

SAS/STAT[®] 14.2 User's Guide

The POWER Procedure

This document is an individual chapter from *SAS/STAT® 14.2 User's Guide*.

The correct bibliographic citation for this manual is as follows: SAS Institute Inc. 2016. *SAS/STAT® 14.2 User's Guide*. Cary, NC: SAS Institute Inc.

SAS/STAT® 14.2 User's Guide

Copyright © 2016, SAS Institute Inc., Cary, NC, USA

All Rights Reserved. Produced in the United States of America.

For a hard-copy book: No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, or otherwise, without the prior written permission of the publisher, SAS Institute Inc.

For a web download or e-book: Your use of this publication shall be governed by the terms established by the vendor at the time you acquire this publication.

The scanning, uploading, and distribution of this book via the Internet or any other means without the permission of the publisher is illegal and punishable by law. Please purchase only authorized electronic editions and do not participate in or encourage electronic piracy of copyrighted materials. Your support of others' rights is appreciated.

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, NC 27513-2414

November 2016

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>.

Chapter 90

The POWER Procedure

Contents

Overview: POWER Procedure	7217
Getting Started: POWER Procedure	7219
Computing Power for a One-Sample t Test	7219
Determining Required Sample Size for a Two-Sample t Test	7221
Syntax: POWER Procedure	7226
PROC POWER Statement	7227
COXREG Statement	7227
CUSTOM Statement	7231
LOGISTIC Statement	7237
MULTREG Statement	7245
ONECORR Statement	7249
ONESAMPLEFREQ Statement	7253
ONESAMPLEMEANS Statement	7261
ONEWAYANOVA Statement	7267
PAIREDFREQ Statement	7272
PAIREDMEANS Statement	7279
PLOT Statement	7288
TWOSAMPLEFREQ Statement	7292
TWOSAMPLEMEANS Statement	7300
TWOSAMPLESURVIVAL Statement	7310
TWOSAMPLEWILCOXON Statement	7321
Details: POWER Procedure	7326
Overview of Power Concepts	7326
Summary of Analyses	7327
Specifying Value Lists in Analysis Statements	7329
Keyword-Lists	7329
Number-Lists	7329
Grouped-Number-Lists	7330
Name-Lists	7331
Grouped-Name-Lists	7331
Sample Size Adjustment Options	7332
Error and Information Output	7333
Displayed Output	7334
ODS Table Names	7335
Computational Resources	7335
Memory	7335

CPU Time	7335
Computational Methods and Formulas	7335
Common Notation	7336
Analyses in the COXREG Statement	7337
Analyses in the CUSTOM Statement	7338
Analyses in the LOGISTIC Statement	7340
Analyses in the MULTREG Statement	7343
Analyses in the ONECORR Statement	7345
Analyses in the ONESAMPLEFREQ Statement	7347
Analyses in the ONESAMPLEMEANS Statement	7365
Analyses in the ONEWAYANOVA Statement	7369
Analyses in the PAIREDFREQ Statement	7370
Analyses in the PAIREDMEANS Statement	7374
Analyses in the TWOSAMPLEFREQ Statement	7378
Analyses in the TWOSAMPLEMEANS Statement	7383
Analyses in the TWOSAMPLESURVIVAL Statement	7389
Analyses in the TWOSAMPLEWILCOXON Statement	7393
ODS Graphics	7395
Examples: POWER Procedure	7396
Example 90.1: One-Way ANOVA	7396
Example 90.2: The Sawtooth Power Function in Proportion Analyses	7401
Example 90.3: Simple AB/BA Crossover Designs	7410
Example 90.4: Noninferiority Test with Lognormal Data	7413
Example 90.5: Multiple Regression and Correlation	7417
Example 90.6: Comparing Two Survival Curves	7421
Example 90.7: Confidence Interval Precision	7423
Example 90.8: Customizing Plots	7426
Assigning Analysis Parameters to Axes	7428
Fine-Tuning a Sample Size Axis	7433
Adding Reference Lines	7437
Linking Plot Features to Analysis Parameters	7440
Choosing Key (Legend) Styles	7445
Modifying Symbol Locations	7449
Example 90.9: Binary Logistic Regression with Independent Predictors	7451
Example 90.10: Wilcoxon-Mann-Whitney Test	7453
Example 90.11: Logistic Regression Using the CUSTOM Statement	7456
References	7463

Overview: POWER Procedure

Power and sample size analysis optimizes the resource usage and design of a study, improving chances of conclusive results with maximum efficiency. The POWER procedure performs prospective power and sample size analyses for a variety of goals, such as the following:

- determining the sample size required to get a significant result with adequate probability (power)
- characterizing the power of a study to detect a meaningful effect
- conducting what-if analyses to assess sensitivity of the power or required sample size to other factors

Here *prospective* indicates that the analysis pertains to planning for a future study. This is in contrast to *retrospective* power analysis for a past study, which is not supported by the procedure.

A variety of statistical analyses are covered:

- t tests, equivalence tests, and confidence intervals for means
- tests, equivalence tests, and confidence intervals for binomial proportions
- multiple regression
- tests of correlation and partial correlation
- one-way analysis of variance
- rank tests for comparing two survival curves
- Cox proportional hazards regression
- logistic regression with binary response
- Wilcoxon-Mann-Whitney (rank-sum) test
- extensions of existing analyses that involve the chi-square, F , t , or normal distribution, or the distribution of the correlation coefficient under multivariate normality

The extensions of existing analyses consist of scalar multipliers and custom values for primary noncentrality and degrees of freedom. Important use cases include the following:

- Wald and likelihood ratio tests in generalized linear models that have nominal, count, or ordinal responses
- sample size inflation that is caused by correlated predictors
- sample size deflation that is caused by correlation between the covariates and the response

Examples of generalized linear models include Poisson regression, logistic regression, and zero-inflated models.

For univariate and multivariate linear models, see Chapter 49, “[The GLMPOWER Procedure](#).”

Input for PROC POWER includes the components considered in study planning:

- design
- statistical model and test
- significance level (alpha)
- surmised effects and variability
- power
- sample size

You designate one of these components by a missing value in the input, in order to identify it as the result parameter. The procedure calculates this result value over one or more scenarios of input values for all other components. Power and sample size are the most common result values, but for some analyses the result can be something else. For example, you can solve for the sample size of a single group for a two-sample t test.

In addition to tabular results, PROC POWER produces graphs. You can produce the most common types of plots easily with default settings and use a variety of options for more customized graphics. For example, you can control the choice of axis variables, axis ranges, number of plotted points, mapping of graphical features (such as color, line style, symbol and panel) to analysis parameters, and legend appearance.

If ODS Graphics is enabled, then PROC POWER uses ODS Graphics to create graphs; otherwise, traditional graphs are produced.

For more information about enabling and disabling ODS Graphics, see the section “[Enabling and Disabling ODS Graphics](#)” on page 607 in Chapter 21, “[Statistical Graphics Using ODS](#).”

For specific information about the statistical graphics and options available with the POWER procedure, see the [PLOT](#) statement and the section “[ODS Graphics](#)” on page 7395.

The POWER procedure is one of several tools available in SAS/STAT software for power and sample size analysis. PROC GLMPOWER supports more complex linear models. SAS[®] Studio provides a user interface and implements many of the analyses supported in the procedures. For more information, see Chapter 49, “[The GLMPOWER Procedure](#),” and *SAS Studio: Task Reference Guide*.

The following sections of this chapter describe how to use PROC POWER and discuss the underlying statistical methodology. The section “[Getting Started: POWER Procedure](#)” on page 7219 introduces PROC POWER with simple examples of power computation for a one-sample t test and sample size determination for a two-sample t test. The section “[Syntax: POWER Procedure](#)” on page 7226 describes the syntax of the procedure. The section “[Details: POWER Procedure](#)” on page 7326 summarizes the methods employed by PROC POWER and provides details on several special topics. The section “[Examples: POWER Procedure](#)” on page 7396 illustrates the use of the POWER procedure with several applications.

For an overview of methodology and SAS tools for power and sample size analysis, see Chapter 18, “[Introduction to Power and Sample Size Analysis](#).” For more discussion and examples, see O’Brien and Casteloe (2007); Casteloe (2000); Casteloe and O’Brien (2001); Muller and Benignus (1992); O’Brien and Muller (1993); Lenth (2001).

Getting Started: POWER Procedure

Computing Power for a One-Sample t Test

Suppose you want to improve the accuracy of a machine used to print logos on sports jerseys. The logo placement has an inherently high variability, but the horizontal alignment of the machine can be adjusted. The operator agrees to pay for a costly adjustment if you can establish a nonzero mean horizontal displacement in either direction with high confidence. You have 150 jerseys at your disposal to measure, and you want to determine your chances of a significant result (power) by using a one-sample t test with a two-sided $\alpha = 0.05$.

You decide that 8 mm is the smallest displacement worth addressing. Hence, you will assume a true mean of 8 in the power computation. Experience indicates that the standard deviation is about 40.

Use the **ONESAMPLEMEANS** statement in the POWER procedure to compute the power. Indicate power as the result parameter by specifying the **POWER=** option with a missing value (.). Specify your conjectures for the mean and standard deviation by using the **MEAN=** and **STDDEV=** options and for the sample size by using the **NTOTAL=** option. The statements required to perform this analysis are as follows:

```
proc power;
  onesamplemeans
    mean      = 8
    ntotal    = 150
    stddev    = 40
    power     = .;
run;
```

Default values for the **TEST=**, **DIST=**, **ALPHA=**, **NULLMEAN=**, and **SIDES=** options specify a two-sided t test for a mean of 0, assuming a normal distribution with a significance level of $\alpha = 0.05$.

Figure 90.1 shows the output.

Figure 90.1 Sample Size Analysis for One-Sample t Test

The POWER Procedure One-Sample t Test for Mean

Fixed Scenario Elements	
Distribution	Normal
Method	Exact
Mean	8
Standard Deviation	40
Total Sample Size	150
Number of Sides	2
Null Mean	0
Alpha	0.05
Computed Power	
Power	
0.682	

The power is about 0.68. In other words, there is about a 2/3 chance that the t test will produce a significant result demonstrating the machine's average off-center displacement. This probability depends on the assumptions for the mean and standard deviation.

Now, suppose you want to account for some of your uncertainty in conjecturing the true mean and standard deviation by evaluating the power for four scenarios, using reasonable low and high values, 5 and 10 for the mean, and 30 and 50 for the standard deviation. Also, you might be able to measure more than 150 jerseys, and you would like to know under what circumstances you could get by with fewer. You want to plot power for sample sizes between 100 and 200 to visualize how sensitive the power is to changes in sample size for these four scenarios of means and standard deviations. The following statements perform this analysis:

```
ods graphics on;

proc power;
  onesamplemeans
    mean      = 5 10
    ntotal    = 150
    stddev    = 30 50
    power     = .;
  plot x=n min=100 max=200;
run;

ods graphics off;
```

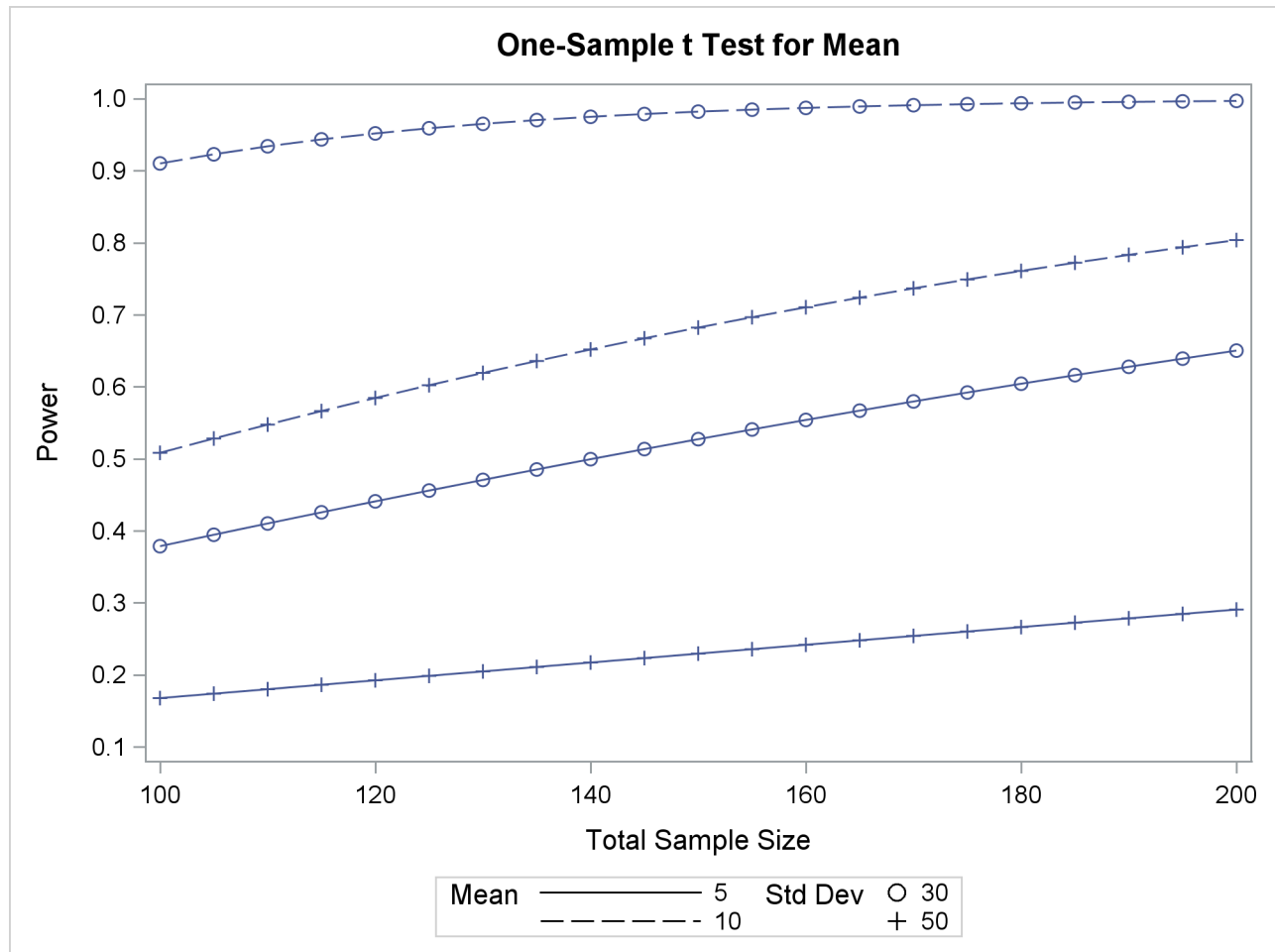
The new mean and standard deviation values are specified by using the **MEAN=** and **STDDEV=** options in the **ONESAMPLEMEANS** statement. The **PLOT** statement with **X=N** produces a plot with sample size on the X axis. (The result parameter, in this case the power, is always plotted on the other axis.) The **MIN=** and **MAX=** options in the **PLOT** statement determine the sample size range. The ODS GRAPHICS ON statement enables ODS Graphics.

Figure 90.2 shows the output, and Figure 90.3 shows the plot.

Figure 90.2 Sample Size Analysis for One-Sample t Test with Input Ranges

The POWER Procedure
One-Sample t Test for Mean

Fixed Scenario Elements				
Distribution		Normal		
Method		Exact		
Total Sample Size		150		
Number of Sides		2		
Null Mean		0		
Alpha		0.05		
Computed Power				
		Std		
Index	Mean	Dev	Power	
1	5	30	0.527	
2	5	50	0.229	
3	10	30	0.982	
4	10	50	0.682	

Figure 90.3 Plot of Power versus Sample Size for One-Sample t Test with Input Ranges

The power ranges from about 0.23 to 0.98 for a sample size of 150 depending on the mean and standard deviation. In Figure 90.3, the line style identifies the mean, and the plotting symbol identifies the standard deviation. The locations of plotting symbols indicate computed powers; the curves are linear interpolations of these points. The plot suggests sufficient power for a mean of 10 and standard deviation of 30 (for any of the sample sizes) but insufficient power for the other three scenarios.

Determining Required Sample Size for a Two-Sample t Test

In this example you want to compare two physical therapy treatments designed to increase muscle flexibility. You need to determine the number of patients required to achieve a power of at least 0.9 to detect a group mean difference in a two-sample t test. You will use $\alpha = 0.05$ (two-tailed).

The mean flexibility with the standard treatment (as measured on a scale of 1 to 20) is well known to be about 13 and is thought to be between 14 and 15 with the new treatment. You conjecture three alternative scenarios for the means:

1. $\mu_1 = 13, \mu_2 = 14$
2. $\mu_1 = 13, \mu_2 = 14.5$
3. $\mu_1 = 13, \mu_2 = 15$

You conjecture two scenarios for the common group standard deviation:

1. $\sigma = 1.2$
2. $\sigma = 1.7$

You also want to try three weighting schemes:

1. equal group sizes (balanced, or 1:1)
2. twice as many patients with the new treatment (1:2)
3. three times as many patients with the new treatment (1:3)

This makes $3 \times 2 \times 3 = 18$ scenarios in all.

Use the **TWOSAMPLEMEANS** statement in the POWER procedure to determine the sample sizes required to give 90% power for each of these 18 scenarios. Indicate total sample size as the result parameter by specifying the **NTOTAL=** option with a missing value (.). Specify your conjectures for the means by using the **GROUPMEANS=** option. Using the “matched” notation (discussed in the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329), enclose the two group means for each scenario in parentheses. Use the **STDDEV=** option to specify scenarios for the common standard deviation. Specify the weighting schemes by using the **GROUPWEIGHTS=** option. You could again use the matched notation. But for illustrative purposes, specify the scenarios for each group weight separately by using the “crossed” notation, with scenarios for each group weight separated by a vertical bar (|). The statements that perform the analysis are as follows:

```
proc power;
  twosamplemeans
    groupmeans    = (13 14) (13 14.5) (13 15)
    stddev        = 1.2 1.7
    groupweights  = 1 | 1 2 3
    power         = 0.9
    ntotal        = .;
run;
```

Default values for the **TEST=**, **DIST=**, **NULLDIFF=**, **ALPHA=**, and **SIDES=** options specify a two-sided t test of group mean difference equal to 0, assuming a normal distribution with a significance level of $\alpha = 0.05$. The results are shown in [Figure 90.4](#).

Figure 90.4 Sample Size Analysis for Two-Sample *t* Test Using Group Means

The POWER Procedure
Two-Sample *t* Test for Mean Difference

Fixed Scenario Elements						
Distribution	Normal					
Method	Exact					
Group 1 Weight	1					
Nominal Power	0.9					
Number of Sides	2					
Null Difference	0					
Alpha	0.05					

Computed N Total						
Index	Mean1	Mean2	Std Dev	Weight2	Actual Power	N Total
1	13	14.0	1.2	1	0.907	64
2	13	14.0	1.2	2	0.908	72
3	13	14.0	1.2	3	0.905	84
4	13	14.0	1.7	1	0.901	124
5	13	14.0	1.7	2	0.905	141
6	13	14.0	1.7	3	0.900	164
7	13	14.5	1.2	1	0.910	30
8	13	14.5	1.2	2	0.906	33
9	13	14.5	1.2	3	0.916	40
10	13	14.5	1.7	1	0.900	56
11	13	14.5	1.7	2	0.901	63
12	13	14.5	1.7	3	0.908	76
13	13	15.0	1.2	1	0.913	18
14	13	15.0	1.2	2	0.927	21
15	13	15.0	1.2	3	0.922	24
16	13	15.0	1.7	1	0.914	34
17	13	15.0	1.7	2	0.921	39
18	13	15.0	1.7	3	0.910	44

The interpretation is that in the best-case scenario (large mean difference of 2, small standard deviation of 1.2, and balanced design), a sample size of $N = 18$ ($n_1 = n_2 = 9$) patients is sufficient to achieve a power of at least 0.9. In the worst-case scenario (small mean difference of 1, large standard deviation of 1.7, and a 1:3 unbalanced design), a sample size of $N = 164$ ($n_1 = 41$, $n_2 = 123$) patients is necessary. The Nominal Power of 0.9 in the “Fixed Scenario Elements” table represents the input target power, and the Actual Power column in the “Computed N Total” table is the power at the sample size (N Total) adjusted to achieve the specified sample weighting exactly.

Note the following characteristics of the analysis, and ways you can modify them if you want:

- The total sample sizes are rounded up to multiples of the weight sums (2 for the 1:1 design, 3 for the 1:2 design, and 4 for the 1:3 design) to ensure that each group size is an integer. To request raw fractional sample size solutions, use the **NFRACTIONAL** option.
- Only the group weight that varies (the one for group 2) is displayed as an output column, while the weight for group 1 appears in the “Fixed Scenario Elements” table. To display the group weights together in output columns, use the matched version of the value list rather than the crossed version.
- If you can specify only differences between group means (instead of their individual values), or if you want to display the mean differences instead of the individual means, use the **MEANDIFF=** option instead of the **GROUPMEANS=** option.

The following statements implement all of these modifications:

```
proc power;
  twosamplemeans
    nfractional
    meandiff      = 1 to 2 by 0.5
    stddev        = 1.2 1.7
    groupweights  = (1 1) (1 2) (1 3)
    power         = 0.9
    ntotal        = .;
run;
```

Figure 90.5 shows the new results.

Figure 90.5 Sample Size Analysis for Two-Sample t Test Using Mean Differences

The POWER Procedure
Two-Sample t Test for Mean Difference

Fixed Scenario Elements	
Distribution	Normal
Method	Exact
Nominal Power	0.9
Number of Sides	2
Null Difference	0
Alpha	0.05

Figure 90.5 *continued*

Computed Ceiling N Total							
Index	Mean Diff	Std Dev	Weight1	Weight2	Fractional N Total	Actual Power	Ceiling N Total
1	1.0	1.2	1	1	62.507429	0.902	63
2	1.0	1.2	1	2	70.065711	0.904	71
3	1.0	1.2	1	3	82.665772	0.901	83
4	1.0	1.7	1	1	123.418482	0.901	124
5	1.0	1.7	1	2	138.598159	0.901	139
6	1.0	1.7	1	3	163.899094	0.900	164
7	1.5	1.2	1	1	28.961958	0.900	29
8	1.5	1.2	1	2	32.308867	0.906	33
9	1.5	1.2	1	3	37.893351	0.901	38
10	1.5	1.7	1	1	55.977156	0.900	56
11	1.5	1.7	1	2	62.717357	0.901	63
12	1.5	1.7	1	3	73.954291	0.900	74
13	2.0	1.2	1	1	17.298518	0.913	18
14	2.0	1.2	1	2	19.163836	0.913	20
15	2.0	1.2	1	3	22.282926	0.910	23
16	2.0	1.7	1	1	32.413512	0.905	33
17	2.0	1.7	1	2	36.195531	0.907	37
18	2.0	1.7	1	3	42.504535	0.903	43

Note that the Nominal Power of 0.9 applies to the raw computed sample size (Fractional N Total), and the Actual Power column applies to the rounded sample size (Ceiling N Total). Some of the adjusted sample sizes in Figure 90.5 are lower than those in Figure 90.4 because underlying group sample sizes are allowed to be fractional (for example, the first Ceiling N Total of 63 corresponding to equal group sizes of 31.5).

Syntax: POWER Procedure

The following statements are available in the POWER procedure:

```

PROC POWER < options > ;
  COXREG < options > ;
  CUSTOM < options > ;
  LOGISTIC < options > ;
  MULTREG < options > ;
  ONECORR < options > ;
  ONESAMPLEFREQ < options > ;
  ONESAMPLEMEANS < options > ;
  ONEWAYANOVA < options > ;
  PAIREDFREQ < options > ;
  PAIREDMEANS < options > ;
  PLOT < plot-options > < / graph-options > ;
  TWOSAMPLEFREQ < options > ;
  TWOSAMPLEMEANS < options > ;
  TWOSAMPLESURVIVAL < options > ;
  TWOSAMPLEWILCOXON < options > ;

```

The statements in the POWER procedure consist of the **PROC POWER** statement, a set of *analysis statements* (for requesting specific power and sample size analyses), and the **PLOT** statement (for producing graphs). The **PROC POWER** statement and at least one of the analysis statements are required. The analysis statements are **COXREG**, **CUSTOM**, **LOGISTIC**, **MULTREG**, **ONECORR**, **ONESAMPLEFREQ**, **ONESAMPLEMEANS**, **ONEWAYANOVA**, **PAIREDFREQ**, **PAIREDMEANS**, **TWOSAMPLEFREQ**, **TWOSAMPLEMEANS**, **TWOSAMPLESURVIVAL**, and **TWOSAMPLEWILCOXON**.

You can use multiple analysis statements and multiple **PLOT** statements. Each analysis statement produces a separate sample size analysis. Each **PLOT** statement refers to the previous analysis statement and generates a separate graph (or set of graphs).

The name of an analysis statement describes the framework of the statistical analysis for which sample size calculations are desired. You use options in the analysis statements to identify the result parameter to compute, to specify the statistical test and computational options, and to provide one or more scenarios for the values of relevant analysis parameters.

Table 90.1 summarizes the basic functions of each statement in PROC POWER. The syntax of each statement in Table 90.1 is described in the following pages.

Table 90.1 Statements in the POWER Procedure

Statement	Description
PROC POWER	Invokes the procedure
COXREG CUSTOM	Cox proportional hazards regression extensions of existing analyses that involve the chi-square, F , t , or normal distribution, or the distribution of the correlation coefficient under multivariate normality

Table 90.1 *continued*

Statement	Description
LOGISTIC	Likelihood ratio chi-square test of a single predictor in logistic regression with binary response
MULTREG	Tests of one or more coefficients in multiple linear regression
ONECORR	Fisher's z test and t test of (partial) correlation
ONESAMPLEFREQ	Tests, confidence interval precision, and equivalence tests of a single binomial proportion
ONESAMPLEMEANS	One-sample t test, confidence interval precision, or equivalence test
ONEWAYANOVA	One-way ANOVA including single-degree-of-freedom contrasts
PAIREDFREQ	McNemar's test for paired proportions
PAIREDMEANS	Paired t test, confidence interval precision, or equivalence test
PLOT	Displays plots for previous sample size analysis
TWOSAMPLEFREQ	Chi-square, likelihood ratio, and Fisher's exact tests for two independent proportions
TWOSAMPLEMEANS	Two-sample t test, confidence interval precision, or equivalence test
TWOSAMPLESURVIVAL	Log-rank, Gehan, and Tarone-Ware tests for comparing two survival curves
TWOSAMPLEWILCOXON	Wilcoxon-Mann-Whitney (rank-sum) test for 2 independent groups

See the section “[Summary of Analyses](#)” on page 7327 for a summary of the analyses available and the syntax required for them.

PROC POWER Statement

PROC POWER < options > ;

The **PROC POWER** statement invokes the POWER procedure. You can specify the following option.

PLOTONLY

specifies that only graphical results from the **PLOT** statement should be produced.

COXREG Statement

COXREG < options > ;

The **COXREG** statement performs power and sample size analyses for the score test of a single scalar predictor in Cox proportional hazards regression for survival data, possibly in the presence of one or more covariates that might be correlated with the tested predictor.

Summary of Options

Table 90.2 summarizes the *options* available in the COXREG statement.

Table 90.2 COXREG Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
SIDES=	Specifies the number of sides and the direction of the statistical test
Specify Effects	
RSQUARE=	Specifies the R^2 value from the regression of the predictor of interest on the remaining predictors
HAZARDRATIO=	Specifies the hazard ratio
Specify Variability	
STDDEV=	Specifies the standard deviation of the predictor variable being tested
Specify Sample Size	
EVENTPROB=	Specifies the probability that an uncensored event occurs
EVENTSTOTAL=	Specifies the expected total number of events
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.3 summarizes the valid result parameters for different analyses in the COXREG statement.

Table 90.3 Summary of Result Parameters in the COXREG Statement

Analyses	Solve For	Syntax
TEST=SCORE	Power Sample size	POWER=. NTOTAL=. EVENTSTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

EVENTPROB=number-list

specifies the probability that an uncensored event occurs. You must specify this option when you use the **NTOTAL=** option, and it is ignored when you use the **EVENTSTOTAL=** option. If you are computing power, the input sample size is multiplied by the event probability to determine the number of events. If you are computing sample size, the internally computed number of events is divided by the event probability to compute the sample size. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

EVENTSTOTAL=number-list**EVENTTOTAL=number-list****EETOTAL=number-list**

specifies the expected number of uncensored events, or requests a solution for this parameter by specifying a missing value (**EVENTSTOTAL=.**). The **NFRACTIONAL** option is automatically enabled when you use the **EVENTSTOTAL=** option. You must use either the **EVENTSTOTAL=** option or the **NTOTAL=** option, and you cannot use both. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

HAZARDRATIO=number-list**HR=number-list**

specifies the hazard ratio for a one-unit increase in the predictor of interest x_1 , holding any other predictors constant. The hazard ratio is equal to $\exp(\beta_1)$, where β_1 is the regression coefficient of x_1 . For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. This option is automatically enabled when you use the **EVENTSTOTAL=** option. For information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option, see the section “[Sample Size Adjustment Options](#)” on page 7332.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). You must use either the **NTOTAL=** option or the **EVENTSTOTAL=** option, and you cannot use both. The number of events is used internally in calculations, and the sample size is the ratio of the number of events (**EVENTSTOTAL**) and the event probability **EVENTPROB**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **ALPHA=**
- **EVENTPROB=**
- **RSQUARE=**
- **HAZARDRATIO=**
- **STDDEV=**

- **NTOTAL=**
- **EVENTSTOTAL=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **COXREG** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **COXREG** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RSQUARE=number-list

specifies the R^2 value from the regression of the predictor of interest on the remaining predictors. The sample size is either multiplied (if you are computing power) or divided (if you are computing sample size) by a factor of $(1 - R^2)$. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating a positive correlation between the tested predictor and survival—that is, a hazard ratio less than 1.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating a negative correlation between the tested predictor and survival—that is, a hazard ratio greater than 1.

By default, **SIDES=2**.

STDDEV=number-list

STD=number-list

specifies the standard deviation of the predictor of interest. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=SCORE

specifies the score test in Cox proportional hazards regression. This is the default test option.

Restrictions on Option Combinations

To specify the sample size, choose one of the following parameterizations:

- sample size (by using the **NTOTAL=** option) and event probability (by using the **EVENTPROB=** option)

- number of events (by using the `EVENTSTOTAL=` option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the `TWOSAMPLESURVIVAL` statement.

Score Test for Cox Proportional Hazards Regression

You can use the `NTOTAL=` and `EVENTPROB=` options, as in the following statements. Default values for the `SIDES=`, `ALPHA=`, and `TEST=` options specify a two-sided score test with a significance level of 0.05.

```
proc power;
  coxreg
    hazardratio = 1.4
    rsquare = 0.15
    stddev = 1.2
    ntotal = 80
    eventprob = 0.8
    power = .
;
run;
```

You can also use the `EVENTSTOTAL=` option, as in the following statements:

```
proc power;
  coxreg
    hazardratio = 1.6
    rsquare = 0.2
    stddev = 1.1
    power = 0.9
    eventstotal = .
;
run;
```

CUSTOM Statement

CUSTOM < options > ;

The `CUSTOM` statement performs power and sample size analyses for extensions of existing analyses that involve the chi-square, F , t , or normal distribution, or the distribution of the correlation coefficient under multivariate normality. The extensions consist of scalar multipliers and custom values for primary noncentrality and degrees of freedom. Important use cases include the following:

- Wald and likelihood ratio tests in generalized linear models that have nominal, count, or ordinal responses
- sample size inflation that is caused by correlated predictors
- sample size deflation that is caused by correlation between the covariates and the response

Examples of generalized linear models include Poisson regression, logistic regression, and zero-inflated models. See [Example 90.11](#) for an example that involves logistic regression.

Summary of Options

Table 90.4 summarizes the *options* available in the CUSTOM statement.

Table 90.4 CUSTOM Statement Options

Option	Description
Define Analysis	
DIST=	Specifies the underlying distributional assumption
Specify Analysis Information	
ALPHA=	Specifies the significance level
MODELDF=	Specifies the model degrees of freedom
SIDES=	Specifies the number of sides and the direction of the statistical test
TESTDF=	Specifies the test degrees of freedom
Specify Effects	
CORR=	Specifies the correlation
CRITMULT=	Specifies the scalar multiplier of the critical value
PNCMULT=	Specifies the scalar multiplier of the primary noncentrality
PRIMNC=	Specifies the primary noncentrality
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the total sample size
Specify Power	
POWER=	Specifies the desired power
Control Ordering in Output	
OUTPUTORDER=	Controls the order of parameters

Table 90.5 shows which *options* apply to each distribution. Options not shown in the table apply to all distributions.

Table 90.5 Options That Can Be Used for Each Distribution in the CUSTOM Statement

	DIST= Value				
Option	CHISQUARE	CORR	F	NORMAL	T
CORR		•			
CRITMULT	•			•	
MODELDF		•	•		•
PNCMULT	•		•	•	•
PRIMNC	•		•	•	•
SIDES		•		•	•
TESTDF	•	•	•		

Table 90.6 summarizes the valid result parameters for different analyses in the **CUSTOM** statement.

Table 90.6 Summary of Result Parameters in the CUSTOM Statement

Solve For	Syntax
Power	POWER=.
Sample size	NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. By default, **ALPHA=0.05**, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CORR=*number-list*

specifies the partial correlation between $p_1 \geq 1$ variables in a set X_1 and a variable Y , adjusting for $p - p_1$ variables in a (possibly empty) set X_{-1} . This option is required for **DIST=****CORR** and not allowed for any other value of the **DIST=** option. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CRITMULT=*number-list*

specifies a scalar multiplier of the critical value in the power computation formula, where values in *number-list* must be nonnegative. This option is allowed only for **DIST=****CHISQUARE** and **DIST=****NORMAL**. Values must be nonnegative. By default, **CRITMULT=1**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DIST=**CHISQUARE** | **CORR** | **F** | **NORMAL** | **T**

specifies the underlying distribution assumed for the test statistic. You must specify one of the following values:

CHISQUARE	specifies the chi-square distribution.
CORR	specifies the distribution of the Pearson correlation coefficient under multivariate normality.
F	specifies the F distribution.
NORMAL	specifies the normal distribution.
T	specifies the t distribution.

MODELDF=*number-list*

specifies the model degrees of freedom. This option is required for **DIST=****CORR**, **DIST=****F**, and **DIST=****T**, and it is not allowed for any other value of the **DIST=** option. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. For information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option, see the section “[Sample Size Adjustment Options](#)” on page 7332.

NTOTAL=number-list

specifies the sample size, or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. You can specify the following values:

INTERNAL arranges the parameters in the output according to the following order of their corresponding *options*:

1. **SIDES=**
2. **TESTDF=**
3. **MODELDF=**
4. **PNCMULT=**
5. **CRITMULT=**
6. **ALPHA=**
7. **PRIMNC=**
8. **CORR=**
9. **NTOTAL=**
10. **POWER=**

SYNTAX arranges the parameters in the output in the same order in which their corresponding *options* are specified in the **CUSTOM** statement.

REVERSE arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **CUSTOM** statement.

By default, **OUTPUTORDER=INTERNAL**.

PNCMULT=number-list

specifies a scalar multiplier of the primary noncentrality in the power computation formula, where values in *number-list* must be nonnegative. This option is allowed only for **DIST=CHISQUARE**, **DIST=F**, **DIST=NORMAL**, and **DIST=T**. By default, **PNCMULT=1**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

POWER=number-list

specifies the desired power of the test, or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability (a number between 0 and 1) rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PRIMNC=*number-list*

specifies the primary noncentrality of the distribution of the test statistic. This option is allowed only for **DIST=CHISQUARE**, **DIST=F**, **DIST=NORMAL**, and **DIST=T**. Values in *number-list* must be nonnegative for **DIST=CHISQUARE** and **DIST=F**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (tails) and the direction of the statistical test. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify one or more of the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

By default, **SIDES=2**.

TESTDF=*number-list*

specifies the test degrees of freedom. This option is allowed only for **DIST=CHISQUARE**, **DIST=CORR**, and **DIST=F**. **TESTDF** must be equal to 1 if **DIST=CORR**; **SIDES=1**, **U**, or **L**; and the **TESTDF=** option is specified. By default, **TESTDF=1**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **CUSTOM** statement.

Chi-Square Distribution

If you assume that the test statistic has a chi-square distribution, then you must specify the **PRIMNC=** option (primary noncentrality) and **TESTDF=** option (test degrees of freedom), as illustrated in the following statements:

```
proc power;
  custom
    dist    = chisquare
    primnc  = 0.23
    testdf  = 1
    ntotal  = 30
    power   = .;
run;
```

The **POWER=.** option requests a solution for the power at a sample size of 30, as specified in the **NTOTAL=** option. You can also specify the **CRITMULT=** option to multiply the critical value in the power formula by a scalar or the **PNCMULT=** option to multiply the primary noncentrality by a scalar.

F Distribution

If you assume that the test statistic has an *F* distribution, then you must specify the **PRIMNC=** option (primary noncentrality), **MODELDF=** option (model degrees of freedom), and **TESTDF=** option (test degrees of freedom), as illustrated in the following statements:

```
proc power;
  custom
    dist    = f
    primnc  = 0.04
    modeldf = 5
    testdf  = 2
    ntotal  = .
    power   = 0.9;
run;
```

The **NTOTAL=.** option requests a solution for the sample size that is required to achieve a power of at least 0.9, as specified in the **POWER=** option. You can also specify the **PNCMULT=** option to multiply the primary noncentrality by a scalar.

t Distribution

If you assume that the test statistic has a *t* distribution, then you must specify the **PRIMNC=** option (primary noncentrality) and **MODELDF=** option (model degrees of freedom), as illustrated in the following statements:

```
proc power;
  custom
    dist    = t
    primnc  = 0.2
    modeldf = 2
    ntotal  = 200
    power   = .;
run;
```

The **POWER=.** option requests a solution for the power at a sample size of 200, as specified in the **NTOTAL=** option. The default of **SIDES=2** specifies a two-sided test. You can also specify the **PNCMULT=** option to multiply the primary noncentrality by a scalar.

