and K. A. Doksum, 319–338. London: World Scientific Publishing Company.
- Genschel, U., and Meeker, W. Q. (2010). "A Comparison of Maximum Likelihood and Median-Rank Regression for Weibull Estimation." *Quality Engineering* 22:236–255.
- Guo, H., Mettas, A., Sarakakis, G., and Niu, P. (2010). "Piecewise NHPP Models with Maximum Likelihood Estimation for Repairable Systems." In *2010 Proceedings - Annual Reliability and Maintainability Symposium (RAMS)*. New York: IEEE Press.
- Hosmer, D. W., Jr., and Lemeshow, S. (1999). *Applied Survival Analysis: Regression Modeling of Time-to-Event Data*. New York: John Wiley & Sons.
- ICH Q1E. (2003). *Evaluation for Stability Data*. Tripartite International Conference on Harmonization Guideline Q1E, Geneva.
https://database.ich.org/sites/default/files/Q1E_Guideline.pdf.
- Kalbfleisch, J. D., and Prentice, R. L. (1980). *The Statistical Analysis of Failure Time Data*. New York: John Wiley & Sons.

- Kalbfleisch, J. D., and Prentice, R. L. (2002). *The Statistical Analysis of Failure Time Data*. 2nd ed. Hoboken, NJ: John Wiley & Sons.
- Kaminskiy, M. P., and Krivtsov, V. V. (2005). "A Simple Procedure for Bayesian Estimation of the Weibull Distribution." *IEEE Transactions on Reliability* 54:612–616.
- Klein, J. P., and Moeschberger, M. L. (1997). *Survival Analysis: Techniques for Censored and Truncated Data*. New York: Springer-Verlag.
- Lawless, J. F. (1982). *Statistical Models and Methods for Lifetime Data*. New York: John Wiley & Sons.
- Lawless, J. F. (2003). *Statistical Models and Methods for Lifetime Data*. 2nd ed. New York: John Wiley & Sons.
- Lee, L., and Lee S. K. (1978). "Some Results on Inference for the Weibull Process." *Technometrics* 20:41–45.
- Liu, P., and Wang, P. (2013). "Competing Failure Modes Modeling with Limited Wearout Failures." In *2013 Proceedings Annual Reliability and Maintainability Symposium (RAMS)*. New York: IEEE Press.
- Meeker, W. Q., and Escobar, L. A. (1998). *Statistical Methods for Reliability Data*. New York: John Wiley & Sons.
- Nair, V. N. (1984). "Confidence Bands for Survival Functions with Censored Data: A Comparative Study." *Technometrics* 26:265–275.
- Nelson, W. B. (1982). *Applied Life Data Analysis*. New York: John Wiley & Sons.
- Nelson, W. B. (1985). "Weibull Analysis of Reliability Data with Few or No Failures." *Journal of Quality Technology* 17:140–146.
- Nelson, W. B. (1990). *Accelerated Testing: Statistical Models, Test Plans, and Data Analyses*. New York: John Wiley & Sons.
- Nelson, W. B. (2003). *Recurrent Events Data Analysis for Product Repairs, Disease Recurrences, and Other Applications*. Philadelphia: Society for Industrial Mathematics.
- Nelson, W. B. (2004). *Accelerated Testing: Statistical Models, Test Plans, and Data Analysis*. New York: John Wiley & Sons.
- Prentice, R. L. (1973). "Exponential Survivals with Censoring and Explanatory Variables." *Biometrika* 60:279–288.
- Rigdon, S. E., and Basu, A. P. (2000). *Statistical Methods for the Reliability of Repairable Systems*. New York: John Wiley & Sons.
- Robert, C. P., and Casella, G. (2004). *Monte Carlo Statistical Methods*. 2nd ed. New York: Springer-Verlag.
- SAS Institute Inc. (2020). "The LIFETEST Procedure." In *SAS/STAT 15.2 User's Guide*. Cary, NC: SAS Institute Inc.
<https://go.documentation.sas.com/api/docsets/statug/15.2/content/lifetest.pdf>.
- Si, S., Dui, H., Zhao, X., Zhang, S., and Sun, S. (2012). "Integrated Importance Measure of Component States Based on Loss of System Performance." *IEEE Transactions on Reliability*. 61:192–202.