Normal Distribution

If you assume that the test statistic has a normal distribution, then you must specify the **PRIMNC=** option (primary noncentrality), as illustrated in the following statements:

```
proc power;
  custom
    dist    = normal
    primnc  = 0.2
    ntotal  = .
    power   = 0.9;
run;
```

The **NTOTAL=.** option requests a solution for the sample size that is required to achieve a power of at least 0.9, as specified in the **POWER=** option. The default of **SIDES=2** specifies a two-sided test. You can also specify the **CRITMULT=** option to multiply the critical value in the power formula by a scalar or the **PNCMULT=** option to multiply the primary noncentrality by a scalar.

Distribution of the Correlation Coefficient under Multivariate Normality

If you plan to use the Pearson correlation coefficient in your data analysis and you assume that the data have a multivariate normal distribution, then you must specify the **CORR=** option (partial correlation), **MODELDF=** option (model degrees of freedom), and **TESTDF=** option (test degrees of freedom), as illustrated in the following statements:

```
proc power;
  custom
    dist      = corr
    corr      = 0.3
    modeldf   = 3
    testdf    = 1
    ntotal    = 130
    power     = .;
run;
```

The **POWER=.** option requests a solution for the power at a sample size of 130, as specified in the **NTOTAL=** option. The default of **SIDES=2** specifies a two-sided test.

LOGISTIC Statement

LOGISTIC <options> ;

The **LOGISTIC** statement performs power and sample size analyses for the likelihood ratio chi-square test of a single predictor in binary logistic regression, possibly in the presence of one or more covariates that might be correlated with the tested predictor.

Summary of Options

Table 90.7 summarizes the *options* available in the **LOGISTIC** statement.

Table 90.7 LOGISTIC Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
COVARIATES=	Specifies the distributions of predictor variables
TESTPREDICTOR=	Specifies the distribution of the predictor variable being tested
VARDIST=	Defines a distribution for a predictor variable
Specify Effects	
CORR=	Specifies the multiple correlation between the predictor and the covariates
COVODDSRATIOS=	Specifies the odds ratios for the covariates
COVREGCOEFFS=	Specifies the regression coefficients for the covariates
DEFAULTUNIT=	Specifies the default change in the predictor variables
INTERCEPT=	Specifies the intercept
RESPONSEPROB=	Specifies the response probability

Table 90.7 *continued*

Option	Description
TESTODDSRATIO=	Specifies the odds ratio being tested
TESTREGCOEFF=	Specifies the regression coefficient for the predictor variable
UNITS=	Specifies the changes in the predictor variables
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Specify Computational Method	
DEFAULTNBINS=	Specifies the default number of categories for each predictor variable
NBINS=	Specifies the number of categories for predictor variables
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.8 summarizes the valid result parameters in the LOGISTIC statement.

Table 90.8 Summary of Result Parameters in the LOGISTIC Statement

Analyses	Solve For	Syntax
TEST=LRCHI	Power	POWER=.
	Sample size	NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CORR=*number-list*

specifies the multiple correlation (ρ) between the tested predictor and the covariates. If you also specify the COVARIATES= option, then the sample size is either multiplied (if you are computing power) or divided (if you are computing sample size) by a factor of $(1 - \rho^2)$. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

COVARIATES=*grouped-name-list*

specifies the distributions of any predictor variables in the model but not being tested, using labels specified with the VARDIST= option. The distributions are assumed to be independent of each other and of the tested predictor. If this option is omitted, then the tested predictor specified by the TESTEDPREDICTOR= option is assumed to be the only predictor in the model. For information about

specifying the *grouped-name-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

COVODDSRATIOS=*grouped-number-list*

specifies the odds ratios for the covariates in the full model (including variables in the **TESTPREDICTOR=** and **COVARIATES=** options). The ordering of the values corresponds to the ordering in the **COVARIATES=** option. If the response variable is coded as $Y = 1$ for success and $Y = 0$ for failure, then the odds ratio for each covariate X is the odds of $Y = 1$ when $X = a$ divided by the odds of $Y = 1$ when $X = b$, where a and b are determined from the **DEFAULTUNIT=** and **UNITS=** options. Values must be greater than zero. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

COVREGCOEFFS=*grouped-number-list*

specifies the regression coefficients for the covariates in the full model including the test predictor (as specified by the **TESTPREDICTOR=** option). The ordering of the values corresponds to the ordering in the **COVARIATES=** option. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

DEFAULTNBINS=*number*

specifies the default number of categories (or “bins”) into which the distribution for each predictor variable is divided in internal calculations. Higher values increase computational time and memory requirements but generally lead to more accurate results. However, if the value is too high, then numerical instability can occur. Lower values are less likely to produce “No solution computed” errors. Each test predictor or covariate that is absent from the **NBINS=** option derives its bin number from the **DEFAULTNBINS=** option. The default value is **DEFAULTNBINS=10**.

There are two variable distributions for which the number of bins can be overridden internally:

- For an ordinal distribution, the number of ordinal values is always used as the number of bins.
- For a binomial distribution, if the requested number of bins is larger than $n + 1$, where n is the sample size parameter of the binomial distribution, then exactly $n + 1$ bins are used.

DEFAULTUNIT=*change-spec*

specifies the default change in the predictor variables assumed for odds ratios specified with the **COVODDSRATIOS=** and **TESTODDSRATIO=** options. Each test predictor or covariate that is absent from the **UNITS=** option derives its change value from the **DEFAULTUNIT=** option. The value must be nonzero. The default value is **DEFAULTUNIT=1**. This option can be used only if at least one of the **COVODDSRATIOS=** and **TESTODDSRATIO=** options is used.

Valid specifications for *change-spec* are as follows:

number defines the odds ratio as the ratio of the response variable odds when $X = a$ to the odds when $X = a - \text{number}$ for any constant a .

<+|->SD defines the odds ratio as the ratio of the odds when $X = a$ to the odds when $X = a - \sigma$ (or $X = a + \sigma$, if **SD** is preceded by a minus sign (–)) for any constant a , where σ is the standard deviation of X (as determined from the **VARDIST=** option).

*multiple****SD** defines the odds ratio as the ratio of the odds when $X = a$ to the odds when $X = a - \text{multiple} * \sigma$ for any constant a , where σ is the standard deviation of X (as determined from the **VARDIST=** option).

PERCENTILES($p1$, $p2$) defines the odds ratio as the ratio of the odds when X is equal to its $p2 \times 100$ th percentile to the odds when X is equal to its $p1 \times 100$ th percentile (where the percentiles are determined from the distribution specified in the **VARDIST=** option). Values for $p1$ and $p2$ must be strictly between 0 and 1.

INTERCEPT=*number-list*

specifies the intercept in the full model (including variables in the **TESTPREDICTOR=** and **COVARIATES=** options). For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NBINS=(“*name*” = *number* <... “*name*” = *number* >)

specifies the number of categories (or “bins”) into which the distribution for each predictor variable (identified by its *name* from the **VARDIST=** option) is divided in internal calculations. Higher values increase computational time and memory requirements but generally lead to more accurate results. However, if the value is too high, then numerical instability can occur. Lower values are less likely to produce “No solution computed” errors. Each predictor variable that is absent from the **NBINS=** option derives its bin number from the **DEFAULTNBINS=** option.

There are two variable distributions for which the **NBINS=** value can be overridden internally:

- For an ordinal distribution, the number of ordinal values is always used as the number of bins.
- For a binomial distribution, if the requested number of bins is larger than $n + 1$, where n is the sample size parameter of the binomial distribution, then exactly $n + 1$ bins are used.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “Sample Size Adjustment Options” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NTOTAL=*number-list*

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). Values must be at least one. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **DEFAULTNBINS=**
- **NBINS=**
- **ALPHA=**
- **RESPONSEPROB=**
- **INTERCEPT=**
- **TESTPREDICTOR=**
- **TESTODDSRATIO=**
- **TESTREGCOEFF=**
- **COVARIATES=**

- COVODDSRATIOS=
- COVREGCOEFFS=
- CORR=
- NTOTAL=
- POWER=

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **LOGISTIC** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **LOGISTIC** statement.

POWER=*number-list*

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RESPONSEPROB=*number-list*

specifies the response probability in the full model when all predictor variables (including variables in the **TESTPREDICTOR=** and **COVARIATES=** options) are equal to their means. The log odds of this probability are equal to the intercept in the full model where all predictor are centered at their means. If the response variable is coded as $Y = 1$ for success and $Y = 0$ for failure, then this probability is equal to the mean of Y in the full model when all X s are equal to their means. Values must be strictly between zero and one. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=LRCHI

specifies the likelihood ratio chi-square test of a single model parameter in binary logistic regression. This is the default test option.

TESTODDSRATIO=*number-list*

specifies the odds ratio for the predictor variable being tested in the full model (including variables in the **TESTPREDICTOR=** and **COVARIATES=** options). If the response variable is coded as $Y = 1$ for success and $Y = 0$ for failure, then the odds ratio for the X being tested is the odds of $Y = 1$ when $X = a$ divided by the odds of $Y = 1$ when $X = b$, where a and b are determined from the **DEFAULTUNIT=** and **UNITS=** options. Values must be greater than zero. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TESTPREDICTOR=*name-list*

specifies the distribution of the predictor variable being tested, using labels specified with the **VARDIST=** option. This distribution is assumed to be independent of the distributions of the covariates as defined in the **COVARIATES=** option. For information about specifying the *name-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TESTREGCOEFF=*number-list*

specifies the regression coefficient for the predictor variable being tested in the full model including the covariates specified by the **COVARIATES=** option. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

UNITS=("name" = *change-spec* <... "name" = *change-spec* >)

specifies the changes in the predictor variables assumed for odds ratios specified with the **COVODDSRATIOS=** and **TESTODDSRATIO=** options. Each predictor variable whose *name* (from the **VARDIST=** option) is absent from the UNITS option derives its change value from the **DEFAULTUNIT=** option. This option can be used only if at least one of the **COVODDSRATIOS=** and **TESTODDSRATIO=** options is used.

Valid specifications for *change-spec* are as follows:

number defines the odds ratio as the ratio of the response variable odds when $X = a$ to the odds when $X = a - \textit{number}$ for any constant a .

<+ | ->**SD** defines the odds ratio as the ratio of the odds when $X = a$ to the odds when $X = a - \sigma$ (or $X = a + \sigma$, if **SD** is preceded by a minus sign (-)) for any constant a , where σ is the standard deviation of X (as determined from the **VARDIST=** option).

*multiple****SD** defines the odds ratio as the ratio of the odds when $X = a$ to the odds when $X = a - \textit{multiple} \times \sigma$ for any constant a , where σ is the standard deviation of X (as determined from the **VARDIST=** option).

PERCENTILES($p1, p2$) defines the odds ratio as the ratio of the odds when X is equal to its $p2 \times 100$ th percentile to the odds when X is equal to its $p1 \times 100$ th percentile (where the percentiles are determined from the distribution specified in the **VARDIST=** option). Values for $p1$ and $p2$ must be strictly between 0 and 1.

Each unit value must be nonzero.

VARDIST("label")=*distribution* (*parameters*)

defines a distribution for a predictor variable.

For the **VARDIST=** option,

label identifies the variable distribution in the output and with the **COVARIATES=** and **TESTPREDICTOR=** options.

distribution specifies the distributional form of the variable.

parameters specifies one or more parameters associated with the distribution.

The *distributions* and *parameters* are named and defined in the same way as the distributions and arguments in the CDF SAS function; for more information, see *SAS Language Reference: Dictionary*. Choices for distributional forms and their parameters are as follows:

ORDINAL ((*values*) : (*probabilities*)) is an ordered categorical distribution. The *values* are any numbers separated by spaces. The *probabilities* are numbers between 0 and 1 (inclusive) separated by spaces. Their sum must be exactly 1. The number of *probabilities* must match the number of *values*.

BETA (a, b <, l, r >) is a beta distribution with shape parameters a and b and optional location parameters l and r . The values of a and b must be greater than 0, and l must be less than r . The default values for l and r are 0 and 1, respectively.

BINOMIAL (p, n) is a binomial distribution with probability of success p and number of independent Bernoulli trials n . The value of p must be greater than 0 and less than 1, and n must be an integer greater than 0. If $n = 1$, then the distribution is binary.

EXPONENTIAL (λ) is an exponential distribution with scale λ , which must be greater than 0.

GAMMA (a, λ) is a gamma distribution with shape a and scale λ . The values of a and λ must be greater than 0.

LAPLACE (θ, λ) is a Laplace distribution with location θ and scale λ . The value of λ must be greater than 0.

LOGISTIC (θ, λ) is a logistic distribution with location θ and scale λ . The value of λ must be greater than 0.

LOGNORMAL (θ, λ) is a lognormal distribution with location θ and scale λ . The value of λ must be greater than 0.

NORMAL (θ, λ) is a normal distribution with mean θ and standard deviation λ . The value of λ must be greater than 0.

POISSON (m) is a Poisson distribution with mean m . The value of m must be greater than 0.

UNIFORM (l, r) is a uniform distribution on the interval $[l, r]$, where $l < r$.

Restrictions on Option Combinations

To specify the intercept in the full model, choose one of the following two parameterizations:

- intercept (using the **INTERCEPT=** options)
- Prob($Y = 1$) when all predictors are equal to their means (using the **RESPONSEPROB=** option)

To specify the effect associated with the predictor variable being tested, choose one of the following two parameterizations:

- odds ratio (using the **TESTODDSRATIO=** options)
- regression coefficient (using the **TESTREGCOEFFS=** option)

To describe the effects of the covariates in the full model, choose one of the following two parameterizations:

- odds ratios (using the **COVODDSRATIOS=** options)
- regression coefficients (using the **COVREGCOEFFS=** options)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **LOGISTIC** statement.

Likelihood Ratio Chi-Square Test for One Predictor

You can express effects in terms of response probability and odds ratios, as in the following statements:

```

proc power;
  logistic
    vardist("x1a") = normal(0, 2)
    vardist("x1b") = normal(0, 3)
    vardist("x2") = poisson(7)
    vardist("x3a") = ordinal((-5 0 5) : (.3 .4 .3))
    vardist("x3b") = ordinal((-5 0 5) : (.4 .3 .3))
    testpredictor = "x1a" "x1b"
    covariates = "x2" | "x3a" "x3b"
    responseprob = 0.15
    testoddsratio = 1.75
    covoddsratios = (2.1 1.4)
    ntotal = 100
    power = .;
run;

```

The **VARDIST=** options define the distributions of the predictor variables. The **TESTPREDICTOR=** option specifies two scenarios for the test predictor distribution, Normal(10,2) and Normal(10,3). The **COVARIATES=** option specifies two covariates, the first with a Poisson(7) distribution. The second covariate has an ordinal distribution on the values -5 , 0 , and 5 with two scenarios for the associated probabilities: $(.3, .4, .3)$ and $(.4, .3, .3)$. The response probability in the full model with all variables equal to zero is specified by the **RESPONSEPROB=** option as 0.15 . The odds ratio for a unit decrease in the tested predictor is specified by the **TESTODDSRATIO=** option to be 1.75 . Corresponding odds ratios for the two covariates in the full model are specified by the **COVODDSRATIOS=** option to be 2.1 and 1.4 . The **POWER=.** option requests a solution for the power at a sample size of 100 as specified by the **NTOTAL=** option.

Default values of the **TEST=** and **ALPHA=** options specify a likelihood ratio test of the first predictor with a significance level of 0.05 . The default of **DEFAULTUNIT=1** specifies that all odds ratios are defined in terms of unit changes in predictors. The default of **DEFAULTNBINS=10** specifies that each of the three predictor variables is discretized into a distribution with 10 categories in internal calculations.

You can also express effects in terms of regression coefficients, as in the following statements:

```

proc power;
  logistic
    vardist("x1a") = normal(0, 2)
    vardist("x1b") = normal(0, 3)
    vardist("x2") = poisson(7)
    vardist("x3a") = ordinal((-5 0 5) : (.3 .4 .3))
    vardist("x3b") = ordinal((-5 0 5) : (.4 .3 .3))
    testpredictor = "x1a" "x1b"
    covariates = "x2" | "x3a" "x3b"
    intercept = -6.928162
    testregcoeff = 0.5596158
    covregcoeffs = (0.7419373 0.3364722)
    ntotal = 100
    power = .;
run;

```

The regression coefficients for the tested predictor (**TESTREGCOEFF=0.5596158**) and covariates (**COVREGCOEFFS=(0.7419373 0.3364722)**) are determined by taking the logarithm of the corresponding odds ratios. The intercept in the full model is specified as -6.928162 by the **INTERCEPT=** option. This number is calculated according to the formula at the end of “Analyses in the LOGISTIC Statement” on page 7340,

which expresses the intercept in terms of the response probability, regression coefficients, and predictor means:

$$\text{Intercept} = \log\left(\frac{0.15}{1 - 0.15}\right) - (0.5596158(0) + 0.7419373(7) + 0.3364722(0))$$

MULTREG Statement

MULTREG < options > ;

The **MULTREG** statement performs power and sample size analyses for Type III *F* tests of sets of predictors in multiple linear regression, assuming either fixed or normally distributed predictors.

Summary of Options

Table 90.9 summarizes the *options* available in the **MULTREG** statement.

Table 90.9 MULTREG Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
MODEL=	Specifies the assumed distribution of the predictors
NFULLPREDICTORS=	Specifies the number of predictors in the full model
NOINT	Specifies a no-intercept model
NREDUCEDPREDICTORS=	Specifies the number of predictors in the reduced model
NTESTPREDICTORS=	Specifies the number of predictors being tested
Specify Effects	
PARTIALCORR=	Specifies the partial correlation
RSQUAREDIFF=	Specifies the difference in R^2
RSQUAREFULL=	Specifies the R^2 of the full model
RSQUAREREDUCED=	Specifies the R^2 of the reduced model
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power
Control Ordering in Output	
OUTPUTORDER=	Controls the order of parameters

Table 90.10 summarizes the valid result parameters in the **MULTREG** statement.

Table 90.10 Summary of Result Parameters in the MULTREG Statement

Analyses	Solve For	Syntax
TEST=TYPE3	Power	POWER=.
	Sample size	NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

MODEL=*keyword-list*

specifies the assumed distribution of the tested predictors. **MODEL=FIXED** indicates a fixed predictor distribution. **MODEL=RANDOM** (the default) indicates a joint multivariate normal distribution for the response and tested predictors. You can use the aliases **CONDITIONAL** for **FIXED** and **UNCONDITIONAL** for **RANDOM**. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

FIXED fixed predictors

RANDOM random (multivariate normal) predictors

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NFULLPREDICTORS=*number-list*

NFULLPRED=*number-list*

specifies the number of predictors in the full model, not counting the intercept. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NOINT

specifies a no-intercept model (for both full and reduced models). By default, the intercept is included in the model. If you want to test the intercept, you can specify the **NOINT** option and simply consider the intercept to be one of the predictors being tested. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NREDUCEDPREDICTORS=*number-list*

NREDUCEDPRED=*number-list*

NREDPRED=*number-list*

specifies the number of predictors in the reduced model, not counting the intercept. This is the same as the difference between values of the **NFULLPREDICTORS=** and **NTESTPREDICTORS=** options. Note that supplying a value of 0 is the same as specifying an *F* test of a Pearson correlation. This option cannot be used at the same time as the **NTESTPREDICTORS=** option. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NTESTPREDICTORS=*number-list*

NTESTPRED=*number-list*

specifies the number of predictors being tested. This is the same as the difference between values of the **NFULLPREDICTORS=** and **NREDUCEDPREDICTORS=** options. Note that supplying identical values for the **NTESTPREDICTORS=** and **NFULLPREDICTORS=** options is the same as specifying an *F* test of a Pearson correlation. This option cannot be used at the same time as the **NREDUCEDPREDICTORS=** option. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NTOTAL=*number-list*

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). The minimum acceptable value for the sample size depends on the **MODEL=**, **NOINT**, **NFULLPREDICTORS=**, **NTESTPREDICTORS=**, and **NREDUCEDPREDICTORS=** options. It ranges from $p + 1$ to $p + 3$, where p is the value of the **NFULLPREDICTORS** option. For further information about minimum **NTOTAL** values, see Table 90.36. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **MODEL=**
- **NFULLPREDICTORS=**
- **NTESTPREDICTORS=**
- **NREDUCEDPREDICTORS=**
- **ALPHA=**
- **PARTIALCORR=**
- **RSQUAREFULL=**
- **RSQUAREREDUCED=**
- **RSQUAREDIFF=**
- **NTOTAL=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **MULTREG** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **MULTREG** statement.

PARTIALCORR=*number-list*

PCORR=*number-list*

specifies the partial correlation between the tested predictors and the response, adjusting for any other predictors in the model. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

POWER=*number-list*

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RSQUAREDIFF=*number-list*

RSQDIFF=*number-list*

specifies the difference in R^2 between the full and reduced models. This is equivalent to the proportion of variation explained by the predictors you are testing. It is also equivalent to the squared semipartial correlation of the tested predictors with the response. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RSQUAREFULL=*number-list*

RSQFULL=*number-list*

specifies the R^2 of the full model, where R^2 is the proportion of variation explained by the model. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RSQUAREREDUCED=*number-list*

RSQREDUCED=*number-list*

RSQRED=*number-list*

specifies the R^2 of the reduced model, where R^2 is the proportion of variation explained by the model. If the reduced model is an empty or intercept-only model (in other words, if **NREDUCEDPREDICTORS=0** or **NTESTPREDICTORS=NFULLPREDICTORS**), then **RSQUAREREDUCED=0** is assumed. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=TYPE3

specifies a Type III F test of a set of predictors adjusting for any other predictors in the model. This is the default test option.

Restrictions on Option Combinations

To specify the number of predictors, use any two of these three options:

- the number of predictors in the full model (**NFULLPREDICTORS=**)
- the number of predictors in the reduced model (**NREDUCEDPREDICTORS=**)
- the number of predictors being tested (**NTESTPREDICTORS=**)

To specify the effect, choose one of the following parameterizations:

- partial correlation (by using the **PARTIALCORR=** option)
- R^2 for the full and reduced models (by using any two of **RSQUAREDIFF=**, **RSQUAREFULL=**, and **RSQUAREREDUCED=**)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **MULTREG** statement.

Type III F Test of a Set of Predictors

You can express effects in terms of partial correlation, as in the following statements. Default values of the **TEST=**, **MODEL=**, and **ALPHA=** options specify a Type III F test with a significance level of 0.05, assuming normally distributed predictors.

```
proc power;
  multreg
    model = random
    nfullpredictors = 7
    ntestpredictors = 3
    partialcorr = 0.35
    ntotal = 100
    power = .;
run;
```

You can also express effects in terms of R^2 :

```
proc power;
  multreg
    model = fixed
    nfullpredictors = 7
    ntestpredictors = 3
    rsquarefull = 0.9
    rsquarediff = 0.1
    ntotal = .
    power = 0.9;
run;
```

ONECORR Statement

ONECORR <options> ;

The **ONECORR** statement performs power and sample size analyses for tests of simple and partial Pearson correlation between two variables. Both Fisher's z test and the t test are supported.

Summary of Options

Table 90.11 summarizes the *options* available in the **ONECORR** statement.

Table 90.11 ONECORR Statement Options

Option	Description
Define Analysis	
DIST=	Specifies the underlying distribution assumed for the test statistic
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
MODEL=	Specifies the assumed distribution of the variables
NPARTIALVARS=	Specifies the number of variables adjusted for in the correlation
NULLCORR=	Specifies the null value of the correlation
SIDES=	Specifies the number of sides and the direction of the statistical test
Specify Effects	
CORR=	Specifies the correlation
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.12 summarizes the valid result parameters in the **ONECORR** statement.

Table 90.12 Summary of Result Parameters in the ONECORR Statement

Analyses	Solve For	Syntax
TEST=PEARSON	Power Sample size	POWER=. NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CORR=*number-list*

specifies the correlation between two variables, possibly adjusting for other variables as determined by the **NPARTIALVARS=** option. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

DIST=FISHERZ | T

specifies the underlying distribution assumed for the test statistic. FISHERZ corresponds to Fisher's z normalizing transformation of the correlation coefficient. T corresponds to the t transformation of the correlation coefficient. Note that **DIST=T** is equivalent to analyses in the **MULTREG** statement with **NTESTPREDICTORS=1**. The default value is FISHERZ.

MODEL=keyword-list

specifies the assumed distribution of the first variable when **DIST=T**. The second variable is assumed to have a normal distribution. **MODEL=FIXED** indicates a fixed distribution. **MODEL=RANDOM** (the default) indicates a joint bivariate normal distribution with the second variable. You can use the aliases **CONDITIONAL** for **FIXED** and **UNCONDITIONAL** for **RANDOM**. This option can be used only for **DIST=T**. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

FIXED fixed variables

RANDOM random (bivariate normal) variables

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPARTIALVARS=number-list**NPVARS=number-list**

specifies the number of variables adjusted for in the correlation between the two primary variables. The default value is 0, which corresponds to a simple correlation. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). Values for the sample size must be at least $p + 3$ when **DIST=T** and **MODEL=CONDITIONAL**, and at least $p + 4$ when either **DIST=FISHER** or when **DIST=T** and **MODEL=UNCONDITIONAL**, where p is the value of the **NPARTIALVARS** option. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLCORR=number-list**NULLC=number-list**

specifies the null value of the correlation. The default value is 0. This option can be used only with the **DIST=FISHERZ** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **MODEL=**
- **SIDES=**

- **NULLCORR=**
- **ALPHA=**
- **NPARTIALVARS=**
- **CORR=**
- **NTOTAL=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **ONECORR** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **ONECORR** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating a correlation greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating a correlation less than the null value.

By default, **SIDES=2**.

TEST=PEARSON

specifies a test of the Pearson correlation coefficient between two variables, possibly adjusting for other variables. This is the default test option.

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **ONECORR** statement.

Fisher’s z Test for Pearson Correlation

The following statements demonstrate a power computation for Fisher’s z test for correlation. Default values of **TEST=PEARSON**, **ALPHA=0.05**, **SIDES=2**, and **NPARTIALVARS=0** are assumed.

```
proc power;
  onecorr dist=fisherz
    nullcorr = 0.15
    corr = 0.35
    ntotal = 180
    power = .;
run;
```

t Test for Pearson Correlation

The following statements demonstrate a sample size computation for the *t* test for correlation. Default values of **TEST=PEARSON**, **MODEL=RANDOM**, **ALPHA=0.05**, and **SIDES=2** are assumed.

```
proc power;
  onecorr dist=t
    npartialvars = 4
    corr = 0.45
    ntotal = .
    power = 0.85;
run;
```

ONESAMPLEFREQ Statement

ONESAMPLEFREQ <options> ;

The **ONESAMPLEFREQ** statement performs power and sample size analyses for exact and approximate tests (including equivalence, noninferiority, and superiority) and confidence interval precision for a single binomial proportion.

Summary of Options

Table 90.13 summarizes the *options* available in the **ONESAMPLEFREQ** statement.

Table 90.13 ONESAMPLEFREQ Statement Options

Option	Description
Define Analysis	
CI=	Specifies an analysis of precision of a confidence interval
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
EQUIVBOUNDS=	Specifies the lower and upper equivalence bounds
LOWER=	Specifies the lower equivalence bound
MARGIN=	Specifies the equivalence or noninferiority or superiority margin
NULLPROPORTION=	Specifies the null proportion
SIDES=	Specifies the number of sides and the direction of the statistical test
UPPER=	Specifies the upper equivalence bound
Specify Effect	
HALFWIDTH=	Specifies the desired confidence interval half-width
PROPORTION=	Specifies the binomial proportion
Specify Variance Estimation	
VAREST=	Specifies how the variance is computed
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size

Table 90.13 *continued*

Option	Description
Specify Power and Related Probabilities	
POWER=	Specifies the desired power of the test
PROBWIDTH=	Specifies the probability of obtaining a confidence interval half-width less than or equal to the value specified by HALFWIDTH=
Choose Computational Method	
METHOD=	Specifies the computational method
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.14 summarizes the valid result parameters for different analyses in the **ONESAMPLEFREQ** statement.

Table 90.14 Summary of Result Parameters in the **ONESAMPLEFREQ** Statement

Analyses	Solve For	Syntax
CI=WILSON	Prob(width)	PROBWIDTH=.
CI=AGRESTICOULL	Prob(width)	PROBWIDTH=.
CI=JEFFREYS	Prob(width)	PROBWIDTH=.
CI=EXACT	Prob(width)	PROBWIDTH=.
CI=WALD	Prob(width)	PROBWIDTH=.
CI=WALD_CORRECT	Prob(width)	PROBWIDTH=.
TEST=ADJZ METHOD=EXACT	Power	POWER=.
TEST=ADJZ METHOD=NORMAL	Power	POWER=.
	Sample size	NTOTAL=.
TEST=EQUIV_ADJZ METHOD=EXACT	Power	POWER=.
TEST=EQUIV_ADJZ METHOD=NORMAL	Power	POWER=.
	Sample size	NTOTAL=.
TEST=EQUIV_EXACT	Power	POWER=.
TEST=EQUIV_Z METHOD=EXACT	Power	POWER=.
TEST=EQUIV_Z METHOD=NORMAL	Power	POWER=.
	Sample size	NTOTAL=.
TEST=EXACT	Power	POWER=.
TEST=Z METHOD=EXACT	Power	POWER=.
TEST=Z METHOD=NORMAL	Power	POWER=.
	Sample size	NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. If the **CI=** and **SIDES=1** options are used, then the value must be less than 0.5. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CI

CI=AGRESTICOULL | AC

CI=JEFFREYS

CI=EXACT | CLOPPERPEARSON | CP

CI=WALD

CI=WALD_CORRECT

CI=WILSON | SCORE

specifies an analysis of precision of a confidence interval for the sample binomial proportion.

The value of the **CI=** option specifies the type of confidence interval. The **CI=**AGRESTICOULL option is a generalization of the “Adjusted Wald / add 2 successes and 2 failures” interval of Agresti and Coull (1998) and is presented in Brown, Cai, and DasGupta (2001). It corresponds to the TABLES / BINOMIAL (AGRESTICOULL) option in PROC FREQ. The **CI=**JEFFREYS option specifies the equal-tailed Jeffreys prior Bayesian interval, which corresponds to the TABLES / BINOMIAL (JEFFREYS) option in PROC FREQ. The **CI=**EXACT option specifies the exact Clopper-Pearson confidence interval based on enumeration, which corresponds to the TABLES / BINOMIAL (EXACT) option in PROC FREQ. The **CI=**WALD option specifies the confidence interval based on the Wald test (also commonly called the *z* test or normal-approximation test), which corresponds to the TABLES / BINOMIAL (WALD) option in PROC FREQ. The **CI=**WALD_CORRECT option specifies the confidence interval based on the Wald test with continuity correction, which corresponds to the TABLES / BINOMIAL (CORRECT WALD) option in PROC FREQ. The **CI=**WILSON option (the default) specifies the confidence interval based on the score statistic, which corresponds to the TABLES / BINOMIAL (WILSON) option in PROC FREQ.

Instead of power, the relevant probability for this analysis is the probability of achieving a desired precision. Specifically, it is the probability that the half-width of the confidence interval will be at most the value specified by the **HALFWIDTH=** option.

EQUIVBOUNDS=*grouped-number-list*

specifies the lower and upper equivalence bounds, representing the same information as the combination of the **LOWER=** and **UPPER=** options but grouping them together. The **EQUIVBOUNDS=** option can be used only with equivalence analyses (**TEST=**EQUIV_ADJZ | EQUIV_EXACT | EQUIV_Z). Values must be strictly between 0 and 1. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

HALFWIDTH=*number-list*

specifies the desired confidence interval half-width. The half-width for a two-sided interval is the length of the confidence interval divided by two. This option can be used only with the **CI=** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

LOWER=number-list

specifies the lower equivalence bound for the binomial proportion. The **LOWER=** option can be used only with equivalence analyses (**TEST=EQUIV_ADJZ** | **EQUIV_EXACT** | **EQUIV_Z**). Values must be strictly between 0 and 1. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

MARGIN=number-list

specifies the equivalence or noninferiority or superiority margin, depending on the analysis.

The **MARGIN=** option can be used with one-sided analyses (**SIDES** = 1 | U | L), in which case it specifies the margin added to the null proportion value in the hypothesis test, resulting in a noninferiority or superiority test (depending on the agreement between the effect and hypothesis directions and the sign of the margin). A test with a null proportion p_0 and a margin m is the same as a test with null proportion $p_0 + m$ and no margin.

The **MARGIN=** option can also be used with equivalence analyses (**TEST=EQUIV_ADJZ** | **EQUIV_EXACT** | **EQUIV_Z**) when the **NULLPROPORTION=** option is used, in which case it specifies the lower and upper equivalence bounds as $p_0 - m$ and $p_0 + m$, where p_0 is the value of the **NULLPROPORTION=** option and m is the value of the **MARGIN=** option.

The **MARGIN=** option cannot be used in conjunction with the **SIDES=2** option. (Instead, specify an equivalence analysis by using **TEST=EQUIV_ADJZ** or **TEST=EQUIV_EXACT** or **TEST=EQUIV_Z**). Also, the **MARGIN=** option cannot be used with the **CI=** option.

Values must be strictly between -1 and 1 . In addition, the sum of **NULLPROPORTION** and **MARGIN** must be strictly between 0 and 1 for one-sided analyses, and the derived lower equivalence bound ($2 * \text{NULLPROPORTION} - \text{MARGIN}$) must be strictly between 0 and 1 for equivalence analyses.

For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

METHOD=EXACT | NORMAL

specifies the computational method. **METHOD=EXACT** (the default) computes exact results by using the binomial distribution. **METHOD=NORMAL** computes approximate results by using the normal approximation to the binomial distribution.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. This option is invalid when the **METHOD=EXACT** option is specified. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLPROPORTION=number-list**NULLP=number-list**

specifies the null proportion. A value of 0.5 corresponds to the sign test. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NULLPROPORTION=**
- **ALPHA=**
- **PROPORTION=**
- **NTOTAL=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **ONESAMPLEFREQ** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **ONESAMPLEFREQ** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROBWIDTH=number-list

specifies the desired probability of obtaining a confidence interval half-width less than or equal to the value specified by the **HALFWIDTH=** option. A missing value (**PROBWIDTH=.**) requests a solution for this probability. Values are expressed as probabilities (for example, 0.9) rather than percentages. This option can be used only with the **CI=** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROPORTION=number-list**P=number-list**

specifies the binomial proportion—that is, the expected proportion of successes in the hypothetical binomial trial. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating a proportion greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating a proportion less than the null value.

If the effect size is zero, then **SIDES=1** is not permitted; instead, specify the direction of the test explicitly in this case with either **SIDES=L** or **SIDES=U**. By default, **SIDES=2**.

TEST= ADJZ | EQUIV_ADJZ | EQUIV_EXACT | EQUIV_Z | EXACT | Z

TEST

specifies the statistical analysis. **TEST=ADJZ** specifies a normal-approximate z test with continuity adjustment. **TEST=EQUIV_ADJZ** specifies a normal-approximate two-sided equivalence test based on the z statistic with continuity adjustment and a TOST (two one-sided tests) procedure. **TEST=EQUIV_EXACT** specifies the exact binomial two-sided equivalence test based on a TOST (two one-sided tests) procedure. **TEST=EQUIV_Z** specifies a normal-approximate two-sided equivalence test based on the z statistic without any continuity adjustment, which is the same as the chi-square statistic, and a TOST (two one-sided tests) procedure. **TEST** or **TEST=EXACT** (the default) specifies the exact binomial test. **TEST=Z** specifies a normal-approximate z test without any continuity adjustment, which is the same as the chi-square test when **SIDES=2**.

UPPER=*number-list*

specifies the upper equivalence bound for the binomial proportion. The **UPPER=** option can be used only with equivalence analyses (**TEST=EQUIV_ADJZ** | **EQUIV_EXACT** | **EQUIV_Z**). Values must be strictly between 0 and 1. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

VAREST=*keyword-list*

specifies how the variance is computed in the test statistic for the **TEST=Z**, **TEST=ADJZ**, **TEST=EQUIV_Z**, and **TEST=EQUIV_ADJZ** analyses. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. Valid *keywords* are as follows:

NULL (the default) estimates the variance by using the null proportion(s) (specified by some combination of the **NULLPROPORTION=**, **MARGIN=**, **LOWER=**, and **UPPER=** options). For **TEST=Z** and **TEST=ADJZ**, the null proportion is the value of the **NULLPROPORTION=** option plus the value of the **MARGIN=** option (if it is used). For **TEST=EQUIV_Z** and **TEST=EQUIV_ADJZ**, there are two null proportions, which correspond to the lower and upper equivalence bounds, one for each test in the TOST (two one-sided tests) procedure.

SAMPLE estimates the variance by using the observed sample proportion.

This option is ignored if the analysis is one other than **TEST=Z**, **TEST=ADJZ**, **TEST=EQUIV_Z**, or **TEST=EQUIV_ADJZ**.

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **ONESAMPLEFREQ** statement.

Exact Test of a Binomial Proportion

The following statements demonstrate a power computation for the exact test of a binomial proportion. Defaults for the **SIDES=** and **ALPHA=** options specify a two-sided test with a 0.05 significance level.

```
proc power;
  onesamplefreq test=exact
    nullproportion = 0.2
    proportion = 0.3
    ntotal = 100
```

```

        power = .;
run;

```

z Test

The following statements demonstrate a sample size computation for the z test of a binomial proportion. Defaults for the **SIDES=**, **ALPHA=**, and **VAREST=** options specify a two-sided test with a 0.05 significance level that uses the null variance estimate.

```

proc power;
  onesamplefreq test=z method=normal
    nullproportion = 0.8
    proportion = 0.85
    sides = u
    ntotal = .
    power = .9;
run;

```

z Test with Continuity Adjustment

The following statements demonstrate a sample size computation for the z test of a binomial proportion with a continuity adjustment. Defaults for the **SIDES=**, **ALPHA=**, and **VAREST=** options specify a two-sided test with a 0.05 significance level that uses the null variance estimate.

```

proc power;
  onesamplefreq test=adjz method=normal
    nullproportion = 0.15
    proportion = 0.1
    sides = l
    ntotal = .
    power = .9;
run;

```

Exact Equivalence Test for a Binomial Proportion

You can specify equivalence bounds by using the **EQUIVBOUNDS=** option, as in the following statements:

```

proc power;
  onesamplefreq test=equiv_exact
    proportion = 0.35
    equivbounds = (0.2 0.4)
    ntotal = 500
    power = .;
run;

```

You can also specify the combination of **NULLPROPORTION=** and **MARGIN=** options:

```

proc power;
  onesamplefreq test=equiv_exact
    proportion = 0.35
    nullproportion = 0.3
    margin = 0.1
    ntotal = 500
    power = .;
run;

```

Finally, you can specify the combination of **LOWER=** and **UPPER=** options:

```
proc power;
  onesamplefreq test=equiv_exact
    proportion = 0.35
    lower = 0.2
    upper = 0.4
    ntotal = 500
    power = .;
run;
```

Note that the three preceding analyses are identical.

Exact Noninferiority Test for a Binomial Proportion

A noninferiority test corresponds to an upper one-sided test with a negative-valued margin, as demonstrated in the following statements:

```
proc power;
  onesamplefreq test=exact
    sides = U
    proportion = 0.15
    nullproportion = 0.1
    margin = -0.02
    ntotal = 130
    power = .;
run;
```

Exact Superiority Test for a Binomial Proportion

A superiority test corresponds to an upper one-sided test with a positive-valued margin, as demonstrated in the following statements:

```
proc power;
  onesamplefreq test=exact
    sides = U
    proportion = 0.15
    nullproportion = 0.1
    margin = 0.02
    ntotal = 130
    power = .;
run;
```

Confidence Interval Precision

The following statements performs a confidence interval precision analysis for the Wilson score-based confidence interval for a binomial proportion. The default value of the **ALPHA=** option specifies a confidence level of 0.95.

```
proc power;
  onesamplefreq ci=wilson
    halfwidth = 0.1
    proportion = 0.3
    ntotal = 70
    probwidth = .;
run;
```

Restrictions on Option Combinations

To specify the equivalence bounds for **TEST=EQUIV_ADJZ**, **TEST=EQUIV_EXACT**, and **TEST=EQUIV_Z**, use any of these three option sets:

- lower and upper equivalence bounds, using the **EQUIVBOUNDS=** option
- lower and upper equivalence bounds, using the **LOWER=** and **UPPER=** options
- null proportion (**NULLPROPORTION=**) and margin (**MARGIN=**)

ONESAMPLEMEANS Statement

ONESAMPLEMEANS < options > ;

The **ONESAMPLEMEANS** statement performs power and sample size analyses for *t* tests, equivalence tests, and confidence interval precision involving one sample.

Summary of Options

Table 90.15 summarizes the *options* available in the **ONESAMPLEMEANS** statement.

Table 90.15 ONESAMPLEMEANS Statement Options

Option	Description
Define Analysis	
CI=	Specifies an analysis of precision of the confidence interval for the mean
DIST=	Specifies the underlying distribution assumed for the test statistic
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
LOWER=	Specifies the lower equivalence bound for the mean
NULLMEAN=	Specifies the null mean
SIDES=	Specifies the number of sides and the direction of the statistical test
UPPER=	Specifies the upper equivalence bound for the mean
Specify Effect	
HALFWIDTH=	Specifies the desired confidence interval half-width
MEAN=	Specifies the mean
Specify Variability	
CV=	Specifies the coefficient of variation
STDDEV=	Specifies the standard deviation
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify Power and Related Probabilities	
POWER=	Specifies the desired power of the test

Table 90.15 *continued*

Option	Description
PROBTYPE=	Specifies the type of probability for the PROBWIDTH= option
PROBWIDTH=	Specifies the probability of obtaining a confidence interval half-width less than or equal to the value specified by HALFWIDTH=
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.16 summarizes the valid result parameters for different analyses in the ONESAMPLEMEANS statement.

Table 90.16 Summary of Result Parameters in the ONESAMPLEMEANS Statement

Analyses	Solve For	Syntax
TEST=T DIST=NORMAL	Power	POWER=.
	Sample size	NTOTAL=.
	Alpha	ALPHA=.
	Mean	MEAN=.
	Standard Deviation	STDDEV=.
TEST=T DIST=LOGNORMAL	Power	POWER=.
	Sample size	NTOTAL=.
TEST=EQUIV	Power	POWER=.
	Sample size	NTOTAL=.
CI=T	Prob(width)	PROBWIDTH=.
	Sample size	NTOTAL=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test or requests a solution for alpha by specifying a missing value (ALPHA=.). The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. If the CI= and SIDES=1 options are used, then the value must be less than 0.5. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CI

CI=T

specifies an analysis of precision of the confidence interval for the mean. Instead of power, the relevant probability for this analysis is the probability of achieving a desired precision. Specifically, it is the probability that the half-width of the confidence interval will be at most the value specified by the HALFWIDTH= option. If neither the CI= option nor the TEST= option is used, the default is TEST=T.

CV=number-list

specifies the coefficient of variation, defined as the ratio of the standard deviation to the mean on the original data scale. You can use this option only with **DIST=LOGNORMAL**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DIST=LOGNORMAL | NORMAL

specifies the underlying distribution assumed for the test statistic. **NORMAL** corresponds to the normal distribution, and **LOGNORMAL** corresponds to the lognormal distribution. The default value is **NORMAL**.

HALFWIDTH=number-list

specifies the desired confidence interval half-width. The half-width is defined as the distance between the point estimate and a finite endpoint. This option can be used only with the **CI=T** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

LOWER=number-list

specifies the lower equivalence bound for the mean. This option can be used only with the **TEST=EQUIV** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

MEAN=number-list

specifies the mean, in the original scale, or requests a solution for the mean by specifying a missing value (**MEAN=.**). The mean is arithmetic if **DIST=NORMAL** and geometric if **DIST=LOGNORMAL**. This option can be used only with the **TEST=T** and **TEST=EQUIV** analyses. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLMEAN=number-list**NULLM=number-list**

specifies the null mean, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). The default value is 0 when **DIST=NORMAL** and 1 when **DIST=LOGNORMAL**. This option can be used only with the **TEST=T** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- SIDES=
- NULLMEAN=
- LOWER=
- UPPER=
- ALPHA=
- MEAN=
- HALFWIDTH=
- STDDEV=
- CV=
- NTOTAL=
- POWER=
- PROBTTYPE=
- PROBWIDTH=

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **ONESAMPLEMEANS** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **ONESAMPLEMEANS** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. This option can be used only with the **TEST=T** and **TEST=EQUIV** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

PROBTTYPE=keyword-list

specifies the type of probability for the **PROBWIDTH=** option. A value of **CONDITIONAL** (the default) indicates the conditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option, given that the true mean is captured by the confidence interval. A value of **UNCONDITIONAL** indicates the unconditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option. You can use the alias **GIVENVALIDITY** for **CONDITIONAL**. The **PROBTTYPE=** option can be used only with the **CI=T** analysis. For information about specifying the *keyword-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CONDITIONAL width probability conditional on interval containing the mean

UNCONDITIONAL unconditional width probability

PROBWIDTH=number-list

specifies the desired probability of obtaining a confidence interval half-width less than or equal to the value specified by the **HALFWIDTH=** option. A missing value (**PROBWIDTH=.**) requests a solution for this probability. The type of probability is controlled with the **PROBTTYPE=** option. Values are expressed as probabilities (for example, 0.9) rather than percentages. This option can be used only with the **CI=T** analysis. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. Valid *keywords* and their interpretation for the **TEST=** analyses are as follows:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating a mean greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating a mean less than the null value.

For confidence intervals, **SIDES=U** refers to an interval between the lower confidence limit and infinity, and **SIDES=L** refers to an interval between minus infinity and the upper confidence limit. For both of these cases and **SIDES=1**, the confidence interval computations are equivalent. The **SIDES=** option can be used only with the **TEST=T** and **CI=T** analyses. By default, **SIDES=2**.

STDDEV=*number-list***STD=***number-list*

specifies the standard deviation, or requests a solution for the standard deviation by specifying a missing value (**STDDEV=.**). You can use this option only with **DIST=NORMAL**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=EQUIV | T**TEST**

specifies the statistical analysis. **TEST=EQUIV** specifies an equivalence test of the mean by using a two one-sided tests (TOST) analysis (Schuirmann 1987). **TEST** or **TEST=T** (the default) specifies a *t* test on the mean. If neither the **TEST=** option nor the **CI=** option is used, the default is **TEST=T**.

UPPER=*number-list*

specifies the upper equivalence bound for the mean, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). This option can be used only with the **TEST=EQUIV** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

Restrictions on Option Combinations

To define the analysis, choose one of the following parameterizations:

- a statistical test (by using the **TEST=** option)
- confidence interval precision (by using the **CI=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the `ONESAMPLEMEANS` statement.

One-Sample *t* Test

The following statements demonstrate a power computation for the one-sample *t* test. Default values for the `DIST=`, `SIDES=`, `NULLMEAN=`, and `ALPHA=` options specify a two-sided test for zero mean with a normal distribution and a significance level of 0.05.

```
proc power;
  onesamplemeans test=t
    mean = 7
    stddev = 3
    ntotal = 50
    power = .;
run;
```

One-Sample *t* Test with Lognormal Data

The following statements demonstrate a sample size computation for the one-sample *t* test for lognormal data. Default values for the `SIDES=`, `NULLMEAN=`, and `ALPHA=` options specify a two-sided test for unit mean with a significance level of 0.05.

```
proc power;
  onesamplemeans test=t dist=lognormal
    mean = 7
    cv = 0.8
    ntotal = .
    power = 0.9;
run;
```

Equivalence Test for Mean of Normal Data

The following statements demonstrate a power computation for the TOST equivalence test for a normal mean. Default values for the `DIST=` and `ALPHA=` options specify a normal distribution and a significance level of 0.05.

```
proc power;
  onesamplemeans test=equiv
    lower = 2
    upper = 7
    mean = 4
    stddev = 3
    ntotal = 100
    power = .;
run;
```

Equivalence Test for Mean of Lognormal Data

The following statements demonstrate a sample size computation for the TOST equivalence test for a lognormal mean. The default of `ALPHA=0.05` specifies a significance level of 0.05.

```

proc power;
  onesamplemeans test=equiv dist=lognormal
    lower = 1
    upper = 5
    mean = 3
    cv = 0.6
    ntotal = .
    power = 0.85;
run;

```

Confidence Interval for Mean

By default **CI=T** analyzes the conditional probability of obtaining the desired precision, given that the interval contains the true mean, as in the following statements. The defaults of **SIDES=2** and **ALPHA=0.05** specify a two-sided interval with a confidence level of 0.95.

```

proc power;
  onesamplemeans ci = t
    halfwidth = 14
    stddev = 8
    ntotal = 50
    probwidth = .;
run;

```

ONEWAYANOVA Statement

ONEWAYANOVA <options> ;

The **ONEWAYANOVA** statement performs power and sample size analyses for one-degree-of-freedom contrasts and the overall *F* test in one-way analysis of variance.

Summary of Options

Table 90.17 summarizes the *options* available in the **ONEWAYANOVA** statement.

Table 90.17 ONEWAYANOVA Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
CONTRAST=	Specifies coefficients for single-degree-of-freedom hypothesis tests
NULLCONTRAST=	Specifies the null value of the contrast
SIDES=	Specifies the number of sides and the direction of the statistical test
Specify Effect	
GROUPMEANS=	Specifies the group means

Table 90.17 *continued*

Option	Description
Specify Variability	
STDDEV=	Specifies the error standard deviation
Specify Sample Size and Allocation	
GROUPNS=	Specifies the group sample sizes
GROUPWEIGHTS=	Specifies the sample size allocation weights for the groups
NFRACTIONAL	Enables fractional input and output for sample sizes
NPERGROUP=	Specifies the common sample size per group
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.18 summarizes the valid result parameters for different analyses in the **ONEWAYANOVA** statement.

Table 90.18 Summary of Result Parameters in the **ONEWAYANOVA** Statement

Analyses	Solve For	Syntax
TEST=CONTRAST	Power Sample size	POWER=. NTOTAL=. NPERGROUP==.
TEST=OVERALL	Power Sample size	POWER=. NTOTAL=. NPERGROUP==.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CONTRAST= (*values*) < (... *values*) >

specifies coefficients for single-degree-of-freedom hypothesis tests. You must provide a coefficient for every mean appearing in the **GROUPMEANS=** option. Specify multiple contrasts either with additional sets of coefficients or with additional **CONTRAST=** options. For example, you can specify two different contrasts of five means by using the following:

```
CONTRAST = (1 -1 0 0 0) (1 0 -1 0 0)
```

GROUPMEANS=*grouped-number-list*

GMEANS=*grouped-number-list*

specifies the group means. This option is used to implicitly set the number of groups. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPNS=*grouped-number-list*

GNS=*grouped-number-list*

specifies the group sample sizes. The number of groups represented must be the same as with the **GROUPMEANS**= option. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPWEIGHTS=*grouped-number-list*

GWEIGHTS=*grouped-number-list*

specifies the sample size allocation weights for the groups. This option controls how the total sample size is divided between the groups. Each set of values across all groups represents relative allocation weights. Additionally, if the **NFRACTIONAL** option is not used, the total sample size is restricted to be equal to a multiple of the sum of the group weights (so that the resulting design has an integer sample size for each group while adhering exactly to the group allocation weights). The number of groups represented must be the same as with the **GROUPMEANS**= option. Values must be integers unless the **NFRACTIONAL** option is used. The default value is 1 for each group, amounting to a balanced design. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPERGROUP=*number-list*

NPERG=*number-list*

specifies the common sample size per group or requests a solution for the common sample size per group by specifying a missing value (**NPERGROUP**=.). Use of this option implicitly specifies a balanced design. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NTOTAL=*number-list*

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL**=.). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLCONTRAST=*number-list*

NULLC=*number-list*

specifies the null value of the contrast. The default value is 0. This option can be used only with the **TEST=CONTRAST** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **CONTRAST=**
- **SIDES=**
- **NULLCONTRAST=**
- **ALPHA=**
- **GROUPMEANS=**
- **STDDEV=**
- **GROUPWEIGHTS=**
- **NTOTAL=**
- **NPERGROUP==**
- **GROUPNS=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **ONEWAYANOVA** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **ONEWAYANOVA** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

You can use this option only with the **TEST=CONTRAST** analysis. By default, **SIDES=2**.

STDDEV=number-list**STD=number-list**

specifies the error standard deviation. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=CONTRAST | OVERALL

specifies the statistical analysis. **TEST=CONTRAST** specifies a one-degree-of-freedom test of a contrast of means. The test is the usual F test for the two-sided case and the usual t test for the one-sided case. **TEST=OVERALL** specifies the overall F test of equality of all means. The default is **TEST=CONTRAST** if the **CONTRAST=** option is used, and **TEST=OVERALL** otherwise.

Restrictions on Option Combinations

To specify the sample size and allocation, choose one of the following parameterizations:

- sample size per group in a balanced design (by using the **NPERGROUP=** option)
- total sample size and allocation weights (by using the **NTOTAL=** and **GROUPWEIGHTS=** options)
- individual group sample sizes (by using the **GROUPNS=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **ONEWAYANOVA** statement.

One-Degree-of-Freedom Contrast

You can use the **NPERGROUP=** option in a balanced design, as in the following statements. Default values for the **SIDES=**, **NULLCONTRAST=**, and **ALPHA=** options specify a two-sided test for a contrast value of 0 with a significance level of 0.05.

```
proc power;
  onewayanova test=contrast
    contrast = (1 0 -1)
    groupmeans = 3 | 7 | 8
    stddev = 4
    npergroup = 50
    power = .;
run;
```

You can also specify an unbalanced design with the **NTOTAL=** and **GROUPWEIGHTS=** options:

```
proc power;
  onewayanova test=contrast
    contrast = (1 0 -1)
    groupmeans = 3 | 7 | 8
    stddev = 4
    groupweights = (1 2 2)
    ntotal = .
    power = 0.9;
run;
```

Another way to specify the sample sizes is with the **GROUPNS=** option:

```
proc power;
  onewayanova test=contrast
    contrast = (1 0 -1)
    groupmeans = 3 | 7 | 8
    stddev = 4
    groupns = (20 40 40)
    power = .;
run;
```

Overall F Test

The following statements demonstrate a power computation for the overall F test in a one-way ANOVA. The default of **ALPHA**=0.05 specifies a significance level of 0.05.

```
proc power;
  onewayanova test=overall
    groupmeans = 3 | 7 | 8
    stddev = 4
    npergroup = 50
    power = .;
run;
```

PAIREDFREQ Statement

PAIREDFREQ <options> ;

The **PAIREDFREQ** statement performs power and sample size analyses for McNemar's test for paired proportions.

Summary of Options

Table 90.19 summarizes the *options* available in the **PAIREDFREQ** statement.

Table 90.19 PAIREDFREQ Statement Options

Option	Description
Define Analysis	
DIST=	Specifies the underlying distribution assumed for the test statistic
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
NULLDISCPRORATIO=	Specifies the null value of the ratio of discordant proportions
SIDES=	Specifies the number of sides and the direction of the statistical test or confidence interval
Specify Effects	
CORR=	Specifies the correlation ϕ between members of a pair
DISCPRORATIO=	Specifies the discordant proportion difference $p_{01} - p_{10}$
DISCPRORATIONS=	Specifies the two discordant proportions, p_{10} and p_{01}

Table 90.19 *continued*

Option	Description
DISCPRORATIO=	Specifies the ratio p_{01}/p_{10}
ODDSRATIO=	Specifies the odds ratio $[p_{.1}/(1 - p_{.1})] / [p_{1.}/(1 - p_{1.})]$
PAIREDPROPORTIONS=	Specifies the two paired proportions, $p_{1.}$ and $p_{.1}$
PROPORTIONDIFF=	Specifies the proportion difference $p_{.1} - p_{1.}$
REFPROPORTION=	Specifies either the reference first proportion $p_{1.}$ or the reference discordant proportion p_{10}
RELATIVERISK=	Specifies the relative risk $p_{.1}/p_{1.}$
TOTALPROPDISC=	Specifies the discordant proportion sum, $p_{10} + p_{01}$
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NPAIRS=	Specifies the total number of proportion pairs
Specify Power	
POWER=	Specifies the desired power of the test
Choose Computational Method	
METHOD=	Specifies the computational method
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.20 summarizes the valid result parameters in the [PAIREDFREQ](#) statement.

Table 90.20 Summary of Result Parameters in the PAIREDFREQ Statement

Analyses	Solve For	Syntax
TEST=MCNEMAR METHOD=CONNOR	Power Sample size	POWER=. NPAIRS=.
TEST=MCNEMAR METHOD=EXACT	Power	POWER=.
TEST=MCNEMAR METHOD=MIETTINEN	Power Sample size	POWER=. NPAIRS=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CORR=*number-list*

specifies the correlation ϕ between members of a pair. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DISCPROPORTIONS=*grouped-number-list*

DISCPS=*grouped-number-list*

specifies the two discordant proportions, p_{10} and p_{01} . For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DISCPROPDIFF=*number-list*

DISCPDIFF=*number-list*

specifies the difference $p_{01} - p_{10}$ between discordant proportions. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DISCPRORATIO=*number-list*

DISCRATIO=*number-list*

specifies the ratio p_{01}/p_{10} of discordant proportions. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DIST=EXACT_COND | NORMAL

specifies the underlying distribution assumed for the test statistic. EXACT_COND corresponds to the exact conditional test, based on the exact binomial distribution of the two types of discordant pairs given the total number of discordant pairs. NORMAL corresponds to the conditional test based on the normal approximation to the binomial distribution of the two types of discordant pairs given the total number of discordant pairs. The default value is EXACT_COND.

METHOD=CONNOR | EXACT | MIETTINEN

specifies the computational method. **METHOD=EXACT** (the default) uses the exact binomial distributions of the total number of discordant pairs and the two types of discordant pairs. **METHOD=CONNOR** uses an approximation from Connor (1987), and **METHOD=MIETTINEN** uses an approximation from Miettinen (1968). The CONNOR and MIETTINEN methods are valid only for **DIST=NORMAL**.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option. This option cannot be used with **METHOD=EXACT**.

NPAIRS=*number-list*

specifies the total number of proportion pairs (concordant and discordant) or requests a solution for the number of pairs by specifying a missing value (**NPAIRS=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLDISCPRORATIO=*number-list*

NULLDISCRATIO=*number-list*

NULLRATIO=*number-list*

NULLR=*number-list*

specifies the null value of the ratio of discordant proportions. The default value is 1. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

ODDSRATIO=*number-list*

OR=*number-list*

specifies the odds ratio $[p_{\cdot 1}/(1 - p_{\cdot 1})] / [p_{1 \cdot}/(1 - p_{1 \cdot})]$. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NULLDISCPRORATIO=**
- **ALPHA=**
- **PAIREDPROPORTIONS=**
- **PROPORTIONDIFF=**
- **ODDSRATIO=**
- **RELATIVERISK=**
- **CORR=**
- **DISCPRORATIONS=**
- **DISCPROPDIFF=**
- **TOTALPROPDISC=**
- **REFPROPORTION=**
- **DISCPRORATIO=**
- **NPAIRS=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **PAIREFREQ** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **PAIREFREQ** statement.

PAIREDPROPORTIONS=*grouped-number-list*

PPROPORTIONS=*grouped-number-list*

PAIREDPS=*grouped-number-list*

PPS=*grouped-number-list*

specifies the two paired proportions, $p_{1\cdot}$ and $p_{\cdot 1}$. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

POWER=*number-list*

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROPORTIONDIFF=*number-list*

PDIF=*number-list*

specifies the proportion difference $p_{\cdot 1} - p_{1\cdot}$. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

REFPROPORTION=*number-list*

REFP=*number-list*

specifies either the reference first proportion p_1 . (when used in conjunction with the **PROPORTION-DIFF=**, **ODDSRATIO=**, or **RELATIVERISK=** option) or the reference discordant proportion p_{10} (when used in conjunction with the **DISCPROPDIFF=** or **DISCPROPRATIO=** option). For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

RELATIVERISK=*number-list*

RR=*number-list*

specifies the relative risk $p_{.1}/p_{1.}$. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

If the effect size is zero, then **SIDES=1** is not permitted; instead, specify the direction of the test explicitly in this case with either **SIDES=L** or **SIDES=U**. By default, **SIDES=2**.

TEST=MCNEMAR

specifies the McNemar test of paired proportions. This is the default test option.

TOTALPROPDISC=*number-list*

TOTALPDISC=*number-list*

PDISC=*number-list*

specifies the sum of the two discordant proportions, $p_{10} + p_{01}$. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

Restrictions on Option Combinations

To specify the proportions, choose one of the following parameterizations:

- discordant proportions (using the **DISCPROPORTIONS=** option)
- difference and sum of discordant proportions (using the **DISCPROPDIFF=** and **TOTALPROPDISC=** options)
- difference of discordant proportions and reference discordant proportion (using the **DISCPROPDIFF=** and **REFPROPORTION=** options)

- ratio of discordant proportions and reference discordant proportion (using the `DISCPROPRATIO=` and `REFPROPORTION=` options)
- ratio and sum of discordant proportions (using the `DISCPROPRATIO=` and `TOTAL-PROPDISC=` options)
- paired proportions and correlation (using the `PAIREDPROPORTIONS=` and `CORR=` options)
- proportion difference, reference proportion, and correlation (using the `PROPORTIONDIFF=`, `REFPROPORTION=`, and `CORR=` options)
- odds ratio, reference proportion, and correlation (using the `ODDSRATIO=`, `REFPROPORTION=`, and `CORR=` options)
- relative risk, reference proportion, and correlation (using the `RELATIVERISK=`, `REFPROPORTION=`, and `CORR=` options)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the `PAIREDFREQ` statement.

McNemar Exact Conditional Test

You can express effects in terms of the individual discordant proportions, as in the following statements. Default values for the `TEST=`, `SIDES=`, `ALPHA=`, and `NULLDISCPROPRATIO=` options specify a two-sided McNemar test for no effect with a significance level of 0.05.

```
proc power;
  pairedfreq dist=exact_cond
    discproportions = 0.15 | 0.45
    npairs = 80
    power = .;
run;
```

You can also express effects in terms of the difference and sum of discordant proportions:

```
proc power;
  pairedfreq dist=exact_cond
    discproppdiff = 0.3
    totalproppdisc = 0.6
    npairs = 80
    power = .;
run;
```

You can also express effects in terms of the difference of discordant proportions and the reference discordant proportion:

```
proc power;
  pairedfreq dist=exact_cond
    discproppdiff = 0.3
    refproportion = 0.15
    npairs = 80
    power = .;
run;
```

You can also express effects in terms of the ratio of discordant proportions and the denominator of the ratio:

```
proc power;
  pairedfreq dist=exact_cond
    discpropratio = 3
    refproportion = 0.15
    npairs = 80
    power = .;
run;
```

You can also express effects in terms of the ratio and sum of discordant proportions:

```
proc power;
  pairedfreq dist=exact_cond
    discpropratio = 3
    totalpropdisc = 0.6
    npairs = 80
    power = .;
run;
```

You can also express effects in terms of the paired proportions and correlation:

```
proc power;
  pairedfreq dist=exact_cond
    pairedproportions = 0.6 | 0.8
    corr = 0.4
    npairs = 45
    power = .;
run;
```

You can also express effects in terms of the proportion difference, reference proportion, and correlation:

```
proc power;
  pairedfreq dist=exact_cond
    proportiondiff = 0.2
    refproportion = 0.6
    corr = 0.4
    npairs = 45
    power = .;
run;
```

You can also express effects in terms of the odds ratio, reference proportion, and correlation:

```
proc power;
  pairedfreq dist=exact_cond
    oddsratio = 2.66667
    refproportion = 0.6
    corr = 0.4
    npairs = 45
    power = .;
run;
```

You can also express effects in terms of the relative risk, reference proportion, and correlation:

```
proc power;
    pairedfreq dist=exact_cond
        relativetisk = 1.33333
        refproportion = 0.6
        corr = 0.4
        npairs = 45
        power = .;
run;
```

McNemar Normal Approximation Test

The following statements demonstrate a sample size computation for the normal-approximate McNemar test. The default value for the **METHOD=** option specifies an exact sample size computation. Default values for the **TEST=**, **SIDES=**, **ALPHA=**, and **NULLDISCPROPRATIO=** options specify a two-sided McNemar test for no effect with a significance level of 0.05.

```
proc power;
    pairedfreq dist=normal method=connor
        discproportions = 0.15 | 0.45
        npairs = .
        power = .9;
run;
```

PAIREDMEANS Statement

PAIREDMEANS <options> ;

The **PAIREDMEANS** statement performs power and sample size analyses for *t* tests, equivalence tests, and confidence interval precision involving paired samples.

Summary of Options

Table 90.19 summarizes the *options* available in the **PAIREDMEANS** statement.

Table 90.21 PAIREDMEANS Statement Options

Option	Description
Define Analysis	
CI=	Specifies an analysis of precision of the confidence interval for the mean difference
DIST=	Specifies the underlying distribution assumed for the test statistic
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
LOWER=	Specifies the lower equivalence bound
NULLDIFF=	Specifies the null mean difference
NULLRATIO=	Specifies the null mean ratio
SIDES=	Specifies the number of sides and the direction of the statistical test or confidence interval

Table 90.21 *continued*

Option	Description
UPPER=	Specifies the upper equivalence bound
Specify Effects	
HALFWIDTH=	Specifies the desired confidence interval half-width
MEANDIFF=	Specifies the mean difference
MEANRATIO=	Specifies the geometric mean ratio, γ_2/γ_1
PAIREDMEANS=	Specifies the two paired means
Specify Variability	
CORR=	Specifies the correlation between members of a pair
CV=	Specifies the common coefficient of variation
PAIREDCVS=	Specifies the coefficient of variation for each member of a pair
PAIREDSTDDEVS=	Specifies the standard deviation of each member of a pair
STDDEV=	Specifies the common standard deviation
Specify Sample Size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NPAIRS=	Specifies the number of pairs
Specify Power and Related Probabilities	
POWER=	Specifies the desired power of the test
PROBTYPE=	Specifies the type of probability for the PROBWIDTH= option
PROBWIDTH=	Specifies the probability of obtaining a confidence interval half-width less than or equal to the value specified by the HALFWIDTH=
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.22 summarizes the valid result parameters for different analyses in the **PAIREDMEANS** statement.

Table 90.22 Summary of Result Parameters in the **PAIREDMEANS** Statement

Analyses	Solve For	Syntax
TEST=DIFF	Power Sample size	POWER=, NPAIRS=.
TEST=RATIO	Power Sample size	POWER=, NPAIRS=.
TEST=EQUIV_DIFF	Power Sample size	POWER=, NPAIRS=.
TEST=EQUIV_RATIO	Power Sample size	POWER=, NPAIRS=.
CI=DIFF	Prob(width) Sample size	PROBWIDTH=, NPAIRS=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. If the **CI=** and **SIDES=1** options are used, then the value must be less than 0.5. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CI

CI=DIFF

specifies an analysis of precision of the confidence interval for the mean difference. Instead of power, the relevant probability for this analysis is the probability of achieving a desired precision. Specifically, it is the probability that the half-width of the observed confidence interval will be at most the value specified by the **HALFWIDTH=** option. If neither the **CI=** option nor the **TEST=** option is used, the default is **TEST=DIFF**.

CORR=*number-list*

specifies the correlation between members of a pair. For tests that assume lognormal data (**DIST=LOGNORMAL**, or **TEST=RATIO** or **TEST=EQUIV_RATIO**), values of the **CORR=** option are restricted to the range (ρ_L, ρ_U) , where

$$\rho_L = \frac{\exp\left(-\left[\log(CV_1^2 + 1)\log(CV_2^2 + 1)\right]^{\frac{1}{2}}\right) - 1}{CV_1 CV_2}$$

$$\rho_U = \frac{\exp\left(\left[\log(CV_1^2 + 1)\log(CV_2^2 + 1)\right]^{\frac{1}{2}}\right) - 1}{CV_1 CV_2}$$

and CV_1 are the CV_2 coefficient of variation values specified by the **CV=** or **PAIREDCVS=** option. See “[Paired *t* Test for Mean Ratio with Lognormal Data \(TEST=RATIO\)](#)” on page 7374 for more information about this restriction on correlation values. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CV=*number-list*

specifies the coefficient of variation that is assumed to be common to both members of a pair. The coefficient of variation is defined as the ratio of the standard deviation to the mean on the original data scale. You can use this option only with **DIST=LOGNORMAL**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

DIST=LOGNORMAL | NORMAL

specifies the underlying distribution assumed for the test statistic. **NORMAL** corresponds the normal distribution, and **LOGNORMAL** corresponds to the lognormal distribution. The default value (also the only acceptable value in each case) is **NORMAL** for **TEST=DIFF**, **TEST=EQUIV_DIFF**, and **CI=DIFF**; and **LOGNORMAL** for **TEST=RATIO** and **TEST=EQUIV_RATIO**.

HALFWIDTH=*number-list*

specifies the desired confidence interval half-width. The half-width is defined as the distance between the point estimate and a finite endpoint. This option can be used only with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

LOWER=number-list

specifies the lower equivalence bound for the mean difference or mean ratio, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). This option can be used only with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

MEANDIFF=number-list

specifies the mean difference, defined as the mean of the difference between the second and first members of a pair, $\mu_2 - \mu_1$. This option can be used only with the **TEST=DIFF** and **TEST=EQUIV_DIFF** analyses. When **TEST=EQUIV_DIFF**, the mean difference is interpreted as the treatment mean minus the reference mean. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

MEANRATIO=number-list

specifies the geometric mean ratio, defined as γ_2/γ_1 . This option can be used only with the **TEST=RATIO** and **TEST=EQUIV_RATIO** analyses. When **TEST=EQUIV_RATIO**, the mean ratio is interpreted as the treatment mean divided by the reference mean. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. See the section “Sample Size Adjustment Options” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPAIRS=number-list

specifies the number of pairs or requests a solution for the number of pairs by specifying a missing value (**NPAIRS=.**). For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NULLDIFF=number-list**NULLD=number-list**

specifies the null mean difference. The default value is 0. This option can be used only with the **TEST=DIFF** analysis. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NULLRATIO=number-list**NULLR=number-list**

specifies the null mean ratio. The default value is 1. This option can be used only with the **TEST=RATIO** analysis. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NULLDIFF=**
- **NULLRATIO=**

- LOWER=
- UPPER=
- ALPHA=
- PAIREDMEANS=
- MEANDIFF=
- MEANRATIO=
- HALFWIDTH=
- STDDEV=
- PAIREDSTDDEVS=
- CV=
- PAIREDCVS=
- CORR=
- NPAIRS=
- POWER=
- PROBTYP=
- PROBWIDTH=

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **PAIREDMEANS** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **PAIREDMEANS** statement.

PAIREDCVS=*grouped-number-list*

specifies the coefficient of variation for each member of a pair. Unlike the **CV=** option, the **PAIREDCVS=** option supports different values for each member of a pair. The coefficient of variation is defined as the ratio of the standard deviation to the mean on the original data scale. Values must be nonnegative (unless both are equal to zero, which is permitted). This option can be used only with **DIST=LOGNORMAL**. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

PAIREDMEANS=*grouped-number-list*

PMEANS=*grouped-number-list*

specifies the two paired means, in the original scale. The means are arithmetic if **DIST=NORMAL** and geometric if **DIST=LOGNORMAL**. This option cannot be used with the **CI=DIFF** analysis. When **TEST=EQUIV_DIFF**, the means are interpreted as the reference mean (first) and the treatment mean (second). For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

PAIREDSTDDEVS=*grouped-number-list*

PAIREDSTD=*grouped-number-list*

PSTDDEVS=*grouped-number-list*

PSTD=*grouped-number-list*

specifies the standard deviation of each member of a pair. Unlike the **STDDEV=** option, the **PAIREDSTDDEVS=** option supports different values for each member of a pair. This option can be used only with **DIST=NORMAL**. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. This option cannot be used with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROBTYPE=keyword-list

specifies the type of probability for the **PROBWIDTH=** option. A value of **CONDITIONAL** (the default) indicates the conditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option, given that the true mean difference is captured by the confidence interval. A value of **UNCONDITIONAL** indicates the unconditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option. you can use the alias **GIVENVALIDITY** for **CONDITIONAL**. The **PROBTYPE=** option can be used only with the **CI=DIFF** analysis. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CONDITIONAL width probability conditional on interval containing the mean

UNCONDITIONAL unconditional width probability

PROBWIDTH=number-list

specifies the desired probability of obtaining a confidence interval half-width less than or equal to the value specified by the **HALFWIDTH=** option. A missing value (**PROBWIDTH=.**) requests a solution for this probability. The type of probability is controlled with the **PROBTYPE=** option. Values are expressed as probabilities (for example, 0.9) rather than percentages. This option can be used only with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. Valid *keywords* and their interpretation for the **TEST=** analyses are as follows:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

For confidence intervals, **SIDES=U** refers to an interval between the lower confidence limit and infinity, and **SIDES=L** refers to an interval between minus infinity and the upper confidence limit. For both of these cases and **SIDES=1**, the confidence interval computations are equivalent. You cannot use the **SIDES=** option with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. By default, **SIDES=2**.

STDDEV=*number-list*

STD=*number-list*

specifies the standard deviation assumed to be common to both members of a pair. This option can be used only with **DIST=NORMAL**. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

TEST=DIFF | EQUIV_DIFF | EQUIV_RATIO | RATIO

TEST

specifies the statistical analysis. **TEST** or **TEST=DIFF** (the default) specifies a paired *t* test on the mean difference. **TEST=EQUIV_DIFF** specifies an additive equivalence test of the mean difference by using a two one-sided tests (TOST) analysis (Schuirmann 1987). **TEST=EQUIV_RATIO** specifies a multiplicative equivalence test of the mean ratio by using a TOST analysis. **TEST=RATIO** specifies a paired *t* test on the geometric mean ratio. If neither the **TEST=** option nor the **CI=** option is used, the default is **TEST=DIFF**.

UPPER=*number-list*

specifies the upper equivalence bound for the mean difference or mean ratio, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). This option can be used only with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

Restrictions on Option Combinations

To define the analysis, choose one of the following parameterizations:

- a statistical test (by using the **TEST=** option)
- confidence interval precision (by using the **CI=** option)

To specify the means, choose one of the following parameterizations:

- individual means (by using the **PAIREDMEANS=** option)
- mean difference (by using the **MEANDIFF=** option)
- mean ratio (by using the **MEANRATIO=** option)

To specify the coefficient of variation, choose one of the following parameterizations:

- common coefficient of variation (by using the **CV=** option)
- individual coefficients of variation (by using the **PAIREDCVS=** option)

To specify the standard deviation, choose one of the following parameterizations:

- common standard deviation (by using the **STDDEV=** option)
- individual standard deviations (by using the **PAIREDSTDDEVS=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **PAIREDMEANS** statement.