- Tobias, P. A., and Trindade, D. C. (1995). *Applied Reliability*. 2nd ed. New York: Van Nostrand Reinhold.
- Tobias, P. A., and Trindade, D. C. (2012). *Applied Reliability*. 3rd ed. Boca Raton, FL: Chapman & Hall/CRC.
- US Department of Defense (1981). *Military Handbook: Reliability Growth Management* (MIL-HDBK-00189). Washington, DC: US Department of Defense.

Appendix **B**

Technology License Notices

- Scintilla is Copyright © 1998–2017 by Neil Hodgson <neilh@scintilla.org>.

All Rights Reserved.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose and without fee is hereby granted, provided that the above copyright notice appear in all copies and that both that copyright notice and this permission notice appear in supporting documentation.

NEIL HODGSON DISCLAIMS ALL WARRANTIES WITH REGARD TO THIS SOFTWARE, INCLUDING ALL IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS, IN NO EVENT SHALL NEIL HODGSON BE LIABLE FOR ANY SPECIAL, INDIRECT OR CONSEQUENTIAL DAMAGES OR ANY DAMAGES WHATSOEVER RESULTING FROM LOSS OF USE, DATA OR PROFITS, WHETHER IN AN ACTION OF CONTRACT, NEGLIGENCE OR OTHER TORTIOUS ACTION, ARISING OUT OF OR IN CONNECTION WITH THE USE OR PERFORMANCE OF THIS SOFTWARE.

- Progress® Telerik® UI for WPF: Copyright © 2008-2019 Progress Software Corporation. All rights reserved. Usage of the included Progress® Telerik® UI for WPF outside of JMP is not permitted.
- ZLIB Compression Library is Copyright © 1995-2005, Jean-Loup Gailly and Mark Adler.
- Made with Natural Earth. Free vector and raster map data @ naturalearthdata.com.
- Packages is Copyright © 2009–2010, Stéphane Sudre (s.sudre.free.fr). All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.

Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.

Neither the name of the WhiteBox nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS “AS IS” AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT OWNER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES

(INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- iODBC software is Copyright © 1995–2006, OpenLink Software Inc and Ke Jin (www.iodbc.org). All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

- Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
- Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.
- Neither the name of OpenLink Software Inc. nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS “AS IS” AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL OPENLINK OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- This program, “bzip2”, the associated library “libbzip2”, and all documentation, are Copyright © 1996–2019 Julian R Seward. All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

1. Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
2. The origin of this software must not be misrepresented; you must not claim that you wrote the original software. If you use this software in a product, an acknowledgment in the product documentation would be appreciated but is not required.
3. Altered source versions must be plainly marked as such, and must not be misrepresented as being the original software.

4. The name of the author may not be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE AUTHOR "AS IS" AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE AUTHOR BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

Julian Seward, jseward@acm.org

bzip2/libbzip2 version 1.0.8 of 13 July 2019

- R software is Copyright © 1999–2012, R Foundation for Statistical Computing.
- MATLAB software is Copyright © 1984-2012, The MathWorks, Inc. Protected by U.S. and international patents. See www.mathworks.com/patents. MATLAB and Simulink are registered trademarks of The MathWorks, Inc. See www.mathworks.com/trademarks for a list of additional trademarks. Other product or brand names may be trademarks or registered trademarks of their respective holders.
- libopc is Copyright © 2011, Florian Reuter. All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

- Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
- Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and / or other materials provided with the distribution.
- Neither the name of Florian Reuter nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS "AS IS" AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT OWNER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF

USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- libxml2 - Except where otherwise noted in the source code (e.g. the files hash.c, list.c and the trio files, which are covered by a similar license but with different Copyright notices) all the files are:

Copyright © 1998–2003 Daniel Veillard. All Rights Reserved.

Permission is hereby granted, free of charge, to any person obtaining a copy of this software and associated documentation files (the “Software”), to deal in the Software without restriction, including without limitation the rights to use, copy, modify, merge, publish, distribute, sublicense, and/or sell copies of the Software, and to permit persons to whom the Software is furnished to do so, subject to the following conditions:

The above copyright notice and this permission notice shall be included in all copies or substantial portions of the Software.

THE SOFTWARE IS PROVIDED “AS IS”, WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO THE WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT. IN NO EVENT SHALL DANIEL VEILLARD BE LIABLE FOR ANY CLAIM, DAMAGES OR OTHER LIABILITY, WHETHER IN AN ACTION OF CONTRACT, TORT OR OTHERWISE, ARISING FROM, OUT OF OR IN CONNECTION WITH THE SOFTWARE OR THE USE OR OTHER DEALINGS IN THE SOFTWARE.

Except as contained in this notice, the name of Daniel Veillard shall not be used in advertising or otherwise to promote the sale, use or other dealings in this Software without prior written authorization from him.

- Regarding the decompression algorithm used for UNIX files:

Copyright © 1985, 1986, 1992, 1993

The Regents of the University of California. All rights reserved.

THIS SOFTWARE IS PROVIDED BY THE REGENTS AND CONTRIBUTORS “AS IS” AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE REGENTS OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

1. Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.

2. Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.

3. Neither the name of the University nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

- Snowball is Copyright © 2001, Dr Martin Porter, Copyright © 2002, Richard Boulton. All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

1. Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.

2. Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.

3. Neither the name of the copyright holder nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS "AS IS" AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT HOLDER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- Pako is Copyright © 2014–2017 by Vitaly Puzrin and Andrei Tuputcyn.

Permission is hereby granted, free of charge, to any person obtaining a copy of this software and associated documentation files (the "Software"), to deal in the Software without restriction, including without limitation the rights to use, copy, modify, merge, publish, distribute, sublicense, and/or sell copies of the Software, and to permit persons to whom the Software is furnished to do so, subject to the following conditions:

The above copyright notice and this permission notice shall be included in all copies or substantial portions of the Software.

THE SOFTWARE IS PROVIDED “AS IS”, WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO THE WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT. IN NO EVENT SHALL THE AUTHORS OR COPYRIGHT HOLDERS BE LIABLE FOR ANY CLAIM, DAMAGES OR OTHER LIABILITY, WHETHER IN AN ACTION OF CONTRACT, TORT OR OTHERWISE, ARISING FROM, OUT OF OR IN CONNECTION WITH THE SOFTWARE OR THE USE OR OTHER DEALINGS IN THE SOFTWARE.

- HDF5 (Hierarchical Data Format 5) Software Library and Utilities are Copyright 2006–2015 by The HDF Group. NCSA HDF5 (Hierarchical Data Format 5) Software Library and Utilities Copyright 1998-2006 by the Board of Trustees of the University of Illinois. All rights reserved. DISCLAIMER: THIS SOFTWARE IS PROVIDED BY THE HDF GROUP AND THE CONTRIBUTORS “AS IS” WITH NO WARRANTY OF ANY KIND, EITHER EXPRESSED OR IMPLIED. In no event shall The HDF Group or the Contributors be liable for any damages suffered by the users arising out of the use of this software, even if advised of the possibility of such damage.
- agl-aglfn technology is Copyright © 2002, 2010, 2015 by Adobe Systems Incorporated. All Rights Reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

- Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
- Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.
- Neither the name of Adobe Systems Incorporated nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS “AS IS” AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT HOLDER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- dmlc/xgboost is Copyright © 2019 SAS Institute.

Licensed under the Apache License, Version 2.0 (the “License”); you may not use this file except in compliance with the License. You may obtain a copy of the License at

<http://www.apache.org/licenses/LICENSE-2.0>

Unless required by applicable law or agreed to in writing, software distributed under the License is distributed on an “AS IS” BASIS, WITHOUT WARRANTIES OR CONDITIONS OF ANY KIND, either express or implied. See the License for the specific language governing permissions and limitations under the License.

- libzip is Copyright © 1999–2019 Dieter Baron and Thomas Klausner.

This file is part of libzip, a library to manipulate ZIP archives. The authors can be contacted at <libzip@nih.at>.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

1. Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
2. Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.
3. The names of the authors may not be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE AUTHORS “AS IS” AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE AUTHORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

- OpenNLP 1.5.3, the pre-trained model (version 1.5 of en-parser-chunking.bin), and dmlc/xgboost Version .90 are licensed under the Apache License 2.0 are Copyright © January 2004 by Apache.org.