Paired *t* Test

You can express effects in terms of the mean difference and variability in terms of a correlation and common standard deviation, as in the following statements. Default values for the **DIST=**, **SIDES=**, **NULLDIFF=**, and **ALPHA=** options specify a two-sided test for no difference with a normal distribution and a significance level of 0.05.

```
proc power;
  pairedmeans test=diff
    meandiff = 7
    corr = 0.4
    stddev = 12
    npairs = 50
    power = .;
run;
```

You can also express effects in terms of individual means and variability in terms of correlation and individual standard deviations:

```
proc power;
  pairedmeans test=diff
    pairedmeans = 8 | 15
    corr = 0.4
    pairedstddevs = (7 12)
    npairs = .
    power = 0.9;
run;
```

Paired *t* Test of Mean Ratio with Lognormal Data

You can express variability in terms of correlation and a common coefficient of variation, as in the following statements. Defaults for the **DIST=**, **SIDES=**, **NULLRATIO=** and **ALPHA=** options specify a two-sided test of mean ratio = 1 assuming a lognormal distribution and a significance level of 0.05.

```
proc power;
  pairedmeans test=ratio
    meanratio = 7
    corr = 0.3
    cv = 1.2
    npairs = 30
    power = .;
run;
```

You can also express variability in terms of correlation and individual coefficients of variation:

```
proc power;
  pairedmeans test=ratio
    meanratio = 7
    corr = 0.3
    pairedcvs = 0.8 | 0.9
```

```

        npairs = 30
        power = .;
run;

```

Additive Equivalence Test for Mean Difference with Normal Data

The following statements demonstrate a sample size computation for a TOST equivalence test for a normal mean difference. Default values for the **DIST=** and **ALPHA=** options specify a normal distribution and a significance level of 0.05.

```

proc power;
    pairedmeans test=equiv_diff
        lower = 2
        upper = 5
        meandiff = 4
        corr = 0.2
        stddev = 8
        npairs = .
        power = 0.9;
run;

```

Multiplicative Equivalence Test for Mean Ratio with Lognormal Data

The following statements demonstrate a power computation for a TOST equivalence test for a lognormal mean ratio. Default values for the **DIST=** and **ALPHA=** options specify a lognormal distribution and a significance level of 0.05.

```

proc power;
    pairedmeans test=equiv_ratio
        lower = 3
        upper = 7
        meanratio = 5
        corr = 0.2
        cv = 1.1
        npairs = 50
        power = .;
run;

```

Confidence Interval for Mean Difference

By default **CI=DIFF** analyzes the conditional probability of obtaining the desired precision, given that the interval contains the true mean difference, as in the following statements. The defaults of **SIDES=2** and **ALPHA=0.05** specify a two-sided interval with a confidence level of 0.95.

```

proc power;
    pairedmeans ci = diff
        halfwidth = 4
        corr = 0.35
        stddev = 8
        npairs = 30
        probwidth = .;
run;

```

PLOT Statement

PLOT *<plot-options>* *</graph-options>* ;

The **PLOT** statement produces a graph or set of graphs for the sample size analysis defined by the previous analysis statement. The *plot-options* define the plot characteristics, and the *graph-options* are SAS/GRAPH-style options. If ODS Graphics is enabled, then the **PLOT** statement uses ODS Graphics to create graphs. For example:

```
ods graphics on;

proc power;
  onesamplemeans
    mean      = 5 10
    ntotal    = 150
    stddev    = 30 50
    power     = .;
  plot x=n min=100 max=200;
run;

ods graphics off;
```

Otherwise, traditional graphics are produced. For example:

```
ods graphics off;

proc power;
  onesamplemeans
    mean      = 5 10
    ntotal    = 150
    stddev    = 30 50
    power     = .;
  plot x=n min=100 max=200;
run;
```

For more information about enabling and disabling ODS Graphics, see the section “[Enabling and Disabling ODS Graphics](#)” on page 607 in Chapter 21, “[Statistical Graphics Using ODS](#).”

Options

You can specify the following *plot-options* in the **PLOT** statement.

INTERPOL=JOIN | NONE

specifies the type of curve to draw through the computed points. The **INTERPOL=JOIN** option connects computed points by straight lines. The **INTERPOL=NONE** option leaves computed points unconnected.

KEY=BYCURVE *<(bycurve-options)>*

KEY=BYFEATURE *<(byfeature-options)>*

KEY=ONCURVES

specifies the style of key (or “legend”) for the plot. The default is **KEY=BYFEATURE**, which specifies a key with a column of entries for each plot feature (line style, color, and/or symbol). Each entry shows

the mapping between a value of the feature and the value(s) of the analysis parameter(s) linked to that feature. The **KEY=BYCURVE** option specifies a key with each row identifying a distinct curve in the plot. The **KEY=ONCURVES** option places a curve-specific label adjacent to each curve.

You can specify the following *byfeature-options* in parentheses after the **KEY=BYCURVE** option.

NUMBERS=OFF | ON

specifies how the key should identify curves. If **NUMBERS=OFF**, then the key includes symbol, color, and line style samples to identify the curves. If **NUMBERS=ON**, then the key includes numbers matching numeric labels placed adjacent to the curves. The default is **NUMBERS=ON**.

POS=BOTTOM | INSET

specifies the position of the key. The **POS=BOTTOM** option places the key below the X axis. The **POS=INSET** option places the key inside the plotting region and attempts to choose the least crowded corner. The default is **POS=BOTTOM**.

You can specify the following *byfeature-options* in parentheses after **KEY=BYFEATURE** option.

POS=BOTTOM | INSET

specifies the position of the key. The **POS=BOTTOM** option places the key below the X axis. The **POS=INSET** option places the key inside the plotting region and attempts to choose the least crowded corner. The default is **POS=BOTTOM**.

MARKERS=ANALYSIS | COMPUTED | NICE | NONE

specifies the locations for plotting symbols.

The **MARKERS=ANALYSIS** option places plotting symbols at locations that correspond to the values of the relevant input parameter from the analysis statement preceding the **PLOT** statement.

The **MARKERS=COMPUTED** option (the default) places plotting symbols at the locations of actual computed points from the sample size analysis.

The **MARKERS=NICE** option places plotting symbols at tick mark locations (corresponding to the argument axis).

The **MARKERS=NONE** option disables plotting symbols.

MAX=number | DATAMAX

specifies the maximum of the range of values for the parameter associated with the “argument” axis (the axis that is *not* representing the parameter being solved for). The default is **DATAMAX**, which specifies the maximum value that occurs for this parameter in the analysis statement that precedes the **PLOT** statement.

MIN=number | DATAMIN

specifies the minimum of the range of values for the parameter associated with the “argument” axis (the axis that is *not* representing the parameter being solved for). The default is **DATAMIN**, which specifies the minimum value that occurs for this parameter in the analysis statement that precedes the **PLOT** statement.

NPOINTS=*number*

NPTS=*number*

specifies the number of values for the parameter associated with the “argument” axis (the axis that is *not* representing the parameter being solved for). You cannot use the **NPOINTS=** and **STEP=** options simultaneously. The default value for typical situations is 20.

STEP=*number*

specifies the increment between values of the parameter associated with the “argument” axis (the axis that is *not* representing the parameter being solved for). You cannot use the **STEP=** and **NPOINTS=** options simultaneously. By default, the **NPOINTS=** option is used instead of the **STEP=** option.

VARY (*feature* < **BY** *parameter-list* > < , ... , *feature* < **BY** *parameter-list* > >)

specifies how plot features should be linked to varying analysis parameters. Available plot *features* are COLOR, LINESSTYLE, PANEL, and SYMBOL. A “panel” refers to a separate plot with a heading identifying the subset of values represented in the plot.

The *parameter-list* is a list of one or more names separated by spaces. Each name must match the name of an analysis option used in the analysis statement preceding the **PLOT** statement. Also, the name must be the *primary* name for the analysis option—that is, the one listed first in the syntax description.

If you omit the < **BY** *parameter-list* > portion for a feature, then one or more multivalued parameters from the analysis will be automatically selected for you.

X=EFFECT | N | POWER

specifies a plot with the requested type of parameter on the X axis and the parameter being solved for on the Y axis. When **X=EFFECT**, the parameter assigned to the X axis is the one most representative of “effect size.” When **X=N**, the parameter assigned to the X axis is the sample size. When **X=POWER**, the parameter assigned to the X axis is the one most representative of “power” (either power itself or a similar probability, such as Prob(Width) for confidence interval analyses). You cannot use the **X=** and **Y=** options simultaneously. The default is **X=POWER**, unless the result parameter is power or Prob(Width), in which case the default is **X=N**.

You can use the **X=N** option only when a scalar sample size parameter is used as input in the analysis. For example, **X=N** can be used with total sample size or sample size per group, or with two group sample sizes when one is being solved for.

Table 90.23 summarizes the parameters that represent effect size in different analyses.

Table 90.23 Effect Size Parameters for Different Analyses

Analysis Statement and Options	Effect Size Parameters
COXREG	Hazard ratio
CUSTOM	Primary noncentrality or correlation
LOGISTIC	None
MULTREG	Partial correlation or R^2 difference
ONECORR	Correlation

Table 90.23 *continued*

Analysis Statement and Options	Effect Size Parameters
ONESAMPLEFREQ TEST	Proportion
ONESAMPLEFREQ CI	CI half-width
ONESAMPLEMEANS TEST=T, ONESAMPLEMEANS TEST=EQUIV	Mean
ONESAMPLEMEANS CI=T	CI half-width
ONEWAYANOVA	None
PAIREDFREQ	Discordant proportion difference or ratio
PAIREDMEANS TEST=DIFF, PAIREDMEANS TEST=EQUIV_DIFF	Mean difference
PAIREDMEANS TEST=RATIO, PAIREDMEANS TEST=EQUIV_RATIO	Mean ratio
PAIREDMEANS CI=DIFF	CI half-width
TWOSAMPLEFREQ	Proportion difference, odds ratio, or relative risk
TWOSAMPLEMEANS TEST=DIFF, TWOSAMPLEMEANS TEST=DIFF_SATT, TWOSAMPLEMEANS TEST=EQUIV_DIFF	Mean difference
TWOSAMPLEMEANS TEST=RATIO, TWOSAMPLEMEANS TEST=EQUIV_RATIO	Mean ratio
TWOSAMPLEMEANS CI=DIFF	CI half-width
TWOSAMPLESURVIVAL	Hazard ratio if used, else none
TWOSAMPLEWILCOXON	None

XOPTS=(*x-options*)

specifies plot characteristics pertaining to the X axis.

You can specify the following *x-options* in parentheses.

CROSSREF=NO | YES

specifies whether the reference lines defined by the **REF=** *x-option* should be crossed with a reference line on the Y axis that indicates the solution point on the curve.

REF=number-list

specifies locations for reference lines extending from the X axis across the entire plotting region. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

Y=EFFECT | N | POWER

specifies a plot with the requested type of parameter on the Y axis and the parameter being solved for on the X axis. When **Y=EFFECT**, the parameter assigned to the Y axis is the one most representative of “effect size.” When **Y=N**, the parameter assigned to the Y axis is the sample size. When **Y=POWER**, the parameter assigned to the Y axis is the one most representative of “power” (either power itself or a similar probability, such as Prob(Width) for confidence interval analyses). You cannot use the **Y=** and **X=** options simultaneously. By default, the **X=** option is used instead of the **Y=** option.

YOPTS=(*y-options*)

specifies plot characteristics pertaining to the Y axis.

You can specify the following *y-options* in parentheses.

CROSSREF=NO | YES

specifies whether the reference lines defined by the **REF= *y-option*** should be crossed with a reference line on the X axis that indicates the solution point on the curve.

REF=*number-list*

specifies locations for reference lines extending from the Y axis across the entire plotting region. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

You can specify the following *graph-options* in the **PLOT** statement after a slash (/).

DESCRIPTION='string'

specifies a descriptive string of up to 40 characters that appears in the “Description” field of the graphics catalog. The description does not appear on the plots. By default, PROC POWER assigns a description either of the form “Y versus X” (for a single-panel plot) or of the form “Y versus X (S),” where Y is the parameter on the Y axis, X is the parameter on the X axis, and S is a description of the subset represented on the current panel of a multipanel plot.

NAME='string'

specifies a name of up to eight characters for the catalog entry for the plot. The default name is PLOT n , where n is the number of the plot statement within the current invocation of PROC POWER. If the name duplicates the name of an existing entry, SAS/GRAPH software adds a number to the duplicate name to create a unique entry—for example, PLOT11 and PLOT12 for the second and third panels of a multipanel plot generated in the first **PLOT** statement in an invocation of PROC POWER.

TWOSAMPLEFREQ Statement

TWOSAMPLEFREQ < *options* > ;

The **TWOSAMPLEFREQ** statement performs power and sample size analyses for tests of two independent proportions. The Farrington-Manning score, Pearson’s chi-square, Fisher’s exact, and likelihood ratio chi-square tests are supported.

Summary of Options

Table 90.24 summarizes the *options* available in the TWOSAMPLEFREQ statement.

Table 90.24 TWOSAMPLEFREQ Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
NULLODDSRATIO=	Specifies the null odds ratio
NULLPROPORTIONDIFF=	Specifies the null proportion difference
NULLRELATIVERISK=	Specifies the null relative risk
SIDES=	Specifies the number of sides and the direction of the statistical test or confidence interval
Specify Effects	
GROUPPROPORTIONS=	Specifies the two independent proportions, p_1 and p_2
ODDSRATIO=	Specifies the odds ratio $[p_2/(1 - p_2)] / [p_1/(1 - p_1)]$
PROPORTIONDIFF=	Specifies the proportion difference $p_2 - p_1$
REFPROPORTION=	Specifies the reference proportion p_1
RELATIVERISK=	Specifies the relative risk p_2/p_1
Specify Sample Size and Allocation	
GROUPNS=	Specifies the two group sample sizes
GROUPWEIGHTS=	Specifies the sample size allocation weights for the two groups
NFRACTIONAL	Enables fractional input and output for sample sizes
NPERGROUP=	Specifies the common sample size per group
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.25 summarizes the valid result parameters for different analyses in the TWOSAMPLEFREQ statement.

Table 90.25 Summary of Result Parameters in the TWOSAMPLEFREQ Statement

Analyses	Solve For	Syntax
TEST=FISHER	Power	POWER=.
	Sample size	NTOTAL=.
TEST=FM	Power	NPERGROUP=.
	Sample size	POWER=.
		NTOTAL=.
		NPERGROUP=.

Table 90.25 *continued*

Analyses	Solve For	Syntax
TEST=FM_RR	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
TEST=LRCHI	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
TEST=PCHI	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPPROPORTIONS=*grouped-number-list*

GPROPORTIONS=*grouped-number-list*

GROUPPS=*grouped-number-list*

GPS=*grouped-number-list*

specifies the two independent proportions, p_1 and p_2 . For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPNS=*grouped-number-list*

GNS=*grouped-number-list*

specifies the two group sample sizes. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPWEIGHTS=*grouped-number-list*

GWEIGHTS=*grouped-number-list*

specifies the sample size allocation weights for the two groups. This option controls how the total sample size is divided between the two groups. Each pair of values for the two groups represents relative allocation weights. Additionally, if the **NFRACTIONAL** option is not used, the total sample size is restricted to be equal to a multiple of the sum of the two group weights (so that the resulting design has an integer sample size for each group while adhering exactly to the group allocation weights). Values must be integers unless the **NFRACTIONAL** option is used. The default value is (1 1), a balanced design with a weight of 1 for each group. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL**NFRAC**

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPERGROUP=number-list**NPERG=number-list**

specifies the common sample size per group or requests a solution for the common sample size per group by specifying a missing value (**NPERGROUP=.**). Use of this option implicitly specifies a balanced design. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NULLODDSRATIO=number-list**NULLOR=number-list**

specifies the null odds ratio. You can specify this option only if you also specify the **ODDSRATIO=** and **TEST=PCHI** options. The **NULLODDSRATIO=** option is inconsistent with **TEST=PCHI**, which tests the proportion difference rather than the odds ratio, and its value is converted internally to a **NULLPROPORTIONDIFF** value by fixing the reference proportion. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. By default, **NULLOR=1**.

NULLPROPORTIONDIFF=number-list**NULLPDIFF=number-list**

specifies the null proportion difference. You can specify this option only if you also specify the **GROUPPROPORTIONS=** or **PROPORTIONDIFF=** option and the **TEST=FM** or **TEST=PCHI** option. If you are using a nondefault null value, then **TEST=FM** is recommended. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. By default, **NULLPDIFF=0**.

NULLRELATIVERISK=number-list**NULLRR=number-list**

specifies the null relative risk. You can specify this option only if you also specify the **GROUPPROPORTIONS=** or **RELATIVERISK=** option and the **TEST=FM_RR** or **TEST=PCHI** option. If you are using a nondefault null value, then **TEST=FM_RR** is recommended. The **NULLRELATIVERISK=** option is inconsistent with **TEST=PCHI**, which tests the proportion difference rather than the relative risk, and its value is converted internally to a **NULLPROPORTIONDIFF** value by fixing the reference proportion. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. By default, **NULLRR=1**.

ODDSRATIO=*number-list*

OR=*number-list*

specifies the odds ratio $[p_2/(1 - p_2)] / [p_1/(1 - p_1)]$. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NULLPROPORTIONDIFF=**
- **NULLODDSRATIO=**
- **NULLRELATIVERISK=**
- **ALPHA=**
- **GROUPPROPORTIONS=**
- **REFPROPORTION=**
- **PROPORTIONDIFF=**
- **ODDSRATIO=**
- **RELATIVERISK=**
- **GROUPWEIGHTS=**
- **NTOTAL=**
- **NPERGROUP=**
- **GROUPNS=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **TWOSAMPLEFREQ** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **TWOSAMPLEFREQ** statement.

POWER=*number-list*

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROPORTIONDIFF=*number-list*

PDIFF=*number-list*

specifies the proportion difference $p_2 - p_1$. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

REFPROPORTION=*number-list*

REFP=*number-list*

specifies the reference proportion p_1 . For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

RELATIVERISK=*number-list*

RR=*number-list*

specifies the relative risk p_2/p_1 . For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

If the effect size is zero, then **SIDES=1** is not permitted; instead, specify the direction of the test explicitly in this case with either **SIDES=L** or **SIDES=U**. By default, **SIDES=2**.

TEST=FISHER | FM | FM_RR | LRCHI | PCHI

specifies the statistical analysis. You can specify the following values:

- FISHER** specifies Fisher’s exact test.
- FM** specifies the score test of Farrington and Manning (1990) for proportion difference.
- FM_RR** specifies the score test of Farrington and Manning (1990) for relative risk.
- LRCHI** specifies the likelihood ratio chi-square test.
- PCHI** specifies Pearson’s chi-square test for proportion difference.

If you are using a nondefault null value for a noninferiority or superiority test, then **TEST=FM** or **TEST=FM_RR** is the most appropriate choice. In the absence of any nondefault null values, the default is **TEST=PCHI**. If you specify at least one nonzero null difference by using the **NULLPROPORTIONDIFF=** option, then the default is **TEST=FM**. If you specify at least one null relative risk not equal to 1 by using the **NULLRR=** option, then the default is **TEST=FM_RR**. For information about the power and sample size computational methods and formulas, see the section “[Analyses in the TWOSAMPLEFREQ Statement](#)” on page 7378.

Restrictions on Option Combinations

To specify the proportions, choose one of the following parameterizations:

- individual proportions (by using the **GROUPPROPORTIONS=** option)
- difference between proportions and reference proportion (by using the **PROPORTIONDIFF=** and **REFPROPORTION=** options)
- odds ratio and reference proportion (by using the **ODDSRATIO=** and **REFPROPORTION=** options)

- relative risk and reference proportion (by using the **RELATIVERISK=** and **REFPROPORTION=** options)

To specify the sample size and allocation, choose one of the following parameterizations:

- sample size per group in a balanced design (by using the **NPERGROUP=** option)
- total sample size and allocation weights (by using the **NTOTAL=** and **GROUPWEIGHTS=** options)
- individual group sample sizes (by using the **GROUPNS=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **TWOSAMPLEFREQ** statement.

Pearson Chi-Square Test for Two Proportions

You can use the **NPERGROUP=** option in a balanced design and express effects in terms of the individual proportions, as in the following statements. Default values for the **SIDES=** and **ALPHA=** options specify a two-sided test with a significance level of 0.05.

```
proc power;
  twosamplefreq test=pchi
    groupproportions = (.15 .25)
    nullproportiondiff = .03
    npergroup = 50
    power = .;
run;
```

You can also specify an unbalanced design by using the **NTOTAL=** and **GROUPWEIGHTS=** options and express effects in terms of the odds ratio. The default value of the **NULLODDSRATIO=** option specifies a test of no effect.

```
proc power;
  twosamplefreq test=pchi
    oddsratio = 2.5
    refproportion = 0.3
    groupweights = (1 2)
    ntotal = .
    power = 0.8;
run;
```

You can also specify sample sizes with the **GROUPNS=** option and express effects in terms of relative risks. The default value of the **NULLRELATIVERISK=** option specifies a test of no effect.

```
proc power;
  twosamplefreq test=pchi
    relativetrisk = 1.5
    refproportion = 0.2
    groupns = 40 | 60
    power = .;
run;
```

You can also express effects in terms of the proportion difference. The default value of the **NULLPROPORTIONDIFF=** option specifies a test of no effect, and the default value of the **GROUPWEIGHTS=** option specifies a balanced design.

```
proc power;
  twosamplefreq test=pchi
    proportiondiff = 0.15
    refproportion = 0.4
    ntotal = 100
    power = .;
run;
```

Farrington-Manning Score Test for Proportion Difference

The following statements demonstrate a sample size computation for the Farrington-Manning score test for the difference of two independent proportions:

```
proc power;
  twosamplefreq test=fm
    proportiondiff = 0.06
    refproportion = 0.32
    nullproportiondiff = -0.02
    sides = u
    ntotal = .
    power = 0.85;
run;
```

Farrington-Manning Score Test for Relative Risk

The following statements demonstrate a sample size computation for the Farrington-Manning score test for the relative risk of two independent proportions:

```
proc power;
  twosamplefreq test=fm_rr
    relativetrisk = 1.1
    refproportion = 0.32
    nullrelativerisk = 0.95
    sides = u
    ntotal = .
    power = 0.9;
run;
```

Fisher's Exact Conditional Test for Two Proportions

The following statements demonstrate a power computation for Fisher's exact conditional test for two proportions. Default values for the **SIDES=** and **ALPHA=** options specify a two-sided test with a significance level of 0.05.

```
proc power;
  twosamplefreq test=fisher
    groupproportions = (.35 .15)
    npergroup = 50
    power = .;
run;
```

Likelihood Ratio Chi-Square Test for Two Proportions

The following statements demonstrate a sample size computation for the likelihood ratio chi-square test for two proportions. Default values for the **SIDES=** and **ALPHA=** options specify a two-sided test with a significance level of 0.05.

```
proc power;
  twosamplefreq test=lrchi
    oddsratio = 2
    refproportion = 0.4
    npergroup = .
    power = 0.9;
run;
```

TWOSAMPLEMEANS Statement

TWOSAMPLEMEANS <options> ;

The **TWOSAMPLEMEANS** statement performs power and sample size analyses for pooled and unpooled *t* tests, equivalence tests, and confidence interval precision involving two independent samples.

Summary of Options

Table 90.26 summarizes the *options* available in the **TWOSAMPLEMEANS** statement.

Table 90.26 TWOSAMPLEMEANS Statement Options

Option	Description
Define Analysis	
CI=	Specifies an analysis of precision of the confidence interval
DIST=	Specifies the underlying distribution assumed for the test statistic
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
LOWER=	Specifies the lower equivalence bound
NULLDIFF=	Specifies the null mean difference
NULLRATIO=	Specifies the null mean ratio
SIDES=	Specifies the number of sides and the direction of the statistical test or confidence interval
UPPER=	Specifies the upper equivalence bound
Specify Effects	
HALFWIDTH=	Specifies the desired confidence interval half-width
GROUPMEANS=	Specifies the two group means
MEANDIFF=	Specifies the mean difference
MEANRATIO=	Specifies the geometric mean ratio, γ_2/γ_1

Table 90.26 *continued*

Option	Description
Specify Variability	
CV=	Specifies the common coefficient of variation
GROUPSTDDEVS=	Specifies the standard deviation of each group
STDDEV=	Specifies the common standard deviation
Specify Sample Size and Allocation	
GROUPNS=	Specifies the two group sample sizes
GROUPWEIGHTS=	Specifies the sample size allocation weights for the two groups
NFRACTIONAL	Enables fractional input and output for sample sizes
NPERGROUP=	Specifies the common sample size per group
NTOTAL=	Specifies the sample size
Specify Power and Related Probabilities	
POWER=	Specifies the desired power of the test
PROBTYPE=	Specifies the type of probability for the PROBWIDTH= option
PROBWIDTH=	Specifies the desired probability of obtaining a confidence interval half-width less than or equal to the value specified
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.27 summarizes the valid result parameters for different analyses in the **TWOSAMPLEMEANS** statement.

Table 90.27 Summary of Result Parameters in the **TWOSAMPLEMEANS** Statement

Analyses	Solve For	Syntax
TEST=DIFF	Power	POWER=.
	Sample size	NTOTAL=.
		NPERGROUP=.
	Group sample size	GROUPNS= <i>n1</i> . GROUPNS= . <i>n2</i> GROUPNS= (<i>n1</i> .) GROUPNS= (. <i>n2</i>)
	Group weight	GROUPWEIGHTS= <i>w1</i> . GROUPWEIGHTS= . <i>w2</i> GROUPWEIGHTS= (<i>w1</i> .) GROUPWEIGHTS= (. <i>w2</i>)
	Alpha	ALPHA=.
	Group mean	GROUPMEANS= <i>mean1</i> . GROUPMEANS= . <i>mean2</i> GROUPMEANS= (<i>mean1</i> .) GROUPMEANS= (. <i>mean2</i>)
	Mean difference	MEANDIFF=.
	Standard deviation	STDDEV=.

Table 90.27 *continued*

Analyses	Solve For	Syntax
TEST=DIFF_SATT	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
TEST=RATIO	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
TEST=EQUIV_DIFF	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
TEST=EQUIV_RATIO	Power Sample size	POWER=. NTOTAL=. NPERGROUP=.
CI=DIFF	Prob(width) Sample size	PROBWIDTH=. NTOTAL=. NPERGROUP=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test or requests a solution for alpha by specifying a missing value (ALPHA=.). The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. If the CI= and SIDES=1 options are used, then the value must be less than 0.5. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CI

CI=DIFF

specifies an analysis of precision of the confidence interval for the mean difference, assuming equal variances. Instead of power, the relevant probability for this analysis is the probability that the interval half-width is at most the value specified by the HALFWIDTH= option. If neither the TEST= option nor the CI= option is used, the default is TEST=DIFF.

CV=*number-list*

specifies the coefficient of variation assumed to be common to both groups. The coefficient of variation is defined as the ratio of the standard deviation to the mean on the original data scale. You can use this option only with DIST=LOGNORMAL. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

DIST=LOGNORMAL | NORMAL

specifies the underlying distribution assumed for the test statistic. NORMAL corresponds to the normal distribution, and LOGNORMAL corresponds to the lognormal distribution. The default value (also the only acceptable value in each case) is NORMAL for TEST=DIFF,

TEST=DIFF_SATT, **TEST=EQUIV_DIFF**, and **CI=DIFF**; and **LOGNORMAL** for **TEST=RATIO** and **TEST=EQUIV_RATIO**.

GROUPMEANS=*grouped-number-list*

GMEANS=*grouped-number-list*

specifies the two group means or requests a solution for one group mean given the other. Means are in the original scale. They are arithmetic if **DIST=NORMAL** and geometric if **DIST=LOGNORMAL**. This option cannot be used with the **CI=DIFF** analysis. When **TEST=EQUIV_DIFF**, the means are interpreted as the reference mean (first) and the treatment mean (second). For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPNS=*grouped-number-list*

GNS=*grouped-number-list*

specifies the two group sample sizes or requests a solution for one group sample size given the other. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPSTDDEVS=*grouped-number-list*

GSTDDEVS=*grouped-number-list*

GROUPSTD=*grouped-number-list*

GSTD=*grouped-number-list*

specifies the standard deviation of each group. Unlike the **STDDEV=** option, the **GROUPSTD-DEVS=** option supports different values for each group. It is valid only for the Satterthwaite *t* test (**TEST=DIFF_SATT** **DIST=NORMAL**). For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPWEIGHTS=*grouped-number-list*

GWEIGHTS=*grouped-number-list*

specifies the sample size allocation weights for the two groups, or requests a solution for one group weight given the other. This option controls how the total sample size is divided between the two groups. Each pair of values for the two groups represents relative allocation weights. Additionally, if the **NFRACTIONAL** option is not used, the total sample size is restricted to be equal to a multiple of the sum of the two group weights (so that the resulting design has an integer sample size for each group while adhering exactly to the group allocation weights). Values must be integers unless the **NFRACTIONAL** option is used. The default value is (1 1), a balanced design with a weight of 1 for each group. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

HALFWIDTH=*number-list*

specifies the desired confidence interval half-width. The half-width is defined as the distance between the point estimate and a finite endpoint. This option can be used only with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

LOWER=*number-list*

specifies the lower equivalence bound for the mean difference or mean ratio, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). Values must be greater than 0

when **DIST=LOGNORMAL**. This option can be used only with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

MEANDIFF=number-list

specifies the mean difference, defined as $\mu_2 - \mu_1$, or requests a solution for the mean difference by specifying a missing value (**MEANDIFF=.**). This option can be used only with the **TEST=DIFF**, **TEST=DIFF_SATT**, and **TEST=EQUIV_DIFF** analyses. When **TEST=EQUIV_DIFF**, the mean difference is interpreted as the treatment mean minus the reference mean. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

MEANRATIO=number-list

specifies the geometric mean ratio, defined as γ_2/γ_1 . This option can be used only with the **TEST=RATIO** and **TEST=EQUIV_RATIO** analyses. When **TEST=EQUIV_RATIO**, the mean ratio is interpreted as the treatment mean divided by the reference mean. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “Sample Size Adjustment Options” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPERGROUP=number-list

NPERG=number-list

specifies the common sample size per group or requests a solution for the common sample size per group by specifying a missing value (**NPERGROUP=.**). Use of this option implicitly specifies a balanced design. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NTOTAL=number-list

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NULLDIFF=number-list

NULLD=number-list

specifies the null mean difference. The default value is 0. This option can be used only with the **TEST=DIFF** and **TEST=DIFF_SATT** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

NULLRATIO=number-list

NULLR=number-list

specifies the null mean ratio. The default value is 1. This option can be used only with the **TEST=RATIO** analysis. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NULLDIFF=**
- **NULLRATIO=**
- **LOWER=**
- **UPPER=**
- **ALPHA=**
- **GROUPMEANS=**
- **MEANDIFF=**
- **MEANRATIO=**
- **HALFWIDTH=**
- **STDDEV=**
- **GROUPSTDDEVS==**
- **CV=**
- **GROUPWEIGHTS=**
- **NTOTAL=**
- **NPERGROUP=**
- **GROUPNS=**
- **POWER=**
- **PROBTYPE=**
- **PROBWIDTH=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **TWOSAMPLEMEANS** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **TWOSAMPLEMEANS** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. This option cannot be used with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

PROBTYPE=keyword-list

specifies the type of probability for the **PROBWIDTH=** option. A value of **CONDITIONAL** (the default) indicates the conditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option, given that the true mean difference is captured by the confidence interval. A value of **UNCONDITIONAL** indicates the unconditional probability that the confidence interval half-width is at most the value specified by the **HALFWIDTH=** option. you can use the alias **GIVENVALIDITY** for **CONDITIONAL**. The **PROBTYPE=** option can be used only with the **CI=DIFF** analysis. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

CONDITIONAL width probability conditional on interval containing the mean

UNCONDITIONAL unconditional width probability

PROBWIDTH=*number-list*

specifies the desired probability of obtaining a confidence interval half-width less than or equal to the value specified by the **HALFWIDTH=** option. A missing value (**PROBWIDTH=.**) requests a solution for this probability. The type of probability is controlled with the **PROBTYPE=** option. Values are expressed as probabilities (for example, 0.9) rather than percentages. This option can be used only with the **CI=DIFF** analysis. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. Valid *keywords* and their interpretation for the **TEST=** analyses are as follows:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

For confidence intervals, **SIDES=U** refers to an interval between the lower confidence limit and infinity, and **SIDES=L** refers to an interval between minus infinity and the upper confidence limit. For both of these cases and **SIDES=1**, the confidence interval computations are equivalent. You cannot use the **SIDES=** option with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. By default, **SIDES=2**.

STDDEV=*number-list*

STD=*number-list*

specifies the standard deviation assumed to be common to both groups, or requests a solution for the common standard deviation with a missing value (**STDDEV=.**). This option can be used only with **DIST=NORMAL**. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TEST=DIFF | DIFF_SATT | EQUIV_DIFF | EQUIV_RATIO | RATIO

TEST

specifies the statistical analysis. **TEST** or **TEST=DIFF** (the default) specifies a pooled *t* test on the mean difference, assuming equal variances. **TEST=DIFF_SATT** specifies a Satterthwaite unpooled *t* test on the mean difference, assuming unequal variances. **TEST=EQUIV_DIFF** specifies an additive equivalence test of the mean difference by using a two one-sided tests (TOST) analysis (Schuirmann 1987). **TEST=EQUIV_RATIO** specifies a multiplicative equivalence test of the mean ratio by using a TOST analysis. **TEST=RATIO** specifies a pooled *t* test on the mean ratio, assuming equal coefficients of variation. If neither the **TEST=** option nor the **CI=** option is used, the default is **TEST=DIFF**.

UPPER=number-list

specifies the upper equivalence bound for the mean difference or mean ratio, in the original scale (whether **DIST=NORMAL** or **DIST=LOGNORMAL**). This option can be used only with the **TEST=EQUIV_DIFF** and **TEST=EQUIV_RATIO** analyses. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

Restrictions on Option Combinations

To define the analysis, choose one of the following parameterizations:

- a statistical test (by using the **TEST=** option)
- confidence interval precision (by using the **CI=** option)

To specify the means, choose one of the following parameterizations:

- individual group means (by using the **GROUPMEANS=** option)
- mean difference (by using the **MEANDIFF=** option)
- mean ratio (by using the **MEANRATIO=** option)

To specify standard deviations in the Satterthwaite *t* test (**TEST=DIFF_SATT**), choose one of the following parameterizations:

- common standard deviation (by using the **STDDEV=** option)
- individual group standard deviations (by using the **GROUPSTDDEVS==** option)

To specify the sample sizes and allocation, choose one of the following parameterizations:

- sample size per group in a balanced design (by using the **NPERGROUP=** option)
- total sample size and allocation weights (by using the **NTOTAL=** and **GROUPWEIGHTS=** options)
- individual group sample sizes (by using the **GROUPNS=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the **TWO SAMPLE MEANS** statement.

Two-Sample *t* Test Assuming Equal Variances

You can use the **NPERGROUP=** option in a balanced design and express effects in terms of the mean difference, as in the following statements. Default values for the **DIST=**, **SIDES=**, **NULLDIFF=**, and **ALPHA=** options specify a two-sided test for no difference with a normal distribution and a significance level of 0.05.

```
proc power;
  twosamplemeans test=diff
    meandiff = 7
    stddev = 12
    npergroup = 50
    power = .;
run;
```

You can also specify an unbalanced design by using the **NTOTAL=** and **GROUPWEIGHTS=** options and express effects in terms of individual group means:

```
proc power;
  twosamplemeans test=diff
    groupmeans = 8 | 15
    stddev = 4
    groupweights = (2 3)
    ntotal = .
    power = 0.9;
run;
```

Another way to specify the sample sizes is with the **GROUPNS=** option:

```
proc power;
  twosamplemeans test=diff
    groupmeans = 8 | 15
    stddev = 4
    groupns = (25 40)
    power = .;
run;
```

Two-Sample Satterthwaite *t* Test Assuming Unequal Variances

The following statements demonstrate a power computation for the two-sample Satterthwaite *t* test allowing unequal variances. Default values for the **DIST=**, **SIDES=**, **NULLDIFF=**, and **ALPHA=** options specify a two-sided test for no difference with a normal distribution and a significance level of 0.05.

```
proc power;
  twosamplemeans test=diff_satt
    meandiff = 3
    groupstddevs = 5 | 8
    groupweights = (1 2)
    ntotal = 60
    power = .;
run;
```

Two-Sample Pooled *t* Test of Mean Ratio with Lognormal Data

The following statements demonstrate a power computation for the pooled *t* test of a lognormal mean ratio. Default values for the **DIST=**, **SIDES=**, **NULLRATIO=**, and **ALPHA=** options specify a two-sided test of mean ratio = 1 assuming a lognormal distribution and a significance level of 0.05.

```
proc power;
  twosamplemeans test=ratio
    meanratio = 7
```

```

    cv = 0.8
    groupns = 50 | 70
    power = .;
run;

```

Additive Equivalence Test for Mean Difference with Normal Data

The following statements demonstrate a sample size computation for the TOST equivalence test for a normal mean difference. A default value of **GROUPWEIGHTS**=(1 1) specifies a balanced design. Default values for the **DIST**= and **ALPHA**= options specify a significance level of 0.05 and an assumption of normally distributed data.

```

proc power;
  twosamplemeans test=equiv_diff
    lower = 2
    upper = 5
    meandiff = 4
    stddev = 8
    ntotal = .
    power = 0.9;
run;

```

Multiplicative Equivalence Test for Mean Ratio with Lognormal Data

The following statements demonstrate a power computation for the TOST equivalence test for a lognormal mean ratio. Default values for the **DIST**= and **ALPHA**= options specify a significance level of 0.05 and an assumption of lognormally distributed data.

```

proc power;
  twosamplemeans test=equiv_ratio
    lower = 3
    upper = 7
    meanratio = 5
    cv = 0.75
    npergroup = 50
    power = .;
run;

```

Confidence Interval for Mean Difference

By default **CI**=DIFF analyzes the conditional probability of obtaining the desired precision, given that the interval contains the true mean difference, as in the following statements. The defaults of **SIDES**=2 and **ALPHA**=0.05 specify a two-sided interval with a confidence level of 0.95.

```

proc power;
  twosamplemeans ci = diff
    halfwidth = 4
    stddev = 8
    groupns = (30 35)
    probwidth = .;
run;

```

TWOSAMPLESURVIVAL Statement

TWOSAMPLESURVIVAL <options> ;

The **TWOSAMPLESURVIVAL** statement performs power and sample size analyses for comparing two survival curves. The log-rank, Gehan, and Tarone-Ware rank tests are supported.

Summary of Options

Table 90.28 summarizes the *options* available in the **TWOSAMPLESURVIVAL** statement.

Table 90.28 TWOSAMPLESURVIVAL Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ACCRUALTIME=	Specifies the accrual time
ALPHA=	Specifies the significance level
FOLLOWUPTIME=	Specifies the follow-up time
SIDES=	Specifies the number of sides and the direction of the statistical test or confidence interval
TOTALTIME=	Specifies the total time
Specify Effects	
CURVE=	Defines a survival curve
GROUPMEDSURVTIMES=	Specifies the median survival times in each group
GROUPSURVEXPHAZARDS=	Specifies exponential hazard rates of the survival curve for each group
GROUPSURVIVAL=	Specifies the survival curve for each group
HAZARDRATIO=	Specifies the hazard ratio
REFSURVEXPHAZARD=	Specifies the exponential hazard rate of the survival curve for the reference group
REFSURVIVAL=	Specifies the survival curve for the reference group
Specify Loss Information	
GROUPLOSS=	Specifies the exponential loss survival curve for each group
GROUPLOSSEXPHAZARDS=	Specifies the exponential hazards of the loss in each group
GROUPMEDLOSSTIMES=	Specifies the median times of the loss in each group
Specify Sample Size and Allocation	
ACCRUALRATEPERGROUP=	Specifies the common accrual rate per group
ACCRUALRATETOTAL=	Specifies the total accrual rate
EVENTSTOTAL=	Specifies the expected total number of events
GROUPACCRUALRATES=	Specifies the accrual rate for each group
GROUPNS=	Specifies the two group sample sizes
GROUPWEIGHTS=	Specifies the sample size allocation weights for the two groups
NFRACTIONAL	Enables fractional input and output for sample sizes
NPERGROUP=	Specifies the common sample size per group
NTOTAL=	Specifies the sample size

Table 90.28 *continued*

Option	Description
Specify Power <code>POWER=</code>	Specifies the desired power of the test
Specify Computational Method <code>NSUBINTERVAL=</code>	Specifies the number of subintervals per unit time
Control Ordering in Output <code>OUTPUTORDER=</code>	Controls the output order of parameters

Table 90.29 summarizes the valid result parameters for different analyses in the `TWOSAMPLESURVIVAL` statement.

Table 90.29 Summary of Result Parameters in the `TWOSAMPLESURVIVAL` Statement

Analyses	Solve For	Syntax
<code>TEST=GEHAN</code>	Power Sample size	<code>POWER=.</code> <code>NTOTAL=.</code> <code>NPERGROUP=.</code> <code>EVENTSTOTAL=.</code> <code>ACCRUALRATETOTAL=.</code> <code>ACCRUALRATEPERGROUP=.</code>
<code>TEST=LOGRANK</code>	Power Sample size	<code>POWER=.</code> <code>NTOTAL=.</code> <code>NPERGROUP=.</code> <code>EVENTSTOTAL=.</code> <code>ACCRUALRATETOTAL=.</code> <code>ACCRUALRATEPERGROUP=.</code>
<code>TEST=TARONEWARE</code>	Power Sample size	<code>POWER=.</code> <code>NTOTAL=.</code> <code>NPERGROUP=.</code> <code>EVENTSTOTAL=.</code> <code>ACCRUALRATETOTAL=.</code> <code>ACCRUALRATEPERGROUP=.</code>

Dictionary of Options

`ACCRUALRATEPERGROUP=number-list`

`ACCRUALRATEPERG=number-list`

`ARPERGROUP=number-list`

`ARPERG=number-list`

specifies the common accrual rate per group or requests a solution for the common accrual rate per group by specifying a missing value (`ACCRUALRATEPERGROUP=.`). The accrual rate per group is

the number of subjects in each group that enters the study per time unit during the accrual period. Use of this option implicitly specifies a balanced design. The **NFRACTIONAL** option is automatically enabled when the **ACCRUALRATEPERGROUP=** option is used. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

ACCRUALRATETOTAL=*number-list*

ARTOTAL=*number-list*

specifies the total accrual rate or requests a solution for the accrual rate by specifying a missing value (**ACCRUALRATETOTAL=.**). The total accrual rate is the total number of subjects that enter the study per time unit during the accrual period. The **NFRACTIONAL** option is automatically enabled when the **ACCRUALRATETOTAL=** option is used. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

ACCRUALTIME=*number-list* | **MAX**

ACCTIME=*number-list* | **MAX**

ACCRUALT=*number-list* | **MAX**

ACCT=*number-list* | **MAX**

specifies the accrual time. Accrual is assumed to occur uniformly from time 0 to the time specified by the **ACCRUALTIME=** option. If the **GROUPSURVIVAL=** or **REFSURVIVAL=** option is used, then the value of the total time (the sum of accrual and follow-up times) must be less than or equal to the largest time in *each* multipoint (piecewise linear) survival curve. If the **ACCRUALRATEPERGROUP=**, **ACCRUALRATETOTAL=**, or **GROUPACCRUALRATES=** option is used, then the accrual time must be greater than 0. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

ACCRUALTIME=MAX can be used when each scenario in the analysis contains at least one piecewise linear survival curve (in the **GROUPSURVIVAL=** or **REFSURVIVAL=** option). It causes the accrual time to be automatically set, separately for each scenario, to the maximum possible time supported by the piecewise linear survival curve(s) in that scenario. It is not compatible with the **FOLLOWUPTIME=MAX** option or the **TOTALTIME=** option.

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

CURVE("label")=*points*

defines a survival curve.

For the **CURVE=** option,

<i>label</i>	identifies the curve in the output and with the GROUPLOSS= , GROUPSURVIVAL= , and REFSURVIVAL= options.
<i>points</i>	specifies one or more (time, survival) pairs on the curve, where the survival value denotes the probability of surviving until at least the specified time.

A single-point curve is interpreted as exponential, and a multipoint curve is interpreted as piecewise linear. Points can be expressed in either of two forms:

- a series of time: survival pairs separated by spaces. For example:

1:0.9 2:0.7 3:0.6

- a DOLIST of times enclosed in parentheses, followed by a colon (:), followed by a DOLIST of survival values enclosed in parentheses. For example:

(1 to 3 by 1) : (0.9 0.7 0.6)

The DOLIST format is the same as in the DATA step.

Points can also be expressed as combinations of the two forms. For example:

1:0.9 2:0.8 (3 to 6 by 1) : (0.7 0.65 0.6 0.55)

The points have the following restrictions:

- The time values must be nonnegative and strictly increasing.
- The survival values must be strictly decreasing.
- The survival value at a time of 0 must be equal to 1.
- If there is only one point, then the time must be greater than 0, and the survival value cannot be 0 or 1.

EVENTSTOTAL=*number-list*

EVENTTOTAL=*number-list*

EETOTAL=*number-list*

specifies the expected total number of events—that is, deaths, whether observed or censored—during the entire study period, or requests a solution for this parameter by specifying a missing value (**EVENTSTOTAL=.**). The **NFRACTIONAL** option is automatically enabled when the **EVENTSTOTAL=** option is used. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

FOLLOWUPTIME=*number-list* | **MAX**

FUTIME=*number-list* | **MAX**

FOLLOWUPT=*number-list* | **MAX**

FUT=*number-list* | **MAX**

specifies the follow-up time, the amount of time in the study past the accrual time. If the **GROUPSURVIVAL=** or **REFSURVIVAL=** option is used, then the value of the total time (the sum of accrual and follow-up times) must be less than or equal to the largest time in *each* multipoint (piecewise linear) survival curve. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

FOLLOWUPTIME=MAX can be used when each scenario in the analysis contains at least one piecewise linear survival curve (in the **GROUPSURVIVAL=** or **REFSURVIVAL=** option). It causes the follow-up time to be automatically set, separately for each scenario, to the maximum possible time supported by the piecewise linear survival curve(s) in that scenario. It is not compatible with the **ACCRUALTIME=MAX** option or the **TOTALTIME=** option.

GROUPACCRUALRATES=*grouped-number-list*

GACCRUALRATES=*grouped-number-list*

GROUPARS=*grouped-number-list*

GARS=*grouped-number-list*

specifies the accrual rate for each group. The groupwise accrual rates are the numbers of subjects in each group that enters the study per time unit during the accrual period. The **NFRACTIONAL** option is automatically enabled when the **GROUPACCRUALRATES=** option is used. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPLOSS=*grouped-name-list*

GLOSS=*grouped-name-list*

specifies the exponential loss survival curve for each group, by using labels specified with the **CURVE=** option. Loss is assumed to follow an exponential curve, indicating the expected rate of censoring (in other words, loss to follow-up) over time. For information about specifying the *grouped-name-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPLOSSEXPHAZARDS=*grouped-number-list*

GLOSSEXPHAZARDS=*grouped-number-list*

GROUPLOSSEXPHS=*grouped-number-list*

GLOSSEXPHS=*grouped-number-list*

specifies the exponential hazards of the loss in each group. Loss is assumed to follow an exponential curve, indicating the expected rate of censoring (in other words, loss to follow-up) over time. If none of the **GROUPLOSSEXPHAZARDS=**, **GROUPLOSS=**, and **GROUPMEDLOSSTIMES=** options are used, the default of **GROUPLOSSEXPHAZARDS=(0 0)** indicates no loss to follow-up. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPMEDLOSSTIMES=*grouped-number-list*

GMEDLOSSTIMES=*grouped-number-list*

GROUPMEDLOSSTS=*grouped-number-list*

GMEDLOSSTS=*grouped-number-list*

specifies the median times of the loss in each group. Loss is assumed to follow an exponential curve, indicating the expected rate of censoring (in other words, loss to follow-up) over time. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPMEDSURVTIMES=*grouped-number-list*

GMEDSURVTIMES=*grouped-number-list*

GROUPMEDSURVTS=*grouped-number-list*

GMEDSURVTS=*grouped-number-list*

specifies the median survival times in each group. When the **GROUPMEDSURVTIMES=** option is used, the survival curve in each group is assumed to be exponential. For information about specifying the *grouped-number-list*, see the section “Specifying Value Lists in Analysis Statements” on page 7329.

GROUPNS=*grouped-number-list*

GNS=*grouped-number-list*

specifies the two group sample sizes. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPSURVEXPHAZARDS=*grouped-number-list*

GSURVEXPHAZARDS=*grouped-number-list*

GROUPSURVEXPHS=*grouped-number-list*

GEXPHS=*grouped-number-list*

specifies exponential hazard rates of the survival curve for each group. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPSURVIVAL=*grouped-name-list*

GSURVIVAL=*grouped-name-list*

GROUPSURV=*grouped-name-list*

GSURV=*grouped-name-list*

specifies the survival curve for each group, by using labels specified with the **CURVE=** option. For information about specifying the *grouped-name-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPWEIGHTS=*grouped-number-list*

GWEIGHTS=*grouped-number-list*

specifies the sample size allocation weights for the two groups. This option controls how the total sample size is divided between the two groups. Each pair of values for the two groups represents relative allocation weights. Additionally, if the **NFRACTIONAL** option is not used, the total sample size is restricted to be equal to a multiple of the sum of the two group weights (so that the resulting design has an integer sample size for each group while adhering exactly to the group allocation weights). Values must be integers unless the **NFRACTIONAL** option is used. The default value is (1 1), a balanced design with a weight of 1 for each group. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

HAZARDRATIO=*number-list*

HR=*number-list*

specifies the hazard ratio of the second group’s survival curve to the first group’s survival curve. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. This option is automatically enabled when any of the following options are used: **ACCRUALRATEPERGROUP=**, **ACCRUALRATETOTAL=**, **EVENTSTOTAL=**, and **GROUPACCRUALRATES=**. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPERGROUP=*number-list*

NPERG=*number-list*

specifies the common sample size per group or requests a solution for the common sample size per group by specifying a missing value (**NPERGROUP**=.). Use of this option implicitly specifies a balanced design. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NSUBINTERVAL=*number-list*

NSUBINTERVALS=*number-list*

NSUB=*number-list*

NSUBS=*number-list*

specifies the number of subintervals per unit time to use in internal calculations. Higher values increase computational time and memory requirements but generally lead to more accurate results. The default value is 12. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NTOTAL=*number-list*

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL**=.). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=**INTERNAL** | **REVERSE** | **SYNTAX**

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER**=**INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES**=
- **ACCRUALTIME**=
- **FOLLOWUPTIME**=
- **TOTALTIME**=
- **NSUBINTERVAL**=
- **ALPHA**=
- **REFSURVIVAL**=
- **GROUPSURVIVAL**=
- **REFSURVEXPHAZARD**=
- **HAZARDRATIO**=
- **GROUPSURVEXPHAZARDS**=
- **GROUPMEDSURVTIMES**=
- **GROUPLOSSEXPHAZARDS**=
- **GROUPLOSS**=
- **GROUPMEDLOSSTIMES**=
- **GROUPWEIGHTS**=
- **NTOTAL**=
- **ACCRUALRATETOTAL**=
- **EVENTSTOTAL**=
- **NPERGROUP**=
- **ACCRUALRATEPERGROUP**=
- **GROUPNS**=

- **GROUPACCRUALRATES=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **TWOSAMPLESURVIVAL** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **TWOSAMPLESURVIVAL** statement.

POWER=*number-list*

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

REFSURVEXPHAZARD=*number-list*

REFSURVEXPH=*number-list*

specifies the exponential hazard rate of the survival curve for the first (reference) group. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

REFSURVIVAL=*name-list*

REFSURV=*name-list*

specifies the survival curve for the first (reference) group, by using labels specified with the **CURVE=** option. For information about specifying the *name-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=*keyword-list*

specifies the number of sides (or tails) and the direction of the statistical test or confidence interval. For information about specifying the *keyword-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis favoring better survival in the second group.
- L** specifies a lower one-sided test, with the alternative hypothesis favoring better survival in the first (reference) group.

By default, **SIDES=2**.

TEST=GEHAN | LOGRANK | TARONEWARE

specifies the statistical analysis. **TEST=GEHAN** specifies the Gehan rank test. **TEST=LOGRANK** (the default) specifies the log-rank test. **TEST=TARONEWARE** specifies the Tarone-Ware rank test.

TOTALTIME=*number-list* | **MAX**

TOTALT=*number-list* | **MAX**

specifies the total time, which is equal to the sum of accrual and follow-up times. If the **GROUPSURVIVAL=** or **REFSURVIVAL=** option is used, then the value of the total time must be less than or equal to the largest time in *each* multipoint (piecewise linear) survival curve. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

TOTALTIME=MAX can be used when each scenario in the analysis contains at least one piecewise linear survival curve (in the **GROUPSURVIVAL=** or **REFSURVIVAL=** option). It causes the total time to be automatically set, separately for each scenario, to the maximum possible time supported by the piecewise linear survival curve(s) in that scenario. It is not compatible with the **ACCRUALTIME=MAX** option or the **FOLLOWUPTIME=MAX** option.

Restrictions on Option Combinations

To specify the survival curves, choose one of the following parameterizations:

- arbitrary piecewise linear or exponential curves (by using the **CURVE=** and **GROUPSURVIVAL=** options)
- curves with proportional hazards (by using the **CURVE=**, **REFSURVIVAL=**, and **HAZARDRATIO=** options)
- exponential curves, by using one of the following parameterizations:
 - median survival times (by using the **GROUPMEDSURVTIMES=** option)
 - the hazard ratio and the hazard of the reference curve (by using the **HAZARDRATIO=** and **REFSURVEXPHAZARD=** options)
 - the individual hazards (by using the **GROUPSURVEXPHAZARDS=** option)

To specify the study time, use any two of the following three options:

- accrual time (by using the **ACCRUALTIME=** option)
- follow-up time (by using the **FOLLOWUPTIME=** option)
- total time, the sum of accrual and follow-up times (by using the **TOTALTIME=** option)

To specify the sample size and allocation, choose one of the following parameterizations:

- sample size per group in a balanced design (by using the **NPERGROUP=** option)
- accrual rate per group in a balanced design (by using the **ACCRUALRATEPERGROUP=** option)
- total sample size and allocation weights (by using the **NTOTAL=** and **GROUPWEIGHTS=** options)
- total accrual rate and allocation weights (by using the **ACCRUALRATETOTAL=** and **GROUPWEIGHTS=** options)
- expected total number of events and allocation weights (by using the **EVENTSTOTAL=** and **GROUPWEIGHTS=** options)
- individual group sample sizes (by using the **GROUPNS=** option)
- individual group accrual rates (by using the **GROUPACCRUALRATES=** option)

The values of parameters that involve expected number of events or accrual rate are converted internally to the analogous sample size parameterization (that is, the `NPERGROUP=`, `NTOTAL=`, or `GROUPNS=` option) for the purpose of sample size adjustments according to the presence or absence of the `NFRACTIONAL` option.

To specify the exponential loss curves, choose one of the following parameterizations:

- a point on the loss curve of each group (by using the `CURVE=` and `GROUPLOSS=` options)
- median loss times (by using the `GROUPMEDLOSSTIMES=` option)
- the individual loss hazards (by using the `GROUPLOSSEXPHAZARDS=` option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the `TWO SAMPLE SURVIVAL` statement.

Log-Rank Test for Two Survival Curves

You can use the `NPERGROUP=` option in a balanced design and specify piecewise linear or exponential survival curves by using the `CURVE=` and `GROUPSURVIVAL=` options, as in the following statements. Default values for the `SIDES=`, `ALPHA=`, `NSUBINTERVAL=`, and `GROUPLOSSEXPHAZARDS=` options specify a two-sided test with a significance level of 0.05, an assumption of no loss to follow-up, and the use of 12 subintervals per unit time in computations.

```
proc power;
  twosamplesurvival test=logrank
    curve("Control")    = (1 2 3):(0.8 0.7 0.6)
    curve("Treatment")  = (5):(.6)
    groupsurvival = "Control" | "Treatment"
    accrualtime = 2
    followuptime = 1
    npergroup = 50
    power = .;
run;
```

In the preceding example, the “Control” curve is piecewise linear (since it has more than one point), and the “Treatment” curve is exponential (since it has only one point).

You can also specify an unbalanced design by using the `NTOTAL=` and `GROUPWEIGHTS=` options and specify piecewise linear or exponential survival curves with proportional hazards by using the `CURVE=`, `REFSURVIVAL=`, and `HAZARDRATIO=` options:

```
proc power;
  twosamplesurvival test=logrank
    curve("Control")    = (1 2 3):(0.8 0.7 0.6)
    refsurvival = "Control"
    hazardratio = 1.5
    accrualtime = 2
    followuptime = 1
    groupweights = (1 2)
    ntotal = .
```

```

        power = 0.8;
run;

```

Instead of computing sample size, you can compute the accrual rate by using the [ACCRUALRATETOTAL=](#) option:

```

proc power;
  twosamplesurvival test=logrank
    curve("Control") = (1 2 3):(0.8 0.7 0.6)
    refsurvival = "Control"
    hazardratio = 1.5
    accrualtime = 2
    followuptime = 1
    groupweights = (1 2)
    accrualratetotal = .
    power = 0.8;
run;

```

or the expected number of events by using the [EVENTSTOTAL=](#) option:

```

proc power;
  twosamplesurvival test=logrank
    curve("Control") = (1 2 3):(0.8 0.7 0.6)
    refsurvival = "Control"
    hazardratio = 1.5
    accrualtime = 2
    followuptime = 1
    groupweights = (1 2)
    eventstotal = .
    power = 0.8;
run;

```

You can also specify sample sizes with the [GROUPNS=](#) option and specify exponential survival curves in terms of median survival times:

```

proc power;
  twosamplesurvival test=logrank
    groupmedsurvtimes = (16 22)
    accrualtime = 6
    totaltime = 18
    groupns = 40 | 60
    power = .;
run;

```

You can also specify exponential survival curves in terms of the hazard ratio and reference hazard. The default value of the [GROUPWEIGHTS=](#) option specifies a balanced design.

```

proc power;
  twosamplesurvival test=logrank
    hazardratio = 1.2
    refsurvexphazard = 0.7
    accrualtime = 2
    totaltime = 4
    ntotal = 100
    power = .;
run;

```

You can also specify exponential survival curves in terms of the individual hazards, as in the following statements:

```
proc power;
  twosamplesurvival test=logrank
    groupsurvexphazards = 0.7 | 0.84
    accrualtime = 2
    totaltime = 4
    ntotal = .
    power = 0.9;
run;
```

Gehan Rank Test for Two Survival Curves

In addition to the log-rank test, you can also specify the Gehan tank test, as in the following statements. Default values for the **SIDES=**, **ALPHA=**, **NSUBINTERVAL=**, and **GROUPLOSSEXPHAZARDS=** options specify a two-sided test with a significance level of 0.05, an assumption of no loss to follow-up, and the use of 12 subintervals per unit time in computations.

```
proc power;
  twosamplesurvival test=gehan
    groupmedsurvtimes = 5 | 7
    accrualtime = 3
    totaltime = 6
    npergroup = .
    power = 0.8;
run;
```

Tarone-Ware Rank Test for Two Survival Curves

You can also specify the Tarone-Ware tank test, as in the following statements. Default values for the **SIDES=**, **ALPHA=**, **NSUBINTERVAL=**, and **GROUPLOSSEXPHAZARDS=** options specify a two-sided test with a significance level of 0.05, an assumption of no loss to follow-up, and the use of 12 subintervals per unit time in computations.

```
proc power;
  twosamplesurvival test=taroneware
    groupmedsurvtimes = 5 | 7
    accrualtime = 3
    totaltime = 6
    npergroup = 100
    power = .;
run;
```

TWOSAMPLEWILCOXON Statement

TWOSAMPLEWILCOXON <options> ;

The **TWOSAMPLEWILCOXON** statement performs power and sample size analyses for the Wilcoxon-Mann-Whitney test (also called the Wilcoxon rank-sum test, Mann-Whitney-Wilcoxon test, or Mann-Whitney U test) for two independent groups.

Note that the O'Brien-Castelloe approach to computing power for the Wilcoxon test is approximate, based on asymptotic behavior as the total sample size gets large. The quality of the power approximation degrades for small sample sizes; conversely, the quality of the sample size approximation degrades if the two distributions are far apart, so that only a small sample is needed to detect a significant difference. But this degradation is rarely a problem in practical situations, in which experiments are usually performed for relatively close distributions.

Summary of Options

Table 90.30 summarizes the *options* available in the **TWOSAMPLEWILCOXON** statement.

Table 90.30 TWOSAMPLEWILCOXON Statement Options

Option	Description
Define Analysis	
TEST=	Specifies the statistical analysis
Specify Analysis Information	
ALPHA=	Specifies the significance level
SIDES=	Specifies the number of sides and the direction of the statistical test
Specify Distributions	
VARDIST=	Defines a distribution for a variable
VARIABLES=	Specifies the distributions of two or more variables
Specify Sample Size and Allocation	
GROUPNS=	Specifies the two group sample sizes
GROUPWEIGHTS=	Specifies the sample size allocation weights for the two groups
NFRACTIONAL	Enables fractional input and output for sample sizes
NPERGROUP=	Specifies the common sample size per group
NTOTAL=	Specifies the sample size
Specify Power	
POWER=	Specifies the desired power of the test
Specify Computational Options	
NBINS=	Specifies the number of categories for each variable
Control Ordering in Output	
OUTPUTORDER=	Controls the output order of parameters

Table 90.31 summarizes the valid result parameters in the **TWOSAMPLEWILCOXON** statement.

Table 90.31 Summary of Result Parameters in the TWOSAMPLEWILCOXON Statement

Analyses	Solve For	Syntax
TEST=WMW	Power	POWER=.
	Sample size	NTOTAL=.
		NPERGROUP=.

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of the statistical test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPNS=*grouped-number-list*

GNS=*grouped-number-list*

specifies the two group sample sizes. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

GROUPWEIGHTS=*grouped-number-list*

GWEIGHTS=*grouped-number-list*

specifies the sample size allocation weights for the two groups. This option controls how the total sample size is divided between the two groups. Each pair of values for the two groups represents relative allocation weights. Additionally, if the **NFRACTIONAL** option is not used, the total sample size is restricted to be equal to a multiple of the sum of the two group weights (so that the resulting design has an integer sample size for each group while adhering exactly to the group allocation weights). Values must be integers unless the **NFRACTIONAL** option is used. The default value is (1 1), a balanced design with a weight of 1 for each group. For information about specifying the *grouped-number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NBINS=*number-list*

specifies the number of categories (or “bins”) each variable’s distribution is divided into (unless it is ordinal, in which case the categories remain intact) in internal calculations. Higher values increase computational time and memory requirements but generally lead to more accurate results. However, if the value is too high, then numerical instability can occur. Lower values are less likely to produce “No solution computed” errors. The default value is 1000. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 7332 for information about the ramifications of the presence (and absence) of the **NFRACTIONAL** option.

NPERGROUP=*number-list*

NPERG=*number-list*

specifies the common sample size per group or requests a solution for the common sample size per group by specifying a missing value (**NPERGROUP=.**). Use of this option implicitly specifies a balanced design. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

NTOTAL=*number-list*

specifies the sample size or requests a solution for the sample size by specifying a missing value (**NTOTAL=.**). For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

OUTPUTORDER=INTERNAL | REVERSE | SYNTAX

controls how the input and default analysis parameters are ordered in the output. **OUTPUTORDER=INTERNAL** (the default) arranges the parameters in the output according to the following order of their corresponding *options*:

- **SIDES=**
- **NBINS=**
- **ALPHA=**
- **VARIABLES=**
- **GROUPWEIGHTS=**
- **NTOTAL=**
- **NPERGROUP=**
- **GROUPNS=**
- **POWER=**

The **OUTPUTORDER=SYNTAX** option arranges the parameters in the output in the same order in which their corresponding options are specified in the **TWOSAMPLEWILCOXON** statement. The **OUTPUTORDER=REVERSE** option arranges the parameters in the output in the reverse of the order in which their corresponding *options* are specified in the **TWOSAMPLEWILCOXON** statement.

POWER=number-list

specifies the desired power of the test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability, a number between 0 and 1, rather than as a percentage. For information about specifying the *number-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

SIDES=keyword-list

specifies the number of sides (or tails) and the direction of the statistical test. You can specify the following *keywords*:

- 1** specifies a one-sided test, with the alternative hypothesis in the same direction as the effect.
- 2** specifies a two-sided test.
- U** specifies an upper one-sided test, with the alternative hypothesis indicating an effect greater than the null value.
- L** specifies a lower one-sided test, with the alternative hypothesis indicating an effect less than the null value.

By default, **SIDES=2**.

TEST=WMW

specifies the Wilcoxon-Mann-Whitney test for two independent groups. This is the default test option.

VARDIST("label")=distribution (parameters)

defines a distribution for a variable.

For the **VARDIST=** option,

- label* identifies the variable distribution in the output and with the **VARIABLES=** option.
- distribution* specifies the distributional form of the variable.

parameters specifies one or more parameters associated with the distribution.

The *distributions* and *parameters* are named and defined in the same way as the distributions and arguments in the CDF SAS function; for more information, see *SAS Language Reference: Dictionary*. Choices for distributional forms and their parameters are as follows:

ORDINAL ((*values*) : (*probabilities*)) is an ordered categorical distribution. The *values* are any numbers separated by spaces. The *probabilities* are numbers between 0 and 1 (inclusive) separated by spaces. Their sum must be exactly 1. The number of *probabilities* must match the number of *values*.

BETA (*a*, *b* < *l*, *r* >) is a beta distribution with shape parameters *a* and *b* and optional location parameters *l* and *r*. The values of *a* and *b* must be greater than 0, and *l* must be less than *r*. The default values for *l* and *r* are 0 and 1, respectively.

BINOMIAL (*p*, *n*) is a binomial distribution with probability of success *p* and number of independent Bernoulli trials *n*. The value of *p* must be greater than 0 and less than 1, and *n* must be an integer greater than 0. If *n* = 1, then the distribution is binary.

EXPONENTIAL (*λ*) is an exponential distribution with scale *λ*, which must be greater than 0.

GAMMA (*a*, *λ*) is a gamma distribution with shape *a* and scale *λ*. The values of *a* and *λ* must be greater than 0.

LAPLACE (*θ*, *λ*) is a Laplace distribution with location *θ* and scale *λ*. The value of *λ* must be greater than 0.

LOGISTIC (*θ*, *λ*) is a logistic distribution with location *θ* and scale *λ*. The value of *λ* must be greater than 0.

LOGNORMAL (*θ*, *λ*) is a lognormal distribution with location *θ* and scale *λ*. The value of *λ* must be greater than 0.

NORMAL (*θ*, *λ*) is a normal distribution with mean *θ* and standard deviation *λ*. The value of *λ* must be greater than 0.

POISSON (*m*) is a Poisson distribution with mean *m*. The value of *m* must be greater than 0.

UNIFORM (*l*, *r*) is a uniform distribution on the interval [*l*, *r*], where *l* < *r*.

VARIABLES=*grouped-name-list*

VAR=*grouped-name-list*

specifies the distributions of two or more variables, using labels specified with the **VARDIST=** option. For information about specifying the *grouped-name-list*, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329.

Restrictions on Option Combinations

To specify the sample size and allocation, choose one of the following parameterizations:

- sample size per group in a balanced design (using the **NPERGROUP=** option)
- total sample size and allocation weights (using the **NTOTAL=** and **GROUPWEIGHTS=** options)
- individual group sample sizes (using the **GROUPNS=** option)

Option Groups for Common Analyses

This section summarizes the syntax for the common analyses that are supported in the `TWOSAMPLEWILCOXON` statement.

Wilcoxon-Mann-Whitney Test for Comparing Two Distributions

The following statements performs a power analysis for Wilcoxon-Mann-Whitney tests comparing an ordinal variable with each other type of distribution. Default values for the `ALPHA=`, `NBINS=`, `SIDES=`, and `TEST=` options specify a two-sided test with a significance level of 0.05 and the use of 1000 categories per distribution when discretization is needed.

```
proc power;
  twosamplewilcoxon
    vardist("myordinal") = ordinal ((0 1 2) : (.2 .3 .5))
    vardist("mybeta1") = beta (1, 2)
    vardist("mybeta2") = beta (1, 2, 0, 2)
    vardist("mybinomial") = binomial (.3, 3)
    vardist("myexponential") = exponential (2)
    vardist("mygamma") = gamma (1.5, 2)
    vardist("mylaplace") = laplace (1, 2)
    vardist("mylogistic") = logistic (1, 2)
    vardist("mylognormal") = lognormal (1, 2)
    vardist("mynormal") = normal (3, 2)
    vardist("mypoisson") = poisson (2)
    vardist("myuniform") = uniform (0, 2)
    variables = "myordinal" | "mybeta1" "mybeta2" "mybinomial"
                  "myexponential" "mygamma" "mylaplace"
                  "mylogistic" "mylognormal" "mynormal"
                  "mypoisson" "myuniform"

    ntotal = 40
    power = .;
run;
```

Details: POWER Procedure

Overview of Power Concepts

In statistical hypothesis testing, you typically express the belief that some effect exists in a population by specifying an alternative hypothesis H_1 . You state a null hypothesis H_0 as the assertion that the effect does *not* exist and attempt to gather evidence to reject H_0 in favor of H_1 . Evidence is gathered in the form of sample data, and a statistical test is used to assess H_0 . If H_0 is rejected but there really is *no* effect, this is called a *Type I error*. The probability of a Type I error is usually designated “alpha” or α , and statistical tests are designed to ensure that α is suitably small (for example, less than 0.05).

If there really is an effect in the population but H_0 is *not* rejected in the statistical test, then a *Type II error* has been made. The probability of a Type II error is usually designated “beta” or β . The probability $1 - \beta$ of avoiding a Type II error—that is, correctly rejecting H_0 and achieving statistical significance—is called the *power*. (NOTE: Another more general definition of power is the probability of rejecting H_0 for any given

set of circumstances, even those that correspond to H_0 being true. The POWER procedure uses this more general definition.)

An important goal in study planning is to ensure an acceptably high level of power. Sample size plays a prominent role in power computations because the focus is often on determining a sufficient sample size to achieve a certain power, or assessing the power for a range of different sample sizes.

Some of the analyses in the POWER procedure focus on *precision* rather than power. An analysis of confidence interval precision is analogous to a traditional power analysis, with “CI Half-Width” taking the place of effect size and “Prob(Width)” taking the place of power. The *CI Half-Width* is the margin of error associated with the confidence interval, the distance between the point estimate and an endpoint. The *Prob(Width)* is the probability of obtaining a confidence interval with *at most* a target half-width.

Summary of Analyses

Table 90.32 gives a summary of the analyses supported in the POWER procedure. The name of the analysis statement reflects the type of data and design. The TEST=, CI=, and DIST= options specify the focus of the statistical hypothesis (in other words, the criterion on which the research question is based) and the test statistic to be used in data analysis.

Table 90.32 Summary of Analyses

Analysis	Statement	Options
Cox proportional hazards regression: score test	COXREG	
Various tests that involve the chi-square, F , t , or normal distribution, or the distribution of the correlation coefficient under multivariate normality	CUSTOM	DIST=CHISQUARE, DIST=CORR, DIST=F, DIST=NORMAL, DIST=T
Logistic regression: likelihood ratio chi-square test	LOGISTIC	
Multiple linear regression: Type III F test	MULTREG	
Correlation: Fisher’s z test	ONECORR	DIST=FISHERZ
Correlation: t test	ONECORR	DIST=T
Binomial proportion: exact test	ONESAMPLEFREQ	TEST=EXACT
Binomial proportion: z test	ONESAMPLEFREQ	TEST=Z
Binomial proportion: z test with continuity adjustment	ONESAMPLEFREQ	TEST=ADJZ
Binomial proportion: exact equivalence test	ONESAMPLEFREQ	TEST=EQUIV_EXACT
Binomial proportion: z equivalence test	ONESAMPLEFREQ	TEST=EQUIV_Z
Binomial proportion: z test with continuity adjustment	ONESAMPLEFREQ	TEST=EQUIV_ADJZ

Table 90.32 continued

Analysis	Statement	Options
Binomial proportion: confidence interval	ONESAMPLEFREQ	CI=AGRESTICOULL, CI=JEFFREYS, CI=EXACT, CI=WALD, CI=WALD_CORRECT, CI=WILSON
One-sample t test	ONESAMPLEMEANS	TEST=T
One-sample t test with lognormal data	ONESAMPLEMEANS	TEST=T
		DIST=LOGNORMAL
One-sample equivalence test for mean of normal data	ONESAMPLEMEANS	TEST=EQUIV
One-sample equivalence test for mean of lognormal data	ONESAMPLEMEANS	TEST=EQUIV DIST=LOGNORMAL
Confidence interval for a mean	ONESAMPLEMEANS	CI=T
One-way ANOVA: one-degree-of-freedom contrast	ONEWAYANOVA	TEST=CONTRAST
One-way ANOVA: overall F test	ONEWAYANOVA	TEST=OVERALL
McNemar exact conditional test	PAIREDFREQ	
McNemar normal approximation test	PAIREDFREQ	DIST=NORMAL
Paired t test	PAIREDMEANS	TEST=DIFF
Paired t test of mean ratio with lognormal data	PAIREDMEANS	TEST=RATIO
Paired additive equivalence of mean difference with normal data	PAIREDMEANS	TEST=EQUIV_DIFF
Paired multiplicative equivalence of mean ratio with lognormal data	PAIREDMEANS	TEST=EQUIV_RATIO
Confidence interval for mean of paired differences	PAIREDMEANS	CI=DIFF
Farrington-Manning score test for proportion difference	TWOSAMPLEFREQ	TEST=FM
Farrington-Manning score test for relative risk	TWOSAMPLEFREQ	TEST=FM_RR
Pearson chi-square test for two independent proportions	TWOSAMPLEFREQ	TEST=PCHI
Fisher's exact test for two independent proportions	TWOSAMPLEFREQ	TEST=FISHER
Likelihood ratio chi-square test for two independent proportions	TWOSAMPLEFREQ	TEST=LRCHI
Two-sample t test assuming equal variances	TWOSAMPLEMEANS	TEST=DIFF
Two-sample Satterthwaite t test assuming unequal variances	TWOSAMPLEMEANS	TEST=DIFF_SATT
Two-sample pooled t test of mean ratio with lognormal data	TWOSAMPLEMEANS	TEST=RATIO

Table 90.32 *continued*

Analysis	Statement	Options
Two-sample additive equivalence of mean difference with normal data	TWOSAMPLEMEANS	TEST=EQUIV_DIFF
Two-sample multiplicative equivalence of mean ratio with lognormal data	TWOSAMPLEMEANS	TEST=EQUIV_RATIO
Two-sample confidence interval for mean difference	TWOSAMPLEMEANS	CI=DIFF
Log-rank test for comparing two survival curves	TWOSAMPLESURVIVAL	TEST=LOGRANK
Gehan rank test for comparing two survival curves	TWOSAMPLESURVIVAL	TEST=GEHAN
Tarone-Ware rank test for comparing two survival curves	TWOSAMPLESURVIVAL	TEST=TARONEWARE
Wilcoxon-Mann-Whitney (rank-sum) test	TWOSAMPLEWILCOXON	

Specifying Value Lists in Analysis Statements

To specify one or more scenarios for an analysis parameter (or set of parameters), you provide a list of values for the statement option that corresponds to the parameter(s). To identify the parameter you want to solve for, you place missing values in the appropriate list.