You may reproduce and distribute copies of the Work or Derivative Works thereof in any medium, with or without modifications, and in Source or Object form, provided that You meet the following conditions:

- You must give any other recipients of the Work or Derivative Works a copy of this License; and
 - You must cause any modified files to carry prominent notices stating that You changed the files; and
 - You must retain, in the Source form of any Derivative Works that You distribute, all copyright, patent, trademark, and attribution notices from the Source form of the Work, excluding those notices that do not pertain to any part of the Derivative Works; and
 - If the Work includes a “NOTICE” text file as part of its distribution, then any Derivative Works that You distribute must include a readable copy of the attribution notices contained within such NOTICE file, excluding those notices that do not pertain to any part of the Derivative Works, in at least one of the following places: within a NOTICE text file distributed as part of the Derivative Works; within the Source form or documentation, if provided along with the Derivative Works; or, within a display generated by the Derivative Works, if and wherever such third-party notices normally appear. The contents of the NOTICE file are for informational purposes only and do not modify the License. You may add Your own attribution notices within Derivative Works that You distribute, alongside or as an addendum to the NOTICE text from the Work, provided that such additional attribution notices cannot be construed as modifying the License.
 - You may add Your own copyright statement to Your modifications and may provide additional or different license terms and conditions for use, reproduction, or distribution of Your modifications, or for any such Derivative Works as a whole, provided Your use, reproduction, and distribution of the Work otherwise complies with the conditions stated in this License.
- LLVM is Copyright © 2003–2019 by the University of Illinois at Urbana-Champaign. Licensed under the Apache License, Version 2.0 (the “License”); you may not use this file except in compliance with the License. You may obtain a copy of the License at:
<http://www.apache.org/licenses/LICENSE-2.0>
Unless required by applicable law or agreed to in writing, software distributed under the License is distributed on an “AS IS” BASIS, WITHOUT WARRANTIES OR CONDITIONS OF ANY KIND, either express or implied. See the License for the specific language governing permissions and limitations under the License.
 - clang is Copyright © 2007–2019 by the University of Illinois at Urbana-Champaign. Licensed under the Apache License, Version 2.0 (the “License”); you may not use this file except in compliance with the License. You may obtain a copy of the License at:
<http://www.apache.org/licenses/LICENSE-2.0>
Unless required by applicable law or agreed to in writing, software distributed under the License is distributed on an “AS IS”, WITHOUT WARRANTIES OR CONDITIONS OF

ANY KIND, either express or implied. See the License for the specific language governing permissions and limitations under the License.

- lld is Copyright © 2011–2019 by the University of Illinois at Urbana-Champaign.

Licensed under the Apache License, Version 2.0 (the “License”); you may not use this file except in compliance with the License. You may obtain a copy of the License at:

<http://www.apache.org/licenses/LICENSE-2.0>

Unless required by applicable law or agreed to in writing, software distributed under the License is distributed on an “AS IS”, WITHOUT WARRANTIES OR CONDITIONS OF ANY KIND, either express or implied. See the License for the specific language governing permissions and limitations under the License.

- libcurl is Copyright © 1996–2020, Daniel Stenberg, daniel@haxx.se, and many contributors, see the THANKS file. All rights reserved.

Permission to use, copy, modify, and distribute this software for any purpose with or without fee is hereby granted, provided that the above copyright notice and this permission notice appear in all copies.

THE SOFTWARE IS PROVIDED “AS IS”, WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO THE WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT OF THIRD PARTY RIGHTS. IN NO EVENT SHALL THE AUTHORS OR COPYRIGHT HOLDERS BE LIABLE FOR ANY CLAIM, DAMAGES OR OTHER LIABILITY, WHETHER IN AN ACTION OF CONTRACT, TORT OR OTHERWISE, ARISING FROM, OUT OF OR IN CONNECTION WITH THE SOFTWARE OR THE USE OR OTHER DEALINGS IN THE SOFTWARE.

Except as contained in this notice, the name of a copyright holder shall not be used in advertising or otherwise to promote the sale, use or other dealings in this Software without prior written authorization of the copyright holder.