There are five basic types of such lists: *keyword-lists*, *number-lists*, *grouped-number-lists*, *name-lists*, and *grouped-name-lists*. Some parameters, such as the direction of a test, have values represented by one or more keywords in a *keyword-list*. Scenarios for scalar-valued parameters, such as power, are represented by a *number-list*. Scenarios for groups of scalar-valued parameters, such as group sample sizes in a multigroup design, are represented by a *grouped-number-list*. Scenarios for named parameters, such as reference survival curves, are represented by a *name-list*. Scenarios for groups of named parameters, such as group survival curves, are represented by a *grouped-name-list*.

The following subsections explain these five basic types of lists.

Keyword-Lists

A *keyword-list* is a list of one or more keywords, separated by spaces. For example, you can specify both two-sided and upper-tailed versions of a one-sample *t* test as follows:

```
SIDES = 2 U
```

Number-Lists

A *number-list* can be one of two things: a series of one or more numbers expressed in the form of one or more DOLISTs, or a missing value indicator (.).

The DOLIST format is the same as in the DATA step language. For example, for the one-sample *t* test you can specify four scenarios (30, 50, 70, and 100) for a total sample size in any of the following ways.

```
NTOTAL = 30 50 70 100
NTOTAL = 30 to 70 by 20 100
```

A missing value identifies a parameter as the result parameter; it is valid only with options representing parameters you can solve for in a given analysis. For example, you can request a solution for NTOTAL as follows:

```
NTOTAL = .
```

Grouped-Number-Lists

A *grouped-number-list* specifies multiple scenarios for numeric values in two or more groups, possibly including missing value indicators to solve for a specific group. The list can assume one of two general forms, a “crossed” version and a “matched” version.

Crossed Grouped-Number-Lists

The crossed version of a grouped number list consists of a series of *number-lists* (see the section “[Number-Lists](#)” on page 7329), one representing each group, with groups separated by a vertical bar (|). The values for each group represent multiple scenarios for that group, and the scenarios for each individual group are crossed to produce the set of all scenarios for the analysis option. For example, you can specify the following six scenarios for the sizes (n_1, n_2) of two groups

```
(20, 30)(20, 40)(20, 50)
(25, 30)(25, 40)(25, 50)
```

as follows:

```
GROUPNS = 20 25 | 30 40 50
```

If the analysis can solve for a value in one group given the other groups, then one of the *number-lists* in a *crossed grouped-number-list* can be a missing value indicator (.). For example, in a two-sample *t* test you can posit three scenarios for the group 2 sample size while solving for the group 1 sample size:

```
GROUPNS = . | 30 40 50
```

Some analyses can involve more than two groups. For example, you can specify $2 \times 3 \times 1 = 6$ scenarios for the means of three groups in a one-way ANOVA as follows:

```
GROUPMEANS = 10 12 | 10 to 20 by 5 | 24
```

Matched Grouped-Number-Lists

The matched version of a grouped number list consists of a series of numeric lists, each enclosed in parentheses. Each list consists of a value for each group and represents a single scenario for the analysis option. Multiple scenarios for the analysis option are represented by multiple lists. For example, you can express the crossed grouped-number-list

```
GROUPNS = 20 25 | 30 40 50
```

alternatively in a matched format:

```
GROUPNS = (20 30) (20 40) (20 50) (25 30) (25 40) (25 50)
```

The matched version is particularly useful when you want to include only a subset of all combinations of individual group values. For example, you might want to pair 20 only with 50, and 25 only with 30 and 40:

```
GROUPNS = (20 50) (25 30) (25 40)
```

If the analysis can solve for a value in one group given the other groups, then you can replace the value for that group with a missing value indicator (.). If used, the missing value indicator must occur in the same group in every scenario. For example, you can solve for the group 1 sample size (as in the section “[Crossed Grouped-Number-Lists](#)” on page 7330) by using a matched format:

```
GROUPNS = (. 30) (. 40) (. 50)
```

Some analyses can involve more than two groups. For example, you can specify two scenarios for the means of three groups in a one-way ANOVA:

```
GROUPMEANS = (15 24 32) (12 25 36)
```

Name-Lists

A *name-list* is a list of one or more names that are enclosed in single or double quotation marks and separated by spaces. For example, you can specify two scenarios for the reference survival curve in a log-rank test as follows:

```
REFSURVIVAL = "Curve A" "Curve B"
```

Grouped-Name-Lists

A *grouped-name-list* specifies multiple scenarios for names in two or more groups. The list can assume one of two general forms, a “crossed” version and a “matched” version.

Crossed Grouped-Name-Lists

The crossed version of a grouped name list consists of a series of *name-lists* (see the section “[Name-Lists](#)” on page 7331), one representing each group, with groups separated by a vertical bar (|). The values for each group represent multiple scenarios for that group, and the scenarios for each individual group are crossed to produce the set of all scenarios for the analysis option. For example, you can specify the following six scenarios for the survival curves (c_1, c_2) of two groups

```
("Curve A", "Curve C")("Curve A", "Curve D")("Curve A", "Curve E")
("Curve B", "Curve C")("Curve B", "Curve D")("Curve B", "Curve E")
```

as follows:

```
GROUPSURVIVAL = "Curve A" "Curve B" | "Curve C" "Curve D"
                "Curve E"
```

Matched Grouped-Name-Lists

The matched version of a grouped name list consists of a series of name lists, each enclosed in parentheses. Each list consists of a name for each group and represents a single scenario for the analysis option. Multiple scenarios for the analysis option are represented by multiple lists. For example, you can express the crossed grouped-name-list

```
GROUPSURVIVAL = "Curve A" "Curve B" | "Curve C" "Curve D"
                  "Curve E"
```

alternatively in a matched format:

```
GROUPSURVIVAL = ("Curve A" "Curve C")
                  ("Curve A" "Curve D")
                  ("Curve A" "Curve E")
                  ("Curve B" "Curve C")
                  ("Curve B" "Curve D")
                  ("Curve B" "Curve E")
```

The matched version is particularly useful when you want to include only a subset of all combinations of individual group values. For example, you might want to pair “Curve A” only with “Curve C”, and “Curve B” only with “Curve D” and “Curve E”:

```
GROUPSURVIVAL = ("Curve A" "Curve C")
                  ("Curve B" "Curve D")
                  ("Curve B" "Curve E")
```

Sample Size Adjustment Options

By default, PROC POWER rounds sample sizes conservatively (down in the input, up in the output) so that all total sizes (and individual group sample sizes, if a multigroup design) are integers. This is generally considered conservative because it selects the closest realistic design providing *at most* the power of the (possibly fractional) input or mathematically optimized design. In addition, in a multigroup design, all group sizes are adjusted to be multiples of the corresponding group weights. For example, if GROUPWEIGHTS = (2 6), then all group 1 sample sizes become multiples of 2, and all group 2 sample sizes become multiples of 6 (and all total sample sizes become multiples of 8).

With the NFRACTIONAL option, sample size input is not rounded, and sample size output (whether total or groupwise) are reported in two versions, a raw “fractional” version and a “ceiling” version rounded up to the nearest integer.

Whenever an input sample size is adjusted, both the original (“nominal”) and adjusted (“actual”) sample sizes are reported. Whenever computed output sample sizes are adjusted, both the original input (“nominal”) power and the achieved (“actual”) power at the adjusted sample size are reported.

Error and Information Output

The Error column in the main output table provides reasons for missing results and flags numerical results that are bounds rather than exact answers. For example, consider the sample size analysis implemented by the following statements:

```
proc power;
  twosamplefreq test=fm_rr
    relativetrisk=1.0001
    refproportion=.4
    nullrelativetrisk=1
    power=.9
    ntotal=.;
run;
```

Figure 90.6 Error Column

The POWER Procedure Farrington-Manning Score Test for Relative Risk

Fixed Scenario Elements		
Distribution	Asymptotic normal	
Method	Normal approximation	
Null Relative Risk	1	
Reference (Group 1) Proportion	0.4	
Relative Risk	1.0001	
Nominal Power	0.9	
Number of Sides	2	
Alpha	0.05	
Group 1 Weight	1	
Group 2 Weight	1	
Computed N Total		
Actual		
Power	N Total	Error
0.473	2.15E+09	Solution is a lower bound

The output in [Figure 90.6](#) reveals that the sample size to achieve a power of 0.9 could not be computed, but that the sample size 2.15E+09 achieves a power of 0.473.

The Info column provides further details about Error column entries, warnings about any boundary conditions detected, and notes about any adjustments to input. Note that the Info column is hidden by default in the main output. You can view it by using the ODS OUTPUT statement to save the output as a data set and the PRINT procedure. For example, the following SAS statements print both the Error and Info columns for a power computation in a two-sample *t* test:

```
proc power;
  twosamplemeans
    meandiff= 0 7
    stddev=2
    ntotal=2 5
    power=.;
```

```

ods output output=Power;
run;

proc print noobs data=Power;
  var MeanDiff NominalNTotal NTotal Power Error Info;
run;

```

The output is shown in Figure 90.7.

Figure 90.7 Error and Info Columns

MeanDiff	NominalNTotal	NTotal	Power	Error	Info
0	2	2	.	Invalid input	N too small / No effect
0	5	4	0.050		Input N adjusted / No effect
7	2	2	.	Invalid input	N too small
7	5	4	0.477		Input N adjusted

The mean difference of 0 specified with the **MEANDIFF=** option leads to a “No effect” message to appear in the Info column. The sample size of 2 specified with the **NTOTAL=** option leads to an “Invalid input” message in the Error column and an “NTotal too small” message in the Info column. The sample size of 5 leads to an “Input N adjusted” message in the Info column because it is rounded down to 4 to produce integer group sizes of 2 per group.

Displayed Output

If you use the **PLOTONLY** option in the **PROC POWER** statement, the procedure displays only graphical output. Otherwise, the displayed output of the POWER procedure includes the following:

- the “Fixed Scenario Elements” table, which shows all applicable single-valued analysis parameters, in the following order: distribution, method, parameters that are input explicitly, and parameters that are supplied with defaults
- an output table that shows the following when applicable (in order): the index of the scenario, all multivalued input, ancillary results, the primary computed result, and error descriptions
- plots (if requested)

For each input parameter, the order of the input values is preserved in the output.

Ancillary results include the following:

- Actual Power, the achieved power, if it differs from the input (Nominal) power value
- Actual Prob(Width), the achieved precision probability, if it differs from the input (Nominal) probability value
- Actual Alpha, the achieved significance level, if it differs from the input (Nominal) alpha value
- fractional sample size, if the **NFRACTIONAL** option is used in the analysis statement

If sample size is the result parameter and the NFRATIONAL option is used in the analysis statement, then both “Fractional” and “Ceiling” sample size results are displayed. Fractional sample sizes correspond to the “Nominal” values of power or precision probability. Ceiling sample sizes are simply the fractional sample sizes rounded up to the nearest integer; they correspond to “Actual” values of power or precision probability.

ODS Table Names

PROC POWER assigns a name to each table that it creates. You can use these names to reference the table when using the Output Delivery System (ODS) to select tables and create output data sets. These names are listed in Table 90.33. For more information about ODS, see Chapter 20, “Using the Output Delivery System.”

Table 90.33 ODS Tables Produced by PROC POWER

ODS Table Name	Description	Statement
FixedElements	Factoid with single-valued analysis parameters	Default*
Output	All input and computed analysis parameters, error messages, and information messages for each scenario	Default
PlotContent	Data contained in plots, including analysis parameters and indices identifying plot features. (NOTE: This table is saved as a data set and not displayed in PROC POWER output.)	PLOT

*Depends on input.

Computational Resources

Memory

In the TWOSAMPLESURVIVAL statement, the amount of required memory is roughly proportional to the product of the number of subintervals (specified by the NSUBINTERVAL= option) and the total time of the study (specified by the ACCRUALTIME=, FOLLOWUPTIME=, and TOTALTIME= options). If you run out of memory, then you can try either specifying a smaller number of subintervals, changing the time scale to a use a longer time unit (for example, years instead of months), or both.

CPU Time

In the Satterthwaite t test analysis (TWOSAMPLEMEANS TEST=DIFF_SATT), the required CPU time grows as the mean difference decreases relative to the standard deviations. In the PAIREDFREQ statement, the required CPU time for the exact power computation (METHOD=EXACT) grows with the sample size.

Computational Methods and Formulas

This section describes the approaches that PROC POWER uses to compute power for each analysis. The first subsection defines some common notation. The following subsections describe the various power analyses,

including discussions of the data, statistical test, and power formula for each analysis. Unless otherwise indicated, computed values for parameters besides power (for example, sample size) are obtained by solving power formulas for the desired parameters.

Common Notation

Table 90.34 displays notation for some of the more common parameters across analyses. The Associated Syntax column shows examples of relevant analysis statement options, where applicable.

Table 90.34 Common Notation

Symbol	Description	Associated Syntax
α	Significance level	ALPHA=
N	Total sample size	NTOTAL=, NPAIRS=
n_i	Sample size in i th group	NPERGROUP=, GROUPNS=
w_i	Allocation weight for i th group (standardized to sum to 1)	GROUPWEIGHTS=
μ	(Arithmetic) mean	MEAN=
μ_i	(Arithmetic) mean in i th group	GROUPMEANS=, PAIREDMEANS=
μ_{diff}	(Arithmetic) mean difference, $\mu_2 - \mu_1$ or $\mu_T - \mu_R$	MEANDIFF=
μ_0	Null mean or mean difference (arithmetic)	NULL=, NULLDIFF=
γ	Geometric mean	MEAN=
γ_i	Geometric mean in i th group	GROUPMEANS=, PAIREDMEANS=
γ_0	Null mean or mean ratio (geometric)	NULL=, NULLRATIO=
σ	Standard deviation (or common standard deviation per group)	STDDEV=
σ_i	Standard deviation in i th group	GROUPSTDDEVS=, PAIREDSTDDEVS=
σ_{diff}	Standard deviation of differences	
CV	Coefficient of variation, defined as the ratio of the standard deviation to the (arithmetic) mean on the original data scale	CV=, PAIREDCVS=
ρ	Correlation	CORR=
μ_T, μ_R	Treatment and reference (arithmetic) means for equivalence test	GROUPMEANS=, PAIREDMEANS=
γ_T, γ_R	Treatment and reference geometric means for equivalence test	GROUPMEANS=, PAIREDMEANS=
θ_L	Lower equivalence bound	LOWER=
θ_U	Upper equivalence bound	UPPER=
$t(\nu, \delta)$	t distribution with $df \ \nu$ and noncentrality δ	
$F(\nu_1, \nu_2, \lambda)$	F distribution with numerator $df \ \nu_1$, denominator $df \ \nu_2$, and noncentrality λ	
$t_{p;\nu}$	p th percentile of t distribution with $df \ \nu$	
$F_{p;\nu_1, \nu_2}$	p th percentile of F distribution with numerator $df \ \nu_1$ and denominator $df \ \nu_2$	

Table 90.34 *continued*

Symbol	Description	Associated Syntax
$\text{Bin}(N, p)$	Binomial distribution with sample size N and proportion p	

A “lower one-sided” test is associated with SIDES=L (or SIDES=1 with the effect smaller than the null value), and an “upper one-sided” test is associated with SIDES=U (or SIDES=1 with the effect larger than the null value).

Owen (1965) defines a function, known as Owen’s Q , that is convenient for representing terms in power formulas for confidence intervals and equivalence tests:

$$Q_v(t, \delta; a, b) = \frac{\sqrt{2\pi}}{\Gamma(\frac{v}{2})2^{\frac{v-2}{2}}} \int_a^b \Phi\left(\frac{tx}{\sqrt{v}} - \delta\right) x^{v-1} \phi(x) dx$$

where $\phi(\cdot)$ and $\Phi(\cdot)$ are the density and cumulative distribution function of the standard normal distribution, respectively.

Analyses in the COXREG Statement

Score Test of a Single Scalar Predictor in Cox Proportional Hazards Regression (TEST=SCORE)

The power-computing formula is based on Hsieh and Lavori (2000, equation (2) and the section “Variance Inflation Factor” on page 556).

Define the following notation for a Cox proportional hazards regression analysis:

- N = #subjects (NTOTAL)
- K = #predictors
- $\mathbf{x} = (x_1, \dots, x_K)'$ = vector of predictors
- x_1 = predictor of interest
- $\mathbf{x}_{-1} = (x_2, \dots, x_K)'$
- $h(t|\mathbf{x})$ = hazard function for survival time given $\mathbf{x}$, evaluated at time t
- $h_0(t)$ = baseline hazard at time t
- h_r = hazard ratio for one-unit increase in x_1 (HAZARDRATIO)
- p_e = Prob(event is uncensored) (EVENTPROB)
- σ = standard deviation of x_1 (STDDEV)
- ρ = Corr($\mathbf{x}_{-1}, x_1$)
- $R^2 = \rho^2 = R^2$ value from regression of x_1 on $\mathbf{x}_{-1}$ (RSQUARE)

The Cox proportional hazards regression model is

$$\begin{aligned} \log(h(t|\mathbf{x})/h_0(t)) &= \beta\mathbf{x} \\ &= \beta_1 x_1 + \dots + \beta_K x_K \end{aligned}$$

You can convert a regression coefficient to a hazard ratio by using the equation $h_r = \exp(\beta_1)$.

The hypothesis test of the first predictor variable is

$$H_0: h_r = 1$$

$$H_1: \begin{cases} h_r \neq 1, & \text{two-sided} \\ h_r < 1, & \text{upper one-sided} \\ h_r > 1, & \text{lower one-sided} \end{cases}$$

The upper and lower one-sided cases are expressed differently than in other analyses. This is because $h_r > 1$ corresponds to a negative correlation between the tested predictor and survival and thus, by the convention used in PROC POWER for regression analyses, the lower side.

The approximate power is

$$\text{power} = \begin{cases} \Phi\left(z_\alpha - \sigma \sqrt{N p_e(1 - R^2)} \log(h_r)\right), & \text{upper one-sided} \\ 1 - \Phi\left(z_{1-\alpha} - \sigma \sqrt{N p_e(1 - R^2)} \log(h_r)\right), & \text{lower one-sided} \\ \Phi\left(z_{\frac{\alpha}{2}} - \sigma \sqrt{N p_e(1 - R^2)} \log(h_r)\right) + 1 - \Phi\left(z_{1-\frac{\alpha}{2}} - \sigma \sqrt{N p_e(1 - R^2)} \log(h_r)\right), & \text{two-sided} \end{cases}$$

For the one-sided cases, a closed-form inversion of the power equation yields an approximate total sample size

$$N = \left(\frac{(z_{\text{power}} + z_{1-\alpha})^2}{p_e(1 - R^2)\sigma^2 \log(h_r)} \right)$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Analyses in the CUSTOM Statement

The first of the following subsections specifies notation for the analyses in the CUSTOM statement, and the remaining subsections show power formulas and special cases equivalent to analyses elsewhere in PROC POWER or PROC GLMPOWER for each type of test statistic distribution.

Notation

Table 90.35 displays notation for the analyses in the CUSTOM statement.

Table 90.35 Common Notation

Symbol	Description	Associated Option
N	Total sample size	NTOTAL=
p_1	Test degrees of freedom	TESTDF=
p	Model degrees of freedom	MODELDF=
λ^*	Primary noncentrality for χ^2 or F distribution	PRIMNC=
δ^*	Primary noncentrality for t or normal distribution	PRIMNC=
$\rho_{YX_1 X_{-1}}$	Partial correlation between $p_1 \geq 1$ variables in a set X_1 and a variable Y , adjusting for $p - p_1$ variables in a (possibly empty) set X_{-1} and assuming multivariate normality	CORR=

Table 90.35 *continued*

Symbol	Description	Associated Option
m_1	Scalar multiplier of primary noncentrality	PNCMULT=
m_2	Scalar multiplier of critical value	CRITMULT=
α	Significance level	ALPHA=

Noncentral Chi-Square Distribution (DIST=CHISQUARE)

The power is computed as

$$\text{power} = P(\chi^2(p_1, Nm_1\lambda^*) \geq m_2\chi_{1-\alpha}^2(p_1))$$

The sample size is computed by numerically inverting the power formula.

The logistic regression analysis in the **LOGISTIC** statement is a special case in which $p_1 = 1$ and $m_2 = 1$. The two-sided Wilcoxon-Mann-Whitney test (**TWOSAMPLEWILCOXON**) is a special case in which $m_1 = 1$.

Correlation Coefficient Distribution Assuming Multivariate Normal Data (DIST=CORR)

The distribution of the correlation coefficient given the value of $\rho_{YX_1|X_{-1}}$ is shown in (Johnson, Kotz, and Balakrishnan 1995, Chapter 32). The one-sided cases are restricted to $p_1 = 1$.

The power is computed as

$$\text{power} = \begin{cases} P \left[R_{YX_1|X_{-1}}^2 \geq \frac{F_{1-\alpha}(p_1, N-p)}{F_{1-\alpha}(p_1, N-p) + \frac{N-p}{p_1}} \mid \rho_{YX_1|X_{-1}} \right] & \text{two-sided} \\ P \left[R_{YX_1|X_{-1}} \geq \frac{t_{1-\alpha}(N-p-1)}{(t_{1-\alpha}^2(N-p-1) + N-p-1)^{\frac{1}{2}}} \mid \rho_{YX_1|X_{-1}} \right] & \text{upper one-sided} \\ P \left[R_{YX_1|X_{-1}} \leq \frac{t_{\alpha}(N-p-1)}{(t_{\alpha}^2(N-p-1) + N-p-1)^{\frac{1}{2}}} \mid \rho_{YX_1|X_{-1}} \right] & \text{lower one-sided} \end{cases}$$

The sample size is computed by numerically inverting the power formula.

The Pearson correlation analysis that assumes joint multivariate normality (**ONECORR DIST=T MODEL=RANDOM**) and the multiple regression analysis that assumes joint multivariate normal predictors (**MULTREG MODEL=RANDOM**) are special cases.

Noncentral F Distribution (DIST=F)

The power is computed as

$$\text{power} = P(F(p_1, N-p, Nm_1\lambda^*) \geq F_{1-\alpha}(p_1, N-p))$$

The sample size is computed by numerically inverting the power formula.

The two-sided contrast in one-way analysis of variance (**ONEWAYANOVA**) is a special case in which $p_1 = 1$ and $m_1 = 1$. The following analyses are special cases in which $m_1 = 1$:

- multiple regression that assumes fixed predictors (**MULTREG MODEL=FIXED**)

- overall F test in one-way analysis of variance (ONEWAYANOVA)
- Type III F test in a fixed-effect univariate linear model (PROC GLMPower)
- multivariate Type III F test in a fixed-effect multivariate linear model (PROC GLMPower)

Normal Distribution ($DIST=NORMAL$)

The power is computed as

$$\text{power} = \begin{cases} \Phi\left(m_2 z_\alpha - N^{\frac{1}{2}} m_1 \delta^*\right) & \text{upper one-sided} \\ 1 - \Phi\left(m_2 z_{1-\alpha} - N^{\frac{1}{2}} m_1 \delta^*\right) & \text{lower one-sided} \\ P\left(\chi^2(1, N(m_1 \delta^*)^2) \geq m_2^2 \chi_{1-\alpha}^2(1)\right) & \text{two-sided} \end{cases}$$

The sample size is computed by numerically inverting the power equation.

Cox proportional hazards regression (COXREG) and the log-rank test (TWO SAMPLESURVIVAL) are special cases in which $m_1 = 1$ and $m_2 = 1$. The Wilcoxon-Mann-Whitney test (TWO SAMPLEWILCOXON) is a special case in which $m_1 = 1$. Other special cases occur among the tests for binomial proportions (ONESAMPLEFREQ, PAIREDFREQ, and TWO SAMPLEFREQ).

Noncentral T Distribution ($DIST=T$)

The power is computed as

$$\text{power} = \begin{cases} P\left(F(1, N-p, N(m_1 \delta^*)^2) \geq F_{1-\alpha}(1, N-p)\right) & \text{two-sided} \\ P\left(t(N-p, N^{\frac{1}{2}} m_1 \delta^*) \geq t_{1-\alpha}(N-p)\right) & \text{upper one-sided} \\ P\left(t(N-p, N^{\frac{1}{2}} m_1 \delta^*) \leq t_\alpha(N-p)\right) & \text{lower one-sided} \end{cases}$$

The sample size is computed by numerically inverting the power formula.

The standard t tests in the ONESAMPLEMEANS and PAIREDMEANS statements are special cases in which $p = 1$ and $m_1 = 1$. The pooled t test (TWO SAMPLEMEANS) is a special case in which $p = 2$ and $m_1 = 1$. The Pearson correlation analysis that assumes fixed X values (ONECORR $DIST=T$ $MODEL=FIXED$) is a special case in which $m_1 = 1$ and $p = p^* + 2$, where p^* is the number of variables that are partialled out. The one-sided contrast in one-way analysis of variance (ONEWAYANOVA) is special case in which $m_1 = 1$ and $p_1 = G$, where G is the number of groups.

Analyses in the LOGISTIC Statement

Likelihood Ratio Chi-Square Test for One Predictor ($TEST=LRCHI$)

The power-computing formula is based on Shieh and O'Brien (1998); Shieh (2000); Self, Mauritsen, and Ohara (1992), and Hsieh (1989).

Define the following notation for a logistic regression analysis:

$$\begin{aligned}
 N &= \text{\#subjects} \quad (\text{NTOTAL}) \\
 K &= \text{\#predictors (not counting intercept)} \\
 \mathbf{x} &= (x_1, \dots, x_K)' = \text{random variables for predictor vector} \\
 \mathbf{x}_{-1} &= (x_2, \dots, x_K)' \\
 \boldsymbol{\mu} &= (\mu_1, \dots, \mu_K)' = E\mathbf{x} = \text{mean predictor vector} \\
 \mathbf{x}_i &= (x_{i1}, \dots, x_{iK})' = \text{predictor vector for subject } i \quad (i \in 1, \dots, N) \\
 Y &= \text{random variable for response (0 or 1)} \\
 Y_i &= \text{response for subject } i \quad (i \in 1, \dots, N) \\
 p_i &= \text{Prob}(Y_i = 1 | \mathbf{x}_i) \quad (i \in 1, \dots, N) \\
 \phi &= \text{Prob}(Y_i = 1 | \mathbf{x}_i = \boldsymbol{\mu}) \quad (\text{RESPONSEPROB}) \\
 U_j &= \text{unit change for } j \text{th predictor} \quad (\text{UNITS}) \\
 \text{OR}_j &= \text{Odds}(Y_i = 1 | x_{ij} = c) / \text{Odds}(Y_i = 1 | x_{ij} = c - U_j) \quad (c \text{ arbitrary}, i \in 1, \dots, N, \\
 &\quad j \in 1, \dots, K) \quad (\text{TESTODDSRATIO if } j = 1, \text{COVODDSRATIOS if } j > 1) \\
 \Psi_0 &= \text{intercept in full model} \quad (\text{INTERCEPT}) \\
 \boldsymbol{\Psi} &= (\Psi_1, \dots, \Psi_K)' = \text{regression coefficients in full model} \\
 &\quad (\Psi_1 = \text{TESTREGCOEFF, others} = \text{COVREGCOEFFS}) \\
 \rho &= \text{Corr}(\mathbf{x}_{-1}, x_1) \quad (\text{CORR}) \\
 c_j &= \text{\#distinct possible values of } x_{ij} \quad (j \in 1, \dots, K)(\text{for any } i) \quad (\text{NBINS}) \\
 x_{gj}^* &= g \text{th possible value of } x_{ij} \quad (g \in 1, \dots, c_j)(j \in 1, \dots, K) \\
 &\quad (\text{for any } i) \quad (\text{VARDIST}) \\
 \pi_{gj} &= \text{Prob}(x_{ij} = x_{gj}^*) \quad (g \in 1, \dots, c_j)(j \in 1, \dots, K) \\
 &\quad (\text{for any } i) \quad (\text{VARDIST}) \\
 C &= \prod_{j=1}^K c_j = \text{\#possible values of } \mathbf{x}_i \quad (\text{for any } i) \\
 \mathbf{x}_m^* &= m \text{th possible value of } \mathbf{x}_i \quad (m \in 1, \dots, C) \\
 \pi_m &= \text{Prob}(\mathbf{x}_i = \mathbf{x}_m^*) \quad (m \in 1, \dots, C)
 \end{aligned}$$

The logistic regression model is

$$\log \left(\frac{p_i}{1 - p_i} \right) = \Psi_0 + \boldsymbol{\Psi}' \mathbf{x}_i$$

The hypothesis test of the first predictor variable is

$$\begin{aligned}
 H_0: \Psi_1 &= 0 \\
 H_1: \Psi_1 &\neq 0
 \end{aligned}$$

Assuming independence among all predictor variables, π_m is defined as follows:

$$\pi_m = \prod_{j=1}^K \pi_{h(m,j)j} \quad (m \in 1, \dots, C)$$

where $h(m, j)$ is calculated according to the following algorithm:

```

z = m;
do j = K to 1;
  h(m, j) = mod(z - 1, c_j) + 1;
  z = floor((z - 1)/c_j) + 1;
end;

```

This algorithm causes the elements of the transposed vector $\{h(m, 1), \dots, h(m, K)\}$ to vary fastest to slowest from right to left as m increases, as shown in the following table of $h(m, j)$ values:

$h(m, j)$	j				
	1	2	$\dots$	$K-1$	K
1	1	1	$\dots$	1	1
1	1	1	$\dots$	1	2
$\vdots$			$\vdots$		
$\vdots$	1	1	$\dots$	1	c_K
$\vdots$	1	1	$\dots$	2	1
$\vdots$	1	1	$\dots$	2	2
$\vdots$			$\vdots$		
m	1	1	$\dots$	2	c_K
$\vdots$			$\vdots$		
$\vdots$	c_1	c_2	$\dots$	c_{K-1}	1
$\vdots$	c_1	c_2	$\dots$	c_{K-1}	2
$\vdots$			$\vdots$		
C	c_1	c_2	$\dots$	c_{K-1}	c_K

The $\mathbf{x}_m^*$ values are determined in a completely analogous manner.

The discretization is handled as follows (unless the distribution is ordinal, or binomial with sample size parameter at least as large as requested number of bins): for x_j , generate c_j quantiles at evenly spaced probability values such that each such quantile is at the midpoint of a bin with probability $\frac{1}{c_j}$. In other words,

$$x_{gj}^* = \left(\frac{g - 0.5}{c_j} \right) \text{th quantile of relevant distribution}$$

$$(g \in 1, \dots, c_j)(j \in 1, \dots, K)$$

$$\pi_{gj} = \frac{1}{c_j} \quad (\text{same for all } g)$$

The primary noncentrality for the power computation is

$$\Delta^* = 2 \sum_{m=1}^C \pi_m [b'(\theta_m) (\theta_m - \theta_m^*) - (b(\theta_m) - b(\theta_m^*))]$$

where

$$\begin{aligned} b'(\theta) &= \frac{\exp(\theta)}{1 + \exp(\theta)} \\ b(\theta) &= \log(1 + \exp(\theta)) \\ \theta_m &= \Psi_0 + \Psi' \mathbf{x}_m^* \\ \theta_m^* &= \Psi_0^* + \Psi^{*'} \mathbf{x}_m^* \end{aligned}$$

where

$$\begin{aligned} \Psi_0^* &= \Psi_0 + \Psi_1 \mu_1 = \text{intercept in reduced model, absorbing the tested predictor} \\ \Psi^* &= (0, \Psi_2, \dots, \Psi_K)' = \text{coefficients in reduced model} \end{aligned}$$

The power is

$$\text{power} = P(\chi^2(1, \Delta^* N(1 - \rho^2)) \geq \chi_{1-\alpha}^2(1))$$

The factor $(1 - \rho^2)$ is the adjustment for correlation between the predictor that is being tested and other predictors, from Hsieh (1989).

Alternative input parameterizations are handled by the following transformations:

$$\begin{aligned} \Psi_0 &= \log\left(\frac{\phi}{1 - \phi}\right) - \Psi' \mu \\ \Psi_j &= \frac{\log(\text{OR}_j)}{U_j} \quad (j \in 1, \dots, K) \end{aligned}$$

Analyses in the MULTREG Statement

Type III F Test in Multiple Regression (TEST=TYPE3)

Maxwell (2000) discusses a number of different ways to represent effect sizes (and to compute exact power based on them) in multiple regression. PROC POWER supports two of these, multiple partial correlation and R^2 in full and reduced models.

Let p denote the total number of predictors in the full model (excluding the intercept), and let Y denote the response variable. You are testing that the coefficients of $p_1 \geq 1$ predictors in a set X_1 are 0, controlling for all of the other predictors X_{-1} , which consists of $p - p_1 \geq 0$ variables.

The hypotheses can be expressed in two different ways. The first is in terms of $\rho_{YX_1|X_{-1}}$, the multiple partial correlation between the predictors in X_1 and the response Y adjusting for the predictors in X_{-1} :

$$\begin{aligned} H_0: \rho_{YX_1|X_{-1}}^2 &= 0 \\ H_1: \rho_{YX_1|X_{-1}}^2 &> 0 \end{aligned}$$

The second is in terms of the multiple correlations in full ($\rho_{Y|(X_1, X_{-1})}$) and reduced ($\rho_{Y|X_{-1}}$) nested models:

$$H_0: \rho_{Y|(X_1, X_{-1})}^2 - \rho_{Y|X_{-1}}^2 = 0$$

$$H_1: \rho_{Y|(X_1, X_{-1})}^2 - \rho_{Y|X_{-1}}^2 > 0$$

Note that the squared values of $\rho_{Y|(X_1, X_{-1})}$ and $\rho_{Y|X_{-1}}$ are the population R^2 values for full and reduced models.

The test statistic can be written in terms of the sample multiple partial correlation $R_{YX_1|X_{-1}}$,

$$F = \begin{cases} (N - 1 - p) \frac{R_{YX_1|X_{-1}}^2}{1 - R_{YX_1|X_{-1}}^2}, & \text{intercept} \\ (N - p) \frac{R_{YX_1|X_{-1}}^2}{1 - R_{YX_1|X_{-1}}^2}, & \text{no intercept} \end{cases}$$

or the sample multiple correlations in full ($R_{Y|(X_1, X_{-1})}$) and reduced ($R_{Y|X_{-1}}$) models,

$$F = \begin{cases} (N - 1 - p) \frac{R_{Y|(X_1, X_{-1})}^2 - R_{Y|X_{-1}}^2}{1 - R_{Y|(X_1, X_{-1})}^2}, & \text{intercept} \\ (N - p) \frac{R_{Y|(X_1, X_{-1})}^2 - R_{Y|X_{-1}}^2}{1 - R_{Y|(X_1, X_{-1})}^2}, & \text{no intercept} \end{cases}$$

The test is the usual Type III F test in multiple regression:

$$\text{Reject } H_0 \text{ if } \begin{cases} F \geq F_{1-\alpha}(p_1, N - 1 - p), & \text{intercept} \\ F \geq F_{1-\alpha}(p_1, N - p), & \text{no intercept} \end{cases}$$

Although the test is invariant to whether the predictors are assumed to be random or fixed, the power is affected by this assumption. If the response and predictors are assumed to have a joint multivariate normal distribution, then the exact power is given by the following formula:

$$\begin{aligned} \text{power} &= \begin{cases} P \left[\left(\frac{N-1-p}{p_1} \right) \left(\frac{R_{YX_1|X_{-1}}^2}{1 - R_{YX_1|X_{-1}}^2} \right) \geq F_{1-\alpha}(p_1, N - 1 - p) \right], & \text{intercept} \\ P \left[\left(\frac{N-p}{p_1} \right) \left(\frac{R_{YX_1|X_{-1}}^2}{1 - R_{YX_1|X_{-1}}^2} \right) \geq F_{1-\alpha}(p_1, N - p) \right], & \text{no intercept} \end{cases} \\ &= \begin{cases} P \left[R_{YX_1|X_{-1}}^2 \geq \frac{F_{1-\alpha}(p_1, N-1-p)}{F_{1-\alpha}(p_1, N-1-p) + \frac{N-1-p}{p_1}} \right], & \text{intercept} \\ P \left[R_{YX_1|X_{-1}}^2 \geq \frac{F_{1-\alpha}(p_1, N-p)}{F_{1-\alpha}(p_1, N-p) + \frac{N-p}{p_1}} \right], & \text{no intercept} \end{cases} \end{aligned}$$

The distribution of $R_{YX_1|X_{-1}}^2$ (for any $\rho_{YX_1|X_{-1}}^2$) is given in Chapter 32 of Johnson, Kotz, and Balakrishnan (1995). Sample size tables are presented in Gatsonis and Sampson (1989).

If the predictors are assumed to have fixed values, then the exact power is given by the noncentral F distribution. The noncentrality parameter is

$$\lambda = N \frac{\rho_{YX_1|X_{-1}}^2}{1 - \rho_{YX_1|X_{-1}}^2}$$

or equivalently,

$$\lambda = N \frac{\rho_{Y|(X_1, X_{-1})}^2 - \rho_{Y|X_{-1}}^2}{1 - \rho_{Y|(X_1, X_{-1})}^2}$$

The power is

$$\text{power} = \begin{cases} P(F(p_1, N - 1 - p, \lambda) \geq F_{1-\alpha}(p_1, N - 1 - p)), & \text{intercept} \\ P(F(p_1, N - p, \lambda) \geq F_{1-\alpha}(p_1, N - p)), & \text{no intercept} \end{cases}$$

The minimum acceptable input value of N depends on several factors, as shown in Table 90.36.

Table 90.36 Minimum Acceptable Sample Size Values in the MULTREG Statement

Predictor Type	Intercept in Model?	$p_1 = 1$?	Minimum N
Random	Yes	Yes	$p + 3$
Random	Yes	No	$p + 2$
Random	No	Yes	$p + 2$
Random	No	No	$p + 1$
Fixed	Yes	Yes or No	$p + 2$
Fixed	No	Yes or No	$p + 1$

Analyses in the ONECORR Statement

Fisher's z Test for Pearson Correlation (TEST=PEARSON DIST=FISHERZ)

Fisher's z transformation (Fisher 1921) of the sample correlation $R_{Y|(X_1, X_{-1})}$ is defined as

$$z = \frac{1}{2} \log \left(\frac{1 + R_{Y|(X_1, X_{-1})}}{1 - R_{Y|(X_1, X_{-1})}} \right)$$

Fisher's z test assumes the approximate normal distribution $N(\mu, \sigma^2)$ for z , where

$$\mu = \frac{1}{2} \log \left(\frac{1 + \rho_{Y|(X_1, X_{-1})}}{1 - \rho_{Y|(X_1, X_{-1})}} \right) + \frac{\rho_{Y|(X_1, X_{-1})}}{2(N - 1 - p^*)}$$

and

$$\sigma^2 = \frac{1}{N - 3 - p^*}$$

where p^* is the number of variables partialled out (Anderson 1984, pp. 132–133) and $\rho_{Y|(X_1, X_{-1})}$ is the partial correlation between Y and X_1 adjusting for the set of zero or more variables X_{-1} .

The test statistic

$$z^* = (N - 3 - p^*)^{\frac{1}{2}} \left[z - \frac{1}{2} \log \left(\frac{1 + \rho_0}{1 - \rho_0} \right) - \frac{\rho_0}{2(N - 1 - p^*)} \right]$$

is assumed to have a normal distribution $N(\delta, \nu)$, where ρ_0 is the null partial correlation and δ and ν are derived from Section 16.33 of Stuart and Ord (1994):

$$\delta = (N - 3 - p^*)^{\frac{1}{2}} \left[\frac{1}{2} \log \left(\frac{1 + \rho_{Y|(X_1, X_{-1})}}{1 - \rho_{Y|(X_1, X_{-1})}} \right) + \frac{\rho_{Y|(X_1, X_{-1})}}{2(N - 1 - p^*)} \left(1 + \frac{5 + \rho_{Y|(X_1, X_{-1})}^2}{4(N - 1 - p^*)} + \frac{11 + 2\rho_{Y|(X_1, X_{-1})}^2 + 3\rho_{Y|(X_1, X_{-1})}^4}{8(N - 1 - p^*)^2} \right) - \frac{1}{2} \log \left(\frac{1 + \rho_0}{1 - \rho_0} \right) - \frac{\rho_0}{2(N - 1 - p^*)} \right]$$

$$\nu = \frac{N - 3 - p^*}{N - 1 - p^*} \left[1 + \frac{4 - \rho_{Y|(X_1, X_{-1})}^2}{2(N - 1 - p^*)} + \frac{22 - 6\rho_{Y|(X_1, X_{-1})}^2 - 3\rho_{Y|(X_1, X_{-1})}^4}{6(N - 1 - p^*)^2} \right]$$

The approximate power is computed as

$$\text{power} = \begin{cases} \Phi \left(\frac{\delta - z_{1-\alpha}}{\nu^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-\delta - z_{1-\alpha}}{\nu^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{\delta - z_{1-\frac{\alpha}{2}}}{\nu^{\frac{1}{2}}} \right) + \Phi \left(\frac{-\delta - z_{1-\frac{\alpha}{2}}}{\nu^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

Because the test is biased, the achieved significance level might differ from the nominal significance level. The actual alpha is computed in the same way as the power, except that the correlation $\rho_{Y|(X_1, X_{-1})}$ is replaced by the null correlation ρ_0 .

***t* Test for Pearson Correlation (TEST=PEARSON DIST=T)**

The two-sided case is identical to multiple regression with an intercept and $p_1 = 1$, which is discussed in the section “Analyses in the MULTREG Statement” on page 7343.

Let p^* denote the number of variables partialled out. For the one-sided cases, the test statistic is

$$t = (N - 2 - p^*)^{\frac{1}{2}} \frac{R_{YX_1|X_{-1}}}{\left(1 - R_{YX_1|X_{-1}}^2\right)^{\frac{1}{2}}}$$

which is assumed to have a null distribution of $t(N - 2 - p^*)$.

If the X and Y variables are assumed to have a joint multivariate normal distribution, then the exact power is given by the following formula:

$$\begin{aligned} \text{power} &= \begin{cases} P \left[(N-2-p^*)^{\frac{1}{2}} \frac{R_{YX_1|X_{-1}}}{(1-R_{YX_1|X_{-1}}^2)^{\frac{1}{2}}} \geq t_{1-\alpha}(N-2-p^*) \right], & \text{upper one-sided} \\ P \left[(N-2-p^*)^{\frac{1}{2}} \frac{R_{YX_1|X_{-1}}}{(1-R_{YX_1|X_{-1}}^2)^{\frac{1}{2}}} \leq t_{\alpha}(N-2-p^*) \right], & \text{lower one-sided} \end{cases} \\ &= \begin{cases} P \left[R_{Y|(X_1, X_{-1})} \geq \frac{t_{1-\alpha}(N-2-p^*)}{(t_{1-\alpha}^2(N-2-p^*) + N-2-p^*)^{\frac{1}{2}}} \right], & \text{upper one-sided} \\ P \left[R_{Y|(X_1, X_{-1})} \leq \frac{t_{\alpha}(N-2-p^*)}{(t_{\alpha}^2(N-2-p^*) + N-2-p^*)^{\frac{1}{2}}} \right], & \text{lower one-sided} \end{cases} \end{aligned}$$

The distribution of $R_{Y|(X_1, X_{-1})}$ (given the underlying true correlation $\rho_{Y|(X_1, X_{-1})}$) is given in Chapter 32 of Johnson, Kotz, and Balakrishnan (1995).

If the X variables are assumed to have fixed values, then the exact power is given by the noncentral t distribution $t(N-2-p^*, \delta)$, where the noncentrality is

$$\delta = N^{\frac{1}{2}} \frac{\rho_{YX_1|X_{-1}}}{(1 - \rho_{YX_1|X_{-1}}^2)^{\frac{1}{2}}}$$

The power is

$$\text{power} = \begin{cases} P(t(N-2-p^*, \delta) \geq t_{1-\alpha}(N-2-p^*)), & \text{upper one-sided} \\ P(t(N-2-p^*, \delta) \leq t_{\alpha}(N-2-p^*)), & \text{lower one-sided} \end{cases}$$

Analyses in the ONESAMPLEFREQ Statement

Exact Test of a Binomial Proportion (TEST=EXACT)

Let X be distributed as $\text{Bin}(N, p)$. The hypotheses for the test of the proportion p are as follows:

$$H_0: p = p_0$$

$$H_1: \begin{cases} p \neq p_0, & \text{two-sided} \\ p > p_0, & \text{upper one-sided} \\ p < p_0, & \text{lower one-sided} \end{cases}$$

The exact test assumes binomially distributed data and requires $N \geq 1$ and $0 < p_0 < 1$. The test statistic is

$$X = \text{number of successes} \sim \text{Bin}(N, p)$$

The significance probability α is split symmetrically for two-sided tests, in the sense that each tail is filled with as much as possible up to $\alpha/2$.

Exact power computations are based on the binomial distribution and computing formulas such as the following from Johnson, Kotz, and Kemp (1992, equation 3.20):

$$P(X \geq C|N, p) = P\left(F_{v_1, v_2} \leq \frac{v_2 p}{v_1(1-p)}\right) \quad \text{where } v_1 = 2C \text{ and } v_2 = 2(N - C + 1)$$

Let C_L and C_U denote lower and upper critical values, respectively. Let α_a denote the achieved (actual) significance level, which for two-sided tests is the sum of the favorable major tail (α_M) and the opposite minor tail (α_m).

For the upper one-sided case,

$$\begin{aligned} C_U &= \min\{C : P(X \geq C|p_0) \leq \alpha\} \\ \text{Reject } H_0 &\text{ if } X \geq C_U \\ \alpha_a &= P(X \geq C_U|p_0) \\ \text{power} &= P(X \geq C_U|p) \end{aligned}$$

For the lower one-sided case,

$$\begin{aligned} C_L &= \max\{C : P(X \leq C|p_0) \leq \alpha\} \\ \text{Reject } H_0 &\text{ if } X \leq C_L \\ \alpha_a &= P(X \leq C_L|p_0) \\ \text{power} &= P(X \leq C_L|p) \end{aligned}$$

For the two-sided case,

$$\begin{aligned} C_L &= \max\{C : P(X \leq C|p_0) \leq \frac{\alpha}{2}\} \\ C_U &= \min\{C : P(X \geq C|p_0) \leq \frac{\alpha}{2}\} \\ \text{Reject } H_0 &\text{ if } X \leq C_L \text{ or } X \geq C_U \\ \alpha_a &= P(X \leq C_L \text{ or } X \geq C_U|p_0) \\ \text{power} &= P(X \leq C_L \text{ or } X \geq C_U|p) \end{aligned}$$

z Test for Binomial Proportion Using Null Variance (TEST=Z VAREST=NULL)

For the normal approximation test, the test statistic is

$$Z(X) = \frac{X - Np_0}{[Np_0(1 - p_0)]^{\frac{1}{2}}}$$

For the METHOD=EXACT option, the computations are the same as described in the section “[Exact Test of a Binomial Proportion \(TEST=EXACT\)](#)” on page 7347 except for the definitions of the critical values.

For the upper one-sided case,

$$C_U = \min\{C : Z(C) \geq z_{1-\alpha}\}$$

For the lower one-sided case,

$$C_L = \max\{C : Z(C) \leq z_\alpha\}$$

For the two-sided case,

$$C_L = \max\{C : Z(C) \leq z_{\frac{\alpha}{2}}\}$$

$$C_U = \min\{C : Z(C) \geq z_{1-\frac{\alpha}{2}}\}$$

For the METHOD=NORMAL option, the test statistic $Z(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{\frac{1}{2}}(p - p_0)}{[p_0(1 - p_0)]^{\frac{1}{2}}}, \frac{p(1 - p)}{p_0(1 - p_0)}\right)$$

The approximate power is computed as

$$\text{power} = \begin{cases} \Phi\left(\frac{z_\alpha + \sqrt{N} \frac{p - p_0}{\sqrt{p_0(1 - p_0)}}}{\sqrt{\frac{p(1 - p)}{p_0(1 - p_0)}}}\right), & \text{upper one-sided} \\ \Phi\left(\frac{z_\alpha - \sqrt{N} \frac{p - p_0}{\sqrt{p_0(1 - p_0)}}}{\sqrt{\frac{p(1 - p)}{p_0(1 - p_0)}}}\right), & \text{lower one-sided} \\ \Phi\left(\frac{z_{\frac{\alpha}{2}} + \sqrt{N} \frac{p - p_0}{\sqrt{p_0(1 - p_0)}}}{\sqrt{\frac{p(1 - p)}{p_0(1 - p_0)}}}\right) + \Phi\left(\frac{z_{\frac{\alpha}{2}} - \sqrt{N} \frac{p - p_0}{\sqrt{p_0(1 - p_0)}}}{\sqrt{\frac{p(1 - p)}{p_0(1 - p_0)}}}\right), & \text{two-sided} \end{cases}$$

The approximate sample size is computed in closed form for the one-sided cases by inverting the power equation,

$$N = \left(\frac{z_{\text{power}} \sqrt{p(1 - p)} + z_{1-\alpha} \sqrt{p_0(1 - p_0)}}{p - p_0}\right)^2$$

and by numerical inversion for the two-sided case.

z Test for Binomial Proportion Using Sample Variance (TEST=Z VAREST=SAMPLE)

For the normal approximation test using the sample variance, the test statistic is

$$Z_s(X) = \frac{X - Np_0}{[N\hat{p}(1 - \hat{p})]^{\frac{1}{2}}}$$

where $\hat{p} = X/N$.

For the METHOD=EXACT option, the computations are the same as described in the section “[Exact Test of a Binomial Proportion \(TEST=EXACT\)](#)” on page 7347 except for the definitions of the critical values.

For the upper one-sided case,

$$C_U = \min\{C : Z_s(C) \geq z_{1-\alpha}\}$$

For the lower one-sided case,

$$C_L = \max\{C : Z_s(C) \leq z_\alpha\}$$

For the two-sided case,

$$C_L = \max\{C : Z_s(C) \leq z_{\frac{\alpha}{2}}\}$$

$$C_U = \min\{C : Z_s(C) \geq z_{1-\frac{\alpha}{2}}\}$$

For the METHOD=NORMAL option, the test statistic $Z_s(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{\frac{1}{2}}(p - p_0)}{[p(1 - p)]^{\frac{1}{2}}}, 1\right)$$

(see Chow, Shao, and Wang (2003, p. 82)).

The approximate power is computed as

$$\text{power} = \begin{cases} \Phi\left(z_\alpha + \sqrt{N} \frac{p - p_0}{\sqrt{p(1-p)}}\right), & \text{upper one-sided} \\ \Phi\left(z_\alpha - \sqrt{N} \frac{p - p_0}{\sqrt{p(1-p)}}\right), & \text{lower one-sided} \\ \Phi\left(z_{\frac{\alpha}{2}} + \sqrt{N} \frac{p - p_0}{\sqrt{p(1-p)}}\right) + \Phi\left(z_{\frac{\alpha}{2}} - \sqrt{N} \frac{p - p_0}{\sqrt{p(1-p)}}\right), & \text{two-sided} \end{cases}$$

The approximate sample size is computed in closed form for the one-sided cases by inverting the power equation,

$$N = p(1 - p) \left(\frac{z_{\text{power}} + z_{1-\alpha}}{p - p_0} \right)^2$$

and by numerical inversion for the two-sided case.

z Test for Binomial Proportion with Continuity Adjustment Using Null Variance (TEST=ADJZ VAREST=NULL)

For the normal approximation test with continuity adjustment, the test statistic is (Pagano and Gauvreau 1993, p. 295):

$$Z_c(X) = \frac{X - Np_0 + 0.5(1_{\{X < Np_0\}}) - 0.5(1_{\{X > Np_0\}})}{[Np_0(1 - p_0)]^{\frac{1}{2}}}$$

For the METHOD=EXACT option, the computations are the same as described in the section “[Exact Test of a Binomial Proportion \(TEST=EXACT\)](#)” on page 7347 except for the definitions of the critical values.

For the upper one-sided case,

$$C_U = \min\{C : Z_c(C) \geq z_{1-\alpha}\}$$

For the lower one-sided case,

$$C_L = \max\{C : Z_c(C) \leq z_\alpha\}$$

For the two-sided case,

$$C_L = \max\{C : Z_c(C) \leq z_{\frac{\alpha}{2}}\}$$

$$C_U = \min\{C : Z_c(C) \geq z_{1-\frac{\alpha}{2}}\}$$

For the METHOD=NORMAL option, the test statistic $Z_c(X)$ is assumed to have the normal distribution $N(\mu, \sigma^2)$, where μ and σ^2 are derived as follows.

For convenience of notation, define

$$k = \frac{1}{2\sqrt{Np_0(1-p_0)}}$$

Then

$$E[Z_c(X)] = 2kNp - 2kNp_0 + kP(X < Np_0) - kP(X > Np_0)$$

and

$$\begin{aligned} \text{Var}[Z_c(X)] &= 4k^2Np(1-p) + k^2[1 - P(X = Np_0)] - k^2[P(X < Np_0) - P(X > Np_0)]^2 \\ &\quad + 4k^2[E(X1_{\{X < Np_0\}}) - E(X1_{\{X > Np_0\}})] - 4k^2Np[P(X < Np_0) - P(X > Np_0)] \end{aligned}$$

The probabilities $P(X = Np_0)$, $P(X < Np_0)$, and $P(X > Np_0)$ and the truncated expectations $E(X1_{\{X < Np_0\}})$ and $E(X1_{\{X > Np_0\}})$ are approximated by assuming the normal approximate distribution of X , $N(Np, Np(1-p))$. Letting $\phi(\cdot)$ and $\Phi(\cdot)$ denote the standard normal PDF and CDF, respectively, and defining d as

$$d = \frac{Np_0 - Np}{[Np(1-p)]^{\frac{1}{2}}}$$

the terms are computed as follows:

$$\begin{aligned} P(X = Np_0) &= 0 \\ P(X < Np_0) &= \Phi(d) \\ P(X > Np_0) &= 1 - \Phi(d) \\ E(X1_{\{X < Np_0\}}) &= Np\Phi(d) - [Np(1-p)]^{\frac{1}{2}}\phi(d) \\ E(X1_{\{X > Np_0\}}) &= Np[1 - \Phi(d)] + [Np(1-p)]^{\frac{1}{2}}\phi(d) \end{aligned}$$

The mean and variance of $Z_c(X)$ are thus approximated by

$$\mu = k[2Np - 2Np_0 + 2\Phi(d) - 1]$$

and

$$\sigma^2 = 4k^2 \left[Np(1-p) + \Phi(d)(1-\Phi(d)) - 2(Np(1-p))^{\frac{1}{2}} \phi(d) \right]$$

The approximate power is computed as

$$\text{power} = \begin{cases} \Phi\left(\frac{z_{\alpha} + \mu}{\sigma}\right), & \text{upper one-sided} \\ \Phi\left(\frac{z_{\alpha} - \mu}{\sigma}\right), & \text{lower one-sided} \\ \Phi\left(\frac{z_{\frac{\alpha}{2}} + \mu}{\sigma}\right) + \Phi\left(\frac{z_{\frac{\alpha}{2}} - \mu}{\sigma}\right), & \text{two-sided} \end{cases}$$

The approximate sample size is computed by numerical inversion.

z Test for Binomial Proportion with Continuity Adjustment Using Sample Variance (TEST=ADJZ VAREST=SAMPLE)

For the normal approximation test with continuity adjustment using the sample variance, the test statistic is

$$Z_{cs}(X) = \frac{X - Np_0 + 0.5(1_{\{X < Np_0\}}) - 0.5(1_{\{X > Np_0\}})}{[N\hat{p}(1-\hat{p})]^{\frac{1}{2}}}$$

where $\hat{p} = X/N$.

For the METHOD=EXACT option, the computations are the same as described in the section “[Exact Test of a Binomial Proportion \(TEST=EXACT\)](#)” on page 7347 except for the definitions of the critical values.

For the upper one-sided case,

$$C_U = \min\{C : Z_{cs}(C) \geq z_{1-\alpha}\}$$

For the lower one-sided case,

$$C_L = \max\{C : Z_{cs}(C) \leq z_{\alpha}\}$$

For the two-sided case,

$$\begin{aligned} C_L &= \max\{C : Z_{cs}(C) \leq z_{\frac{\alpha}{2}}\} \\ C_U &= \min\{C : Z_{cs}(C) \geq z_{1-\frac{\alpha}{2}}\} \end{aligned}$$

For the METHOD=NORMAL option, the test statistic $Z_{cs}(X)$ is assumed to have the normal distribution $N(\mu, \sigma^2)$, where μ and σ^2 are derived as follows.

For convenience of notation, define

$$k = \frac{1}{2\sqrt{Np(1-p)}}$$

Then

$$E[Z_{cs}(X)] \approx 2kNp - 2kNp_0 + kP(X < Np_0) - kP(X > Np_0)$$

and

$$\begin{aligned} \text{Var}[Z_{cs}(X)] &\approx 4k^2Np(1-p) + k^2[1 - P(X = Np_0)] - k^2[P(X < Np_0) - P(X > Np_0)]^2 \\ &\quad + 4k^2[E(X1_{\{X < Np_0\}}) - E(X1_{\{X > Np_0\}})] - 4k^2Np[P(X < Np_0) - P(X > Np_0)] \end{aligned}$$

The probabilities $P(X = Np_0)$, $P(X < Np_0)$, and $P(X > Np_0)$ and the truncated expectations $E(X1_{\{X < Np_0\}})$ and $E(X1_{\{X > Np_0\}})$ are approximated by assuming the normal-approximate distribution of X , $N(Np, Np(1-p))$. Letting $\phi(\cdot)$ and $\Phi(\cdot)$ denote the standard normal PDF and CDF, respectively, and defining d as

$$d = \frac{Np_0 - Np}{[Np(1-p)]^{\frac{1}{2}}}$$

the terms are computed as follows:

$$\begin{aligned} P(X = Np_0) &= 0 \\ P(X < Np_0) &= \Phi(d) \\ P(X > Np_0) &= 1 - \Phi(d) \\ E(X1_{\{X < Np_0\}}) &= Np\Phi(d) - [Np(1-p)]^{\frac{1}{2}}\phi(d) \\ E(X1_{\{X > Np_0\}}) &= Np[1 - \Phi(d)] + [Np(1-p)]^{\frac{1}{2}}\phi(d) \end{aligned}$$

The mean and variance of $Z_{cs}(X)$ are thus approximated by

$$\mu = k[2Np - 2Np_0 + 2\Phi(d) - 1]$$

and

$$\sigma^2 = 4k^2 \left[Np(1-p) + \Phi(d)(1 - \Phi(d)) - 2(Np(1-p))^{\frac{1}{2}}\phi(d) \right]$$

The approximate power is computed as

$$\text{power} = \begin{cases} \Phi\left(\frac{z_{\alpha} + \mu}{\sigma}\right), & \text{upper one-sided} \\ \Phi\left(\frac{z_{\alpha} - \mu}{\sigma}\right), & \text{lower one-sided} \\ \Phi\left(\frac{z_{\frac{\alpha}{2}} + \mu}{\sigma}\right) + \Phi\left(\frac{z_{\frac{\alpha}{2}} - \mu}{\sigma}\right), & \text{two-sided} \end{cases}$$

The approximate sample size is computed by numerical inversion.

Exact Equivalence Test of a Binomial Proportion (TEST=EQUIV_EXACT)

The hypotheses for the equivalence test are

$$H_0: p < \theta_L \quad \text{or} \quad p > \theta_U$$

$$H_1: \theta_L \leq p \leq \theta_U$$

where θ_L and θ_U are the lower and upper equivalence bounds, respectively.

The analysis is the two one-sided tests (TOST) procedure as described in Chow, Shao, and Wang (2003) on p. 84, but using exact critical values as on p. 116 instead of normal-based critical values.

Two different hypothesis tests are carried out:

$$H_{a0}: p < \theta_L$$

$$H_{a1}: p \geq \theta_L$$

and

$$H_{b0}: p > \theta_U$$

$$H_{b1}: p \leq \theta_U$$

If H_{a0} is rejected in favor of H_{a1} and H_{b0} is rejected in favor of H_{b1} , then H_0 is rejected in favor of H_1 .

The test statistic for each of the two tests (H_{a0} versus H_{a1} and H_{b0} versus H_{b1}) is

$$X = \text{number of successes} \sim \text{Bin}(N, p)$$

Let C_U denote the critical value of the exact upper one-sided test of H_{a0} versus H_{a1} , and let C_L denote the critical value of the exact lower one-sided test of H_{b0} versus H_{b1} . These critical values are computed in the section “**Exact Test of a Binomial Proportion (TEST=EXACT)**” on page 7347. Both of these tests are rejected if and only if $C_U \leq X \leq C_L$. Thus, the exact power of the equivalence test is

$$\begin{aligned} \text{power} &= P(C_U \leq X \leq C_L) \\ &= P(X \geq C_U) - P(X \geq C_L + 1) \end{aligned}$$

The probabilities are computed using Johnson and Kotz (1970, equation 3.20).

z Equivalence Test for Binomial Proportion Using Null Variance (TEST=EQUIV_Z VAREST=NULL)

The hypotheses for the equivalence test are

$$H_0: p < \theta_L \quad \text{or} \quad p > \theta_U$$

$$H_1: \theta_L \leq p \leq \theta_U$$

where θ_L and θ_U are the lower and upper equivalence bounds, respectively.

The analysis is the two one-sided tests (TOST) procedure as described in Chow, Shao, and Wang (2003) on p. 84, but using the null variance instead of the sample variance.

Two different hypothesis tests are carried out:

$$H_{a0}: p < \theta_L$$

$$H_{a1}: p \geq \theta_L$$

and

$$H_{b0}: p > \theta_U$$

$$H_{b1}: p \leq \theta_U$$

If H_{a0} is rejected in favor of H_{a1} and H_{b0} is rejected in favor of H_{b1} , then H_0 is rejected in favor of H_1 .

The test statistic for the test of H_{a0} versus H_{a1} is

$$Z_L(X) = \frac{X - N\theta_L}{[N\theta_L(1 - \theta_L)]^{\frac{1}{2}}}$$

The test statistic for the test of H_{b0} versus H_{b1} is

$$Z_U(X) = \frac{X - N\theta_U}{[N\theta_U(1 - \theta_U)]^{\frac{1}{2}}}$$

For the METHOD=EXACT option, let C_U denote the critical value of the exact upper one-sided test of H_{a0} versus H_{a1} using $Z_L(X)$. This critical value is computed in the section “z Test for Binomial Proportion Using Null Variance (TEST=Z VAREST=NULL)” on page 7348. Similarly, let C_L denote the critical value of the exact lower one-sided test of H_{b0} versus H_{b1} using $Z_U(X)$. Both of these tests are rejected if and only if $C_U \leq X \leq C_L$. Thus, the exact power of the equivalence test is

$$\begin{aligned} \text{power} &= P(C_U \leq X \leq C_L) \\ &= P(X \geq C_U) - P(X \geq C_L + 1) \end{aligned}$$

The probabilities are computed using Johnson and Kotz (1970, equation 3.20).

For the METHOD=NORMAL option, the test statistic $Z_L(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{\frac{1}{2}}(p - \theta_L)}{[\theta_L(1 - \theta_L)]^{\frac{1}{2}}}, \frac{p(1 - p)}{\theta_L(1 - \theta_L)}\right)$$

and the test statistic $Z_U(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{\frac{1}{2}}(p - \theta_U)}{[\theta_U(1 - \theta_U)]^{\frac{1}{2}}}, \frac{p(1 - p)}{\theta_U(1 - \theta_U)}\right)$$

(see Chow, Shao, and Wang (2003, p. 84)). The approximate power is computed as

$$\text{power} = \Phi\left(\frac{z_\alpha - \sqrt{N} \frac{p - \theta_U}{\sqrt{\theta_U(1 - \theta_U)}}}{\sqrt{\frac{p(1 - p)}{\theta_U(1 - \theta_U)}}}\right) + \Phi\left(\frac{z_\alpha + \sqrt{N} \frac{p - \theta_L}{\sqrt{\theta_L(1 - \theta_L)}}}{\sqrt{\frac{p(1 - p)}{\theta_L(1 - \theta_L)}}}\right) - 1$$

The approximate sample size is computed by numerically inverting the power formula, using the sample size estimate N_0 of Chow, Shao, and Wang (2003, p. 85) as an initial guess:

$$N_0 = p(1 - p) \left(\frac{z_{1-\alpha} + z_{(1+\text{power})/2}}{0.5(\theta_U - \theta_L) - |p - 0.5(\theta_L + \theta_U)|} \right)^2$$

z Equivalence Test for Binomial Proportion Using Sample Variance (TEST=EQUIV_Z VAREST=SAMPLE)

The hypotheses for the equivalence test are

$$H_0: p < \theta_L \quad \text{or} \quad p > \theta_U$$

$$H_1: \theta_L \leq p \leq \theta_U$$

where θ_L and θ_U are the lower and upper equivalence bounds, respectively.

The analysis is the two one-sided tests (TOST) procedure as described in Chow, Shao, and Wang (2003) on p. 84.

Two different hypothesis tests are carried out:

$$H_{a0}: p < \theta_L$$

$$H_{a1}: p \geq \theta_L$$

and

$$H_{b0}: p > \theta_U$$

$$H_{b1}: p \leq \theta_U$$

If H_{a0} is rejected in favor of H_{a1} and H_{b0} is rejected in favor of H_{b1} , then H_0 is rejected in favor of H_1 .

The test statistic for the test of H_{a0} versus H_{a1} is

$$Z_{sL}(X) = \frac{X - N\theta_L}{[N\hat{p}(1 - \hat{p})]^{1/2}}$$

where $\hat{p} = X/N$.

The test statistic for the test of H_{b0} versus H_{b1} is

$$Z_{sU}(X) = \frac{X - N\theta_U}{[N\hat{p}(1 - \hat{p})]^{1/2}}$$

For the METHOD=EXACT option, let C_U denote the critical value of the exact upper one-sided test of H_{a0} versus H_{a1} using $Z_{sL}(X)$. This critical value is computed in the section “z Test for Binomial Proportion Using Sample Variance (TEST=Z VAREST=SAMPLE)” on page 7349. Similarly, let C_L denote the critical value of the exact lower one-sided test of H_{b0} versus H_{b1} using $Z_{sU}(X)$. Both of these tests are rejected if and only if $C_U \leq X \leq C_L$. Thus, the exact power of the equivalence test is

$$\begin{aligned} \text{power} &= P(C_U \leq X \leq C_L) \\ &= P(X \geq C_U) - P(X \geq C_L + 1) \end{aligned}$$

The probabilities are computed using Johnson and Kotz (1970, equation 3.20).

For the METHOD=NORMAL option, the test statistic $Z_{sL}(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{1/2}(p - \theta_L)}{[p(1 - p)]^{1/2}}, 1\right)$$

and the test statistic $Z_{sU}(X)$ is assumed to have the normal distribution

$$N\left(\frac{N^{\frac{1}{2}}(p - \theta_U)}{[p(1-p)]^{\frac{1}{2}}}, 1\right)$$

(see Chow, Shao, and Wang (2003), p. 84).

The approximate power is computed as

$$\text{power} = \Phi\left(z_\alpha - \sqrt{N} \frac{p - \theta_U}{\sqrt{p(1-p)}}\right) + \Phi\left(z_\alpha + \sqrt{N} \frac{p - \theta_L}{\sqrt{p(1-p)}}\right) - 1$$

The approximate sample size is computed by numerically inverting the power formula, using the sample size estimate N_0 of Chow, Shao, and Wang (2003, p. 85) as an initial guess:

$$N_0 = p(1-p) \left(\frac{z_{1-\alpha} + z_{(1+\text{power})/2}}{0.5(\theta_U - \theta_L) - |p - 0.5(\theta_L + \theta_U)|} \right)^2$$

z Equivalence Test for Binomial Proportion with Continuity Adjustment Using Null Variance (TEST=EQUIV_ADJZ VAREST=NULL)

The hypotheses for the equivalence test are

$$H_0: p < \theta_L \quad \text{or} \quad p > \theta_U$$

$$H_1: \theta_L \leq p \leq \theta_U$$

where θ_L and θ_U are the lower and upper equivalence bounds, respectively.

The analysis is the two one-sided tests (TOST) procedure as described in Chow, Shao, and Wang (2003) on p. 84, but using the null variance instead of the sample variance.

Two different hypothesis tests are carried out:

$$H_{a0}: p < \theta_L$$

$$H_{a1}: p \geq \theta_L$$

and

$$H_{b0}: p > \theta_U$$

$$H_{b1}: p \leq \theta_U$$

If H_{a0} is rejected in favor of H_{a1} and H_{b0} is rejected in favor of H_{b1} , then H_0 is rejected in favor of H_1 .

The test statistic for the test of H_{a0} versus H_{a1} is

$$Z_{cL}(X) = \frac{X - N\theta_L + 0.5(1_{\{X < N\theta_L\}}) - 0.5(1_{\{X > N\theta_L\}})}{\left[N\hat{\theta}_L(1 - \hat{\theta}_L)\right]^{\frac{1}{2}}}$$

where $\hat{p} = X/N$.

The test statistic for the test of H_{b0} versus H_{b1} is

$$Z_{cU}(X) = \frac{X - N\theta_U + 0.5(1_{\{X < N\theta_U\}}) - 0.5(1_{\{X > N\theta_U\}})}{\left[N\hat{\theta}_U(1 - \hat{\theta}_U)\right]^{\frac{1}{2}}}$$

For the METHOD=EXACT option, let C_U denote the critical value of the exact upper one-sided test of H_{a0} versus H_{a1} using $Z_{cL}(X)$. This critical value is computed in the section “z Test for Binomial Proportion with Continuity Adjustment Using Null Variance (TEST=ADJZ VAREST=NULL)” on page 7350. Similarly, let C_L denote the critical value of the exact lower one-sided test of H_{b0} versus H_{b1} using $Z_{cU}(X)$. Both of these tests are rejected if and only if $C_U \leq X \leq C_L$. Thus, the exact power of the equivalence test is

$$\begin{aligned} \text{power} &= P(C_U \leq X \leq C_L) \\ &= P(X \geq C_U) - P(X \geq C_L + 1) \end{aligned}$$

The probabilities are computed using Johnson and Kotz (1970, equation 3.20).

For the METHOD=NORMAL option, the test statistic $Z_{cL}(X)$ is assumed to have the normal distribution $N(\mu_L, \sigma_L^2)$, and $Z_{cU}(X)$ is assumed to have the normal distribution $N(\mu_U, \sigma_U^2)$, where μ_L , μ_U , σ_L^2 , and σ_U^2 are derived as follows.

For convenience of notation, define

$$\begin{aligned} k_L &= \frac{1}{2\sqrt{N\theta_L(1 - \theta_L)}} \\ k_U &= \frac{1}{2\sqrt{N\theta_U(1 - \theta_U)}} \end{aligned}$$

Then

$$\begin{aligned} E[Z_{cL}(X)] &\approx 2k_L Np - 2k_L N\theta_L + k_L P(X < N\theta_L) - k_L P(X > N\theta_L) \\ E[Z_{cU}(X)] &\approx 2k_U Np - 2k_U N\theta_U + k_U P(X < N\theta_U) - k_U P(X > N\theta_U) \end{aligned}$$

and

$$\begin{aligned} \text{Var}[Z_{cL}(X)] &\approx 4k_L^2 Np(1 - p) + k_L^2 [1 - P(X = N\theta_L)] - k_L^2 [P(X < N\theta_L) - P(X > N\theta_L)]^2 \\ &\quad + 4k_L^2 [E(X1_{\{X < N\theta_L\}}) - E(X1_{\{X > N\theta_L\}})] - 4k_L^2 Np [P(X < N\theta_L) - P(X > N\theta_L)] \\ \text{Var}[Z_{cU}(X)] &\approx 4k_U^2 Np(1 - p) + k_U^2 [1 - P(X = N\theta_U)] - k_U^2 [P(X < N\theta_U) - P(X > N\theta_U)]^2 \\ &\quad + 4k_U^2 [E(X1_{\{X < N\theta_U\}}) - E(X1_{\{X > N\theta_U\}})] - 4k_U^2 Np [P(X < N\theta_U) - P(X > N\theta_U)] \end{aligned}$$

The probabilities $P(X = N\theta_L)$, $P(X < N\theta_L)$, $P(X > N\theta_L)$, $P(X = N\theta_U)$, $P(X < N\theta_U)$, and $P(X > N\theta_U)$ and the truncated expectations $E(X1_{\{X < N\theta_L\}})$, $E(X1_{\{X > N\theta_L\}})$, $E(X1_{\{X < N\theta_U\}})$, and $E(X1_{\{X > N\theta_U\}})$ are approximated by assuming the normal-approximate distribution of X , $N(Np, Np(1 - p))$.

Letting $\phi(\cdot)$ and $\Phi(\cdot)$ denote the standard normal PDF and CDF, respectively, and defining d_L and d_U as

$$d_L = \frac{N\theta_L - Np}{[Np(1-p)]^{\frac{1}{2}}}$$

$$d_U = \frac{N\theta_U - Np}{[Np(1-p)]^{\frac{1}{2}}}$$

the terms are computed as follows:

$$\begin{aligned} P(X = N\theta_L) &= 0 \\ P(X = N\theta_U) &= 0 \\ P(X < N\theta_L) &= \Phi(d_L) \\ P(X < N\theta_U) &= \Phi(d_U) \\ P(X > N\theta_L) &= 1 - \Phi(d_L) \\ P(X > N\theta_U) &= 1 - \Phi(d_U) \\ E(X1_{\{X < N\theta_L\}}) &= Np\Phi(d_L) - [Np(1-p)]^{\frac{1}{2}}\phi(d_L) \\ E(X1_{\{X < N\theta_U\}}) &= Np\Phi(d_U) - [Np(1-p)]^{\frac{1}{2}}\phi(d_U) \\ E(X1_{\{X > N\theta_L\}}) &= Np[1 - \Phi(d_L)] + [Np(1-p)]^{\frac{1}{2}}\phi(d_L) \\ E(X1_{\{X > N\theta_U\}}) &= Np[1 - \Phi(d_U)] + [Np(1-p)]^{\frac{1}{2}}\phi(d_U) \end{aligned}$$

The mean and variance of $Z_{cL}(X)$ and $Z_{cU}(X)$ are thus approximated by

$$\begin{aligned} \mu_L &= k_L [2Np - 2N\theta_L + 2\Phi(d_L) - 1] \\ \mu_U &= k_U [2Np - 2N\theta_U + 2\Phi(d_U) - 1] \end{aligned}$$

and

$$\begin{aligned} \sigma_L^2 &= 4k_L^2 \left[Np(1-p) + \Phi(d_L)(1 - \Phi(d_L)) - 2(Np(1-p))^{\frac{1}{2}}\phi(d_L) \right] \\ \sigma_U^2 &= 4k_U^2 \left[Np(1-p) + \Phi(d_U)(1 - \Phi(d_U)) - 2(Np(1-p))^{\frac{1}{2}}\phi(d_U) \right] \end{aligned}$$

The approximate power is computed as

$$\text{power} = \Phi\left(\frac{z_\alpha - \mu_U}{\sigma_U}\right) + \Phi\left(\frac{z_\alpha + \mu_L}{\sigma_L}\right) - 1$$

The approximate sample size is computed by numerically inverting the power formula.

z Equivalence Test for Binomial Proportion with Continuity Adjustment Using Sample Variance
(TEST=EQUIV_ADJZ VAREST=SAMPLE)

The hypotheses for the equivalence test are

$$\begin{aligned} H_0: p < \theta_L \quad \text{or} \quad p > \theta_U \\ H_1: \theta_L \leq p \leq \theta_U \end{aligned}$$

where θ_L and θ_U are the lower and upper equivalence bounds, respectively.

The analysis is the two one-sided tests (TOST) procedure as described in Chow, Shao, and Wang (2003) on p. 84.

Two different hypothesis tests are carried out:

$$H_{a0}: p < \theta_L$$

$$H_{a1}: p \geq \theta_L$$

and

$$H_{b0}: p > \theta_U$$

$$H_{b1}: p \leq \theta_U$$

If H_{a0} is rejected in favor of H_{a1} and H_{b0} is rejected in favor of H_{b1} , then H_0 is rejected in favor of H_1 .

The test statistic for the test of H_{a0} versus H_{a1} is

$$Z_{csL}(X) = \frac{X - N\theta_L + 0.5(1_{\{X < N\theta_L\}}) - 0.5(1_{\{X > N\theta_L\}})}{[N\hat{p}(1 - \hat{p})]^{\frac{1}{2}}}$$

where $\hat{p} = X/N$.

The test statistic for the test of H_{b0} versus H_{b1} is

$$Z_{csU}(X) = \frac{X - N\theta_U + 0.5(1_{\{X < N\theta_U\}}) - 0.5(1_{\{X > N\theta_U\}})}{[N\hat{p}(1 - \hat{p})]^{\frac{1}{2}}}$$

For the METHOD=EXACT option, let C_U denote the critical value of the exact upper one-sided test of H_{a0} versus H_{a1} using $Z_{csL}(X)$. This critical value is computed in the section “[z Test for Binomial Proportion with Continuity Adjustment Using Sample Variance \(TEST=ADJZ VAREST=SAMPLE\)](#)” on page 7352. Similarly, let C_L denote the critical value of the exact lower one-sided test of H_{b0} versus H_{b1} using $Z_{csU}(X)$. Both of these tests are rejected if and only if $C_U \leq X \leq C_L$. Thus, the exact power of the equivalence test is

$$\begin{aligned} \text{power} &= P(C_U \leq X \leq C_L) \\ &= P(X \geq C_U) - P(X \geq C_L + 1) \end{aligned}$$

The probabilities are computed using Johnson and Kotz (1970, equation 3.20).

For the METHOD=NORMAL option, the test statistic $Z_{csL}(X)$ is assumed to have the normal distribution $N(\mu_L, \sigma_L^2)$, and $Z_{csU}(X)$ is assumed to have the normal distribution $N(\mu_U, \sigma_U^2)$, where μ_L , μ_U , σ_L^2 and σ_U^2 are derived as follows.

For convenience of notation, define

$$k = \frac{1}{2\sqrt{Np(1-p)}}$$

Then

$$\begin{aligned} E[Z_{csL}(X)] &\approx 2kNp - 2kN\theta_L + kP(X < N\theta_L) - kP(X > N\theta_L) \\ E[Z_{csU}(X)] &\approx 2kNp - 2kN\theta_U + kP(X < N\theta_U) - kP(X > N\theta_U) \end{aligned}$$

and

$$\begin{aligned} \text{Var}[Z_{csL}(X)] &\approx 4k^2Np(1-p) + k^2[1 - P(X = N\theta_L)] - k^2[P(X < N\theta_L) - P(X > N\theta_L)]^2 \\ &\quad + 4k^2[E(X1_{\{X < N\theta_L\}}) - E(X1_{\{X > N\theta_L\}})] - 4k^2Np[P(X < N\theta_L) - P(X > N\theta_L)] \\ \text{Var}[Z_{csU}(X)] &\approx 4k^2Np(1-p) + k^2[1 - P(X = N\theta_U)] - k^2[P(X < N\theta_U) - P(X > N\theta_U)]^2 \\ &\quad + 4k^2[E(X1_{\{X < N\theta_U\}}) - E(X1_{\{X > N\theta_U\}})] - 4k^2Np[P(X < N\theta_U) - P(X > N\theta_U)] \end{aligned}$$

The probabilities $P(X = N\theta_L)$, $P(X < N\theta_L)$, $P(X > N\theta_L)$, $P(X = N\theta_U)$, $P(X < N\theta_U)$, and $P(X > N\theta_U)$ and the truncated expectations $E(X1_{\{X < N\theta_L\}})$, $E(X1_{\{X > N\theta_L\}})$, $E(X1_{\{X < N\theta_U\}})$, and $E(X1_{\{X > N\theta_U\}})$ are approximated by assuming the normal-approximate distribution of X , $N(Np, Np(1-p))$. Letting $\phi(\cdot)$ and $\Phi(\cdot)$ denote the standard normal PDF and CDF, respectively, and defining d_L and d_U as

$$\begin{aligned} d_L &= \frac{N\theta_L - Np}{[Np(1-p)]^{\frac{1}{2}}} \\ d_U &= \frac{N\theta_U - Np}{[Np(1-p)]^{\frac{1}{2}}} \end{aligned}$$

the terms are computed as follows:

$$\begin{aligned} P(X = N\theta_L) &= 0 \\ P(X = N\theta_U) &= 0 \\ P(X < N\theta_L) &= \Phi(d_L) \\ P(X < N\theta_U) &= \Phi(d_U) \\ P(X > N\theta_L) &= 1 - \Phi(d_L) \\ P(X > N\theta_U) &= 1 - \Phi(d_U) \\ E(X1_{\{X < N\theta_L\}}) &= Np\Phi(d_L) - [Np(1-p)]^{\frac{1}{2}}\phi(d_L) \\ E(X1_{\{X < N\theta_U\}}) &= Np\Phi(d_U) - [Np(1-p)]^{\frac{1}{2}}\phi(d_U) \\ E(X1_{\{X > N\theta_L\}}) &= Np[1 - \Phi(d_L)] + [Np(1-p)]^{\frac{1}{2}}\phi(d_L) \\ E(X1_{\{X > N\theta_U\}}) &= Np[1 - \Phi(d_U)] + [Np(1-p)]^{\frac{1}{2}}\phi(d_U) \end{aligned}$$

The mean and variance of $Z_{csL}(X)$ and $Z_{csU}(X)$ are thus approximated by

$$\begin{aligned} \mu_L &= k[2Np - 2N\theta_L + 2\Phi(d_L) - 1] \\ \mu_U &= k[2Np - 2N\theta_U + 2\Phi(d_U) - 1] \end{aligned}$$

and

$$\begin{aligned}\sigma_L^2 &= 4k^2 \left[Np(1-p) + \Phi(d_L)(1 - \Phi(d_L)) - 2(Np(1-p))^{\frac{1}{2}} \phi(d_L) \right] \\ \sigma_U^2 &= 4k^2 \left[Np(1-p) + \Phi(d_U)(1 - \Phi(d_U)) - 2(Np(1-p))^{\frac{1}{2}} \phi(d_U) \right]\end{aligned}$$

The approximate power is computed as

$$\text{power} = \Phi\left(\frac{z_\alpha - \mu_U}{\sigma_U}\right) + \Phi\left(\frac{z_\alpha + \mu_L}{\sigma_L}\right) - 1$$

The approximate sample size is computed by numerically inverting the power formula.

Wilson Score Confidence Interval for Binomial Proportion (CI=WILSON)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$\frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \pm \frac{z_{1-\alpha/2} N^{\frac{1}{2}}}{N + z_{1-\alpha/2}^2} \left(\hat{p}(1 - \hat{p}) + \frac{z_{1-\alpha/2}^2}{4N} \right)^{\frac{1}{2}}$$

So the half-width for the two-sided $100(1 - \alpha)\%$ confidence interval is

$$\text{half-width} = \frac{z_{1-\alpha/2} N^{\frac{1}{2}}}{N + z_{1-\alpha/2}^2} \left(\hat{p}(1 - \hat{p}) + \frac{z_{1-\alpha/2}^2}{4N} \right)^{\frac{1}{2}}$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

For references and more details about this and all other confidence intervals associated with the **CI=** option, see “Binomial Proportion” on page 2774 in Chapter 41, “The FREQ Procedure.”

Agresti-Coull “Add k Successes and Failures” Confidence Interval for Binomial Proportion (CI=AGRESTICOULL)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$\frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \pm z_{1-\alpha/2} \left(\frac{\frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \left(1 - \frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \right)}{N + z_{1-\alpha/2}^2} \right)^{\frac{1}{2}}$$

So the half-width for the two-sided $100(1 - \alpha)\%$ confidence interval is

$$\text{half-width} = z_{1-\alpha/2} \left(\frac{\frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \left(1 - \frac{X + \frac{z_{1-\alpha/2}^2}{2}}{N + z_{1-\alpha/2}^2} \right)}{N + z_{1-\alpha/2}^2} \right)^{\frac{1}{2}}$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

Jeffreys Confidence Interval for Binomial Proportion (CI=JEFFREYS)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$[L_J(X), U_J(X)]$$

where

$$L_J(X) = \begin{cases} 0, & X = 0 \\ \text{Beta}_{\alpha/2; X+1/2, N-X+1/2}, & X > 0 \end{cases}$$

and

$$U_J(X) = \begin{cases} \text{Beta}_{1-\alpha/2; X+1/2, N-X+1/2}, & X < N \\ 1, & X = N \end{cases}$$

The half-width of this two-sided $100(1 - \alpha)\%$ confidence interval is defined as half the width of the full interval:

$$\text{half-width} = \frac{1}{2} (U_J(X) - L_J(X))$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

Exact Clopper-Pearson Confidence Interval for Binomial Proportion (CI=EXACT)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$[L_E(X), U_E(X)]$$

where

$$L_E(X) = \begin{cases} 0, & X = 0 \\ \text{Beta}_{\alpha/2; X, N-X+1}, & X > 0 \end{cases}$$

and

$$U_E(X) = \begin{cases} \text{Beta}_{1-\alpha/2; X+1, N-X}, & X < N \\ 1, & X = N \end{cases}$$

The half-width of this two-sided $100(1 - \alpha)\%$ confidence interval is defined as half the width of the full interval:

$$\text{half-width} = \frac{1}{2} (U_E(X) - L_E(X))$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

Wald Confidence Interval for Binomial Proportion (CI=WALD)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$\hat{p} \pm z_{1-\alpha/2} \left(\frac{\hat{p}(1 - \hat{p})}{N} \right)^{\frac{1}{2}}$$

So the half-width for the two-sided $100(1 - \alpha)\%$ confidence interval is

$$\text{half-width} = z_{1-\alpha/2} \left(\frac{\hat{p}(1 - \hat{p})}{N} \right)^{\frac{1}{2}}$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

Continuity-Corrected Wald Confidence Interval for Binomial Proportion (CI=WALD_CORRECT)

The two-sided $100(1 - \alpha)\%$ confidence interval for p is

$$\hat{p} \pm \left[z_{1-\alpha/2} \left(\frac{\hat{p}(1-\hat{p})}{N} \right)^{\frac{1}{2}} + \frac{1}{2N} \right]$$

So the half-width for the two-sided $100(1 - \alpha)\%$ confidence interval is

$$\text{half-width} = z_{1-\alpha/2} \left(\frac{\hat{p}(1-\hat{p})}{N} \right)^{\frac{1}{2}} + \frac{1}{2N}$$

Prob(Width) is calculated exactly by adding up the probabilities of observing each $X \in \{1, \dots, N\}$ that produces a confidence interval whose half-width is at most a target value h :

$$\text{Prob(Width)} = \sum_{i=0}^N P(X = i) 1_{\text{half-width} < h}$$

Analyses in the ONESAMPLEMEANS Statement**One-Sample t Test (TEST=T)**

The hypotheses for the one-sample t test are

$$H_0: \mu = \mu_0$$

$$H_1: \begin{cases} \mu \neq \mu_0, & \text{two-sided} \\ \mu > \mu_0, & \text{upper one-sided} \\ \mu < \mu_0, & \text{lower one-sided} \end{cases}$$

The test assumes normally distributed data and requires $N \geq 2$. The test statistics are

$$t = N^{\frac{1}{2}} \left(\frac{\bar{x} - \mu_0}{s} \right) \sim t(N-1, \delta)$$

$$t^2 \sim F(1, N-1, \delta^2)$$

where $\bar{x}$ is the sample mean, s is the sample standard deviation, and

$$\delta = N^{\frac{1}{2}} \left(\frac{\mu - \mu_0}{\sigma} \right)$$

The test is

$$\text{Reject } H_0 \text{ if } \begin{cases} t^2 \geq F_{1-\alpha}(1, N-1), & \text{two-sided} \\ t \geq t_{1-\alpha}(N-1), & \text{upper one-sided} \\ t \leq t_{\alpha}(N-1), & \text{lower one-sided} \end{cases}$$

Exact power computations for t tests are discussed in O'Brien and Muller (1993, Section 8.2), although not specifically for the one-sample case. The power is based on the noncentral t and F distributions:

$$\text{power} = \begin{cases} P(F(1, N-1, \delta^2) \geq F_{1-\alpha}(1, N-1)), & \text{two-sided} \\ P(t(N-1, \delta) \geq t_{1-\alpha}(N-1)), & \text{upper one-sided} \\ P(t(N-1, \delta) \leq t_{\alpha}(N-1)), & \text{lower one-sided} \end{cases}$$

Solutions for N , α , and δ are obtained by numerically inverting the power equation. Closed-form solutions for other parameters, in terms of δ , are as follows:

$$\begin{aligned} \mu &= \delta \sigma N^{-\frac{1}{2}} + \mu_0 \\ \sigma &= \begin{cases} \delta^{-1} N^{\frac{1}{2}} (\mu - \mu_0), & |\delta| > 0 \\ \text{undefined}, & \text{otherwise} \end{cases} \end{aligned}$$

One-Sample t Test with Lognormal Data (TEST=T DIST=LOGNORMAL)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section “One-Sample t Test (TEST=T)” on page 7365 then apply.

In contrast to the usual t test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means. This is because the transformation of a null arithmetic mean of lognormal data to the normal scale depends on the unknown coefficient of variation, resulting in an ill-defined hypothesis on the log-transformed data. Geometric means transform cleanly and are more natural for lognormal data.

The hypotheses for the one-sample t test with lognormal data are

$$\begin{aligned} H_0: \frac{\gamma}{\gamma_0} &= 1 \\ H_1: \begin{cases} \frac{\gamma}{\gamma_0} \neq 1, & \text{two-sided} \\ \frac{\gamma}{\gamma_0} > 1, & \text{upper one-sided} \\ \frac{\gamma}{\gamma_0} < 1, & \text{lower one-sided} \end{cases} \end{aligned}$$

Let μ^* and σ^* be the (arithmetic) mean and standard deviation of the normal distribution of the log-transformed data. The hypotheses can be rewritten as follows:

$$\begin{aligned} H_0: \mu^* &= \log(\gamma_0) \\ H_1: \begin{cases} \mu^* \neq \log(\gamma_0), & \text{two-sided} \\ \mu^* > \log(\gamma_0), & \text{upper one-sided} \\ \mu^* < \log(\gamma_0), & \text{lower one-sided} \end{cases} \end{aligned}$$

where $\mu^* = \log(\gamma)$.

The test assumes lognormally distributed data and requires $N \geq 2$.

The power is

$$\text{power} = \begin{cases} P(F(1, N-1, \delta^2) \geq F_{1-\alpha}(1, N-1)), & \text{two-sided} \\ P(t(N-1, \delta) \geq t_{1-\alpha}(N-1)), & \text{upper one-sided} \\ P(t(N-1, \delta) \leq t_{\alpha}(N-1)), & \text{lower one-sided} \end{cases}$$

where

$$\delta = N^{\frac{1}{2}} \left(\frac{\mu^* - \log(\gamma_0)}{\sigma^*} \right)$$

$$\sigma^* = [\log(\text{CV}^2 + 1)]^{\frac{1}{2}}$$

Equivalence Test for Mean of Normal Data (TEST=EQUIV DIST=NORMAL)

The hypotheses for the equivalence test are

$$H_0: \mu < \theta_L \quad \text{or} \quad \mu > \theta_U$$

$$H_1: \theta_L \leq \mu \leq \theta_U$$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987). The test assumes normally distributed data and requires $N \geq 2$. Phillips (1990) derives an expression for the exact power assuming a two-sample balanced design; the results are easily adapted to a one-sample design:

$$\text{power} = Q_{N-1} \left((-t_{1-\alpha}(N-1)), \frac{\mu - \theta_U}{\sigma N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right) -$$

$$Q_{N-1} \left((t_{1-\alpha}(N-1)), \frac{\mu - \theta_L}{\sigma N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right)$$

where $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section “[Common Notation](#)” on page 7336.

Equivalence Test for Mean of Lognormal Data (TEST=EQUIV DIST=LOGNORMAL)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section “[Equivalence Test for Mean of Normal Data \(TEST=EQUIV DIST=NORMAL\)](#)” on page 7367 then apply.

In contrast to the additive equivalence test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means. This is because the transformation of an arithmetic mean of lognormal data to the normal scale depends on the unknown coefficient of variation, resulting in an ill-defined hypothesis on the log-transformed data. Geometric means transform cleanly and are more natural for lognormal data.

The hypotheses for the equivalence test are

$$H_0: \gamma \leq \theta_L \quad \text{or} \quad \gamma \geq \theta_U$$

$$H_1: \theta_L < \gamma < \theta_U$$

where $0 < \theta_L < \theta_U$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987) on the log-transformed data. The test assumes lognormally distributed data and requires $N \geq 2$. Diletti, Hauschke, and Steinijans (1991) derive an expression for the exact power assuming a crossover design; the results are easily adapted to a one-sample design:

$$\text{power} = Q_{N-1} \left((-t_{1-\alpha}(N-1)), \frac{\log(\gamma) - \log(\theta_U)}{\sigma^* N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right) - \\ Q_{N-1} \left((t_{1-\alpha}(N-1)), \frac{\log(\gamma) - \log(\theta_L)}{\sigma^* N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right)$$

where

$$\sigma^* = [\log(\text{CV}^2 + 1)]^{\frac{1}{2}}$$

is the standard deviation of the log-transformed data, and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section "Common Notation" on page 7336.

Confidence Interval for Mean (CI=T)

This analysis of precision applies to the standard t -based confidence interval:

$$\begin{cases} \left[\bar{x} - t_{1-\frac{\alpha}{2}}(N-1) \frac{s}{\sqrt{N}}, \bar{x} + t_{1-\frac{\alpha}{2}}(N-1) \frac{s}{\sqrt{N}} \right], & \text{two-sided} \\ \left[\bar{x} - t_{1-\alpha}(N-1) \frac{s}{\sqrt{N}}, \infty \right), & \text{upper one-sided} \\ \left(-\infty, \bar{x} + t_{1-\alpha}(N-1) \frac{s}{\sqrt{N}} \right], & \text{lower one-sided} \end{cases}$$

where $\bar{x}$ is the sample mean and s is the sample standard deviation. The "half-width" is defined as the distance from the point estimate $\bar{x}$ to a finite endpoint,

$$\text{half-width} = \begin{cases} t_{1-\frac{\alpha}{2}}(N-1) \frac{s}{\sqrt{N}}, & \text{two-sided} \\ t_{1-\alpha}(N-1) \frac{s}{\sqrt{N}}, & \text{one-sided} \end{cases}$$

A "valid" confidence interval captures the true mean. The exact probability of obtaining at most the target confidence interval half-width h , unconditional or conditional on validity, is given by Beal (1989):

$$\begin{aligned} \Pr(\text{half-width} \leq h) &= \begin{cases} P \left(\chi^2(N-1) \leq \frac{h^2 N(N-1)}{\sigma^2(t_{1-\frac{\alpha}{2}}^2(N-1))} \right), & \text{two-sided} \\ P \left(\chi^2(N-1) \leq \frac{h^2 N(N-1)}{\sigma^2(t_{1-\alpha}^2(N-1))} \right), & \text{one-sided} \end{cases} \\ \Pr(\text{half-width} \leq h | \text{validity}) &= \begin{cases} \left(\frac{1}{1-\alpha} \right) 2 \left[Q_{N-1} \left((t_{1-\frac{\alpha}{2}}(N-1)), 0; 0, b_1 \right) - Q_{N-1}(0, 0; 0, b_1) \right], & \text{two-sided} \\ \left(\frac{1}{1-\alpha} \right) Q_{N-1} \left((t_{1-\alpha}(N-1)), 0; 0, b_1 \right), & \text{one-sided} \end{cases} \end{aligned}$$

where

$$b_1 = \frac{h(N-1)^{\frac{1}{2}}}{\sigma(t_{1-\frac{\alpha}{c}}(N-1))N^{-\frac{1}{2}}}$$

c = number of sides

and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section "Common Notation" on page 7336.

A "quality" confidence interval is both sufficiently narrow (half-width $\leq h$) and valid:

$$\begin{aligned}\Pr(\text{quality}) &= \Pr(\text{half-width} \leq h \text{ and validity}) \\ &= \Pr(\text{half-width} \leq h | \text{validity})(1 - \alpha)\end{aligned}$$

Analyses in the ONEWAYANOVA Statement

One-Degree-of-Freedom Contrast (TEST=CONTRAST)

The hypotheses are

$$\begin{aligned}H_0: c_1\mu_1 + \cdots + c_G\mu_G &= c_0 \\ H_1: \begin{cases} c_1\mu_1 + \cdots + c_G\mu_G \neq c_0, & \text{two-sided} \\ c_1\mu_1 + \cdots + c_G\mu_G > c_0, & \text{upper one-sided} \\ c_1\mu_1 + \cdots + c_G\mu_G < c_0, & \text{lower one-sided} \end{cases}\end{aligned}$$

where G is the number of groups, $\{c_1, \dots, c_G\}$ are the contrast coefficients, and c_0 is the null contrast value.

The test is the usual F test for a contrast in one-way ANOVA. It assumes normal data with common group variances and requires $N \geq G + 1$ and $n_i \geq 1$.

O'Brien and Muller (1993, Section 8.2.3.2) give the exact power as

$$\text{power} = \begin{cases} P(F(1, N - G, \delta^2) \geq F_{1-\alpha}(1, N - G)), & \text{two-sided} \\ P(t(N - G, \delta) \geq t_{1-\alpha}(N - G)), & \text{upper one-sided} \\ P(t(N - G, \delta) \leq t_{\alpha}(N - G)), & \text{lower one-sided} \end{cases}$$

where

$$\delta = N^{\frac{1}{2}} \left(\frac{\sum_{i=1}^G c_i \mu_i - c_0}{\sigma \left(\sum_{i=1}^G \frac{c_i^2}{w_i} \right)^{\frac{1}{2}}} \right)$$

Overall F Test (TEST=OVERALL)

The hypotheses are

$$\begin{aligned}H_0: \mu_1 &= \mu_2 = \cdots = \mu_G \\ H_1: \mu_i &\neq \mu_j \text{ for some } i, j\end{aligned}$$

where G is the number of groups.

The test is the usual overall F test for equality of means in one-way ANOVA. It assumes normal data with common group variances and requires $N \geq G + 1$ and $n_i \geq 1$.

O'Brien and Muller (1993, Section 8.2.3.1) give the exact power as

$$\text{power} = P(F(G - 1, N - G, \lambda) \geq F_{1-\alpha}(G - 1, N - G))$$

where the noncentrality is

$$\lambda = N \left(\frac{\sum_{i=1}^G w_i (\mu_i - \bar{\mu})^2}{\sigma^2} \right)$$

and

$$\bar{\mu} = \sum_{i=1}^G w_i \mu_i$$

Analyses in the PAIREDFREQ Statement

Overview of Conditional McNemar tests

Notation:

		Case		
		Failure	Success	
Control	Failure	n_{00}	n_{01}	$n_{0\cdot}$
	Success	n_{10}	n_{11}	$n_{1\cdot}$
		$n_{\cdot 0}$	$n_{\cdot 1}$	N

$$n_{00} = \#\{\text{control=failure, case=failure}\}$$

$$n_{01} = \#\{\text{control=failure, case=success}\}$$

$$n_{10} = \#\{\text{control=success, case=failure}\}$$

$$n_{11} = \#\{\text{control=success, case=success}\}$$

$$N = n_{00} + n_{01} + n_{10} + n_{11}$$

$$n_D = n_{01} + n_{10} \equiv \#\text{discordant pairs}$$

$$\hat{\pi}_{ij} = \frac{n_{ij}}{N}$$

$$\pi_{ij} = \text{theoretical population value of } \hat{\pi}_{ij}$$

$$\pi_{1\cdot} = \pi_{10} + \pi_{11}$$

$$\pi_{\cdot 1} = \pi_{01} + \pi_{11}$$

$$\phi = \text{Corr}(\text{control observation, case observation}) \quad (\text{within a pair})$$

$$\text{DPR} = \text{"discordant proportion ratio"} = \frac{\pi_{01}}{\pi_{10}}$$

$$\text{DPR}_0 = \text{null DPR}$$

Power formulas are given here in terms of the discordant proportions π_{10} and π_{01} . If the input is specified in terms of $\{\pi_{1\cdot}, \pi_{\cdot 1}, \phi\}$, then it can be converted into values for $\{\pi_{10}, \pi_{01}\}$ as follows:

$$\pi_{01} = \pi_{\cdot 1}(1 - \pi_{1\cdot}) - \phi((1 - \pi_{1\cdot})\pi_{1\cdot}(1 - \pi_{\cdot 1})\pi_{\cdot 1})^{\frac{1}{2}}$$

$$\pi_{10} = \pi_{01} + \pi_{1\cdot} - \pi_{\cdot 1}$$

All McNemar tests covered in PROC POWER are *conditional*, meaning that n_D is assumed fixed at its observed value.

For the usual $\text{DPR}_0 = 1$, the hypotheses are

$$H_0: \pi_{\cdot 1} = \pi_1.$$

$$H_1: \begin{cases} \pi_{\cdot 1} \neq \pi_1, & \text{two-sided} \\ \pi_{\cdot 1} > \pi_1, & \text{upper one-sided} \\ \pi_{\cdot 1} < \pi_1, & \text{lower one-sided} \end{cases}$$

The test statistic for both tests covered in PROC POWER ($\text{DIST}=\text{EXACT_COND}$ and $\text{DIST}=\text{NORMAL}$) is the McNemar statistic Q_M , which has the following form when $\text{DPR}_0 = 1$:

$$Q_{M_0} = \frac{(n_{01} - n_{10})^2}{n_{01} + n_{10}}$$

For the conditional McNemar tests, this is equivalent to the square of the $Z(X)$ statistic for the test of a single proportion (normal approximation to binomial), where the proportion is $\frac{\pi_{01}}{\pi_{01} + \pi_{10}}$, the null is 0.5, and “ N ” is n_D (see, for example, Schork and Williams 1980):

$$\begin{aligned} Z(X) &= \frac{n_{01} - n_D(0.5)}{[n_D 0.5(1 - 0.5)]^{\frac{1}{2}}} \sim N \left(\frac{n_D^{\frac{1}{2}} (\frac{\pi_{01}}{\pi_{01} + \pi_{10}} - 0.5)}{[0.5(1 - 0.5)]^{\frac{1}{2}}}, \frac{\frac{\pi_{01}}{\pi_{01} + \pi_{10}} (1 - \frac{\pi_{01}}{\pi_{01} + \pi_{10}})}{0.5(1 - 0.5)} \right) \\ &= \frac{n_{01} - (n_{01} + n_{10})(0.5)}{[(n_{01} + n_{10})0.5(1 - 0.5)]^{\frac{1}{2}}} \\ &= \frac{n_{01} - n_{10}}{[n_{01} + n_{10}]^{\frac{1}{2}}} \\ &= \sqrt{Q_{M_0}} \end{aligned}$$

This can be generalized to a custom null for $\frac{\pi_{01}}{\pi_{01} + \pi_{10}}$, which is equivalent to specifying a custom null DPR :

$$\left[\frac{\pi_{01}}{\pi_{01} + \pi_{10}} \right]_0 \equiv \left[\frac{1}{1 + \frac{1}{\frac{\pi_{01}}{\pi_{10}}}} \right]_0 \equiv \frac{1}{1 + \frac{1}{\text{DPR}_0}}$$

So, a conditional McNemar test (asymptotic or exact) with a custom null is equivalent to the test of a single proportion $p_1 \equiv \frac{\pi_{01}}{\pi_{01} + \pi_{10}}$ with a null value $p_0 \equiv \frac{1}{1 + \frac{1}{\text{DPR}_0}}$, with a sample size of n_D :

$$H_0: p_1 = p_0$$

$$H_1: \begin{cases} p_1 \neq p_0, & \text{two-sided} \\ p_1 > p_0, & \text{one-sided U} \\ p_1 < p_0, & \text{one-sided L} \end{cases}$$

which is equivalent to

$$H_0: \text{DPR} = \text{DPR}_0$$

$$H_1: \begin{cases} \text{DPR} \neq \text{DPR}_0, & \text{two-sided} \\ \text{DPR} > \text{DPR}_0, & \text{one-sided U} \\ \text{DPR} < \text{DPR}_0, & \text{one-sided L} \end{cases}$$

The general form of the test statistic is thus

$$Q_M = \frac{(n_{01} - n_D p_0)^2}{n_D p_0 (1 - p_0)}$$

The two most common conditional McNemar tests assume either the exact conditional distribution of Q_M (covered by the DIST=EXACT_COND analysis) or a standard normal distribution for Q_M (covered by the DIST=NORMAL analysis).

McNemar Exact Conditional Test (TEST=MCNEMAR DIST=EXACT_COND)

For DIST=EXACT_COND, the power is calculated assuming that the test is conducted by using the exact conditional distribution of Q_M (conditional on n_D). The power is calculated by first computing the conditional power for each possible n_D . The unconditional power is computed as a weighted average over all possible outcomes of n_D :

$$\text{power} = \sum_{n_D=0}^N P(n_D) P(\text{Reject } p_1 = p_0 | n_D)$$

where $n_D \sim \text{Bin}(\pi_{01} + \pi_{10}, N)$, and $P(\text{Reject } p_1 = p_0 | n_D)$ is calculated by using the exact method in the section “Exact Test of a Binomial Proportion (TEST=EXACT)” on page 7347.

The achieved significance level, reported as “Actual Alpha” in the analysis, is computed in the same way except by using the actual alpha of the one-sample test in place of its power:

$$\text{actual alpha} = \sum_{n_D=0}^N P(n_D) \alpha^*(p_1, p_0 | n_D)$$

where $\alpha^*(p_1, p_0 | n_D)$ is the actual alpha calculated by using the exact method in the section “Exact Test of a Binomial Proportion (TEST=EXACT)” on page 7347 with proportion p_1 , null p_0 , and sample size n_D .

McNemar Normal Approximation Test (TEST=MCNEMAR DIST=NORMAL)

For DIST=NORMAL, power is calculated assuming the test is conducted by using the normal-approximate distribution of Q_M (conditional on n_D).

For the METHOD=EXACT option, the power is calculated in the same way as described in the section “McNemar Exact Conditional Test (TEST=MCNEMAR DIST=EXACT_COND)” on page 7372, except that $P(\text{Reject } p_1 = p_0 | n_D)$ is calculated by using the exact method in the section “z Test for Binomial Proportion Using Null Variance (TEST=Z VAREST=NULL)” on page 7348. The achieved significance level is calculated in the same way as described at the end of the section “McNemar Exact Conditional Test (TEST=MCNEMAR DIST=EXACT_COND)” on page 7372.

For the METHOD=MIETTINEN option, approximate sample size for the one-sided cases is computed according to equation (5.6) in Miettinen (1968):

$$N = \frac{\left\{ z_{1-\alpha}(p_{10} + p_{01}) + z_{power} \left[(p_{10} + p_{01})^2 - \frac{1}{4}(p_{01} - p_{10})^2(3 + p_{10} + p_{01}) \right]^{\frac{1}{2}} \right\}^2}{(p_{10} + p_{01})(p_{01} - p_{10})^2}$$

Approximate power for the one-sided cases is computed by solving the sample size equation for power, and approximate power for the two-sided case follows easily by summing the one-sided powers each at $\alpha/2$:

$$\text{power} = \begin{cases} \Phi \left(\frac{(p_{01} - p_{10})[N(p_{10} + p_{01})]^{\frac{1}{2}} - z_{1-\alpha}(p_{10} + p_{01})}{[(p_{10} + p_{01})^2 - \frac{1}{4}(p_{01} - p_{10})^2(3 + p_{10} + p_{01})]^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-(p_{01} - p_{10})[N(p_{10} + p_{01})]^{\frac{1}{2}} - z_{1-\alpha}(p_{10} + p_{01})}{[(p_{10} + p_{01})^2 - \frac{1}{4}(p_{01} - p_{10})^2(3 + p_{10} + p_{01})]^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{(p_{01} - p_{10})[N(p_{10} + p_{01})]^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}(p_{10} + p_{01})}{[(p_{10} + p_{01})^2 - \frac{1}{4}(p_{01} - p_{10})^2(3 + p_{10} + p_{01})]^{\frac{1}{2}}} \right) + \\ \Phi \left(\frac{-(p_{01} - p_{10})[N(p_{10} + p_{01})]^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}(p_{10} + p_{01})}{[(p_{10} + p_{01})^2 - \frac{1}{4}(p_{01} - p_{10})^2(3 + p_{10} + p_{01})]^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

The two-sided solution for N is obtained by numerically inverting the power equation.

In general, compared to METHOD=CONNOR, the METHOD=MIETTINEN approximation tends to be slightly more accurate but can be slightly anticonservative in the sense of underestimating sample size and overestimating power (Lachin 1992, p. 1250).

For the METHOD=CONNOR option, approximate sample size for the one-sided cases is computed according to equation (3) in Connor (1987):

$$N = \frac{\left\{ z_{1-\alpha}(p_{10} + p_{01})^{\frac{1}{2}} + z_{power} [p_{10} + p_{01} - (p_{01} - p_{10})^2]^{\frac{1}{2}} \right\}^2}{(p_{01} - p_{10})^2}$$

Approximate power for the one-sided cases is computed by solving the sample size equation for power, and approximate power for the two-sided case follows easily by summing the one-sided powers each at $\alpha/2$:

$$\text{power} = \begin{cases} \Phi \left(\frac{(p_{01} - p_{10})N^{\frac{1}{2}} - z_{1-\alpha}(p_{10} + p_{01})^{\frac{1}{2}}}{[p_{10} + p_{01} - (p_{01} - p_{10})^2]^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-(p_{01} - p_{10})N^{\frac{1}{2}} - z_{1-\alpha}(p_{10} + p_{01})^{\frac{1}{2}}}{[p_{10} + p_{01} - (p_{01} - p_{10})^2]^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{(p_{01} - p_{10})N^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}(p_{10} + p_{01})^{\frac{1}{2}}}{[p_{10} + p_{01} - (p_{01} - p_{10})^2]^{\frac{1}{2}}} \right) + \\ \Phi \left(\frac{-(p_{01} - p_{10})N^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}(p_{10} + p_{01})^{\frac{1}{2}}}{[p_{10} + p_{01} - (p_{01} - p_{10})^2]^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

The two-sided solution for N is obtained by numerically inverting the power equation.

In general, compared to METHOD=MIETTINEN, the METHOD=CONNOR approximation tends to be slightly less accurate but slightly conservative in the sense of overestimating sample size and underestimating power (Lachin 1992, p. 1250).

Analyses in the PAIREDMEANS Statement

Paired t Test (TEST=DIFF)

The hypotheses for the paired t test are

$$H_0: \mu_{\text{diff}} = \mu_0$$

$$H_1: \begin{cases} \mu_{\text{diff}} \neq \mu_0, & \text{two-sided} \\ \mu_{\text{diff}} > \mu_0, & \text{upper one-sided} \\ \mu_{\text{diff}} < \mu_0, & \text{lower one-sided} \end{cases}$$

The test assumes normally distributed data and requires $N \geq 2$. The test statistics are

$$t = N^{\frac{1}{2}} \left(\frac{\bar{d} - \mu_0}{s_d} \right) \sim t(N - 1, \delta)$$

$$t^2 \sim F(1, N - 1, \delta^2)$$

where $\bar{d}$ and s_d are the sample mean and standard deviation of the differences and

$$\delta = N^{\frac{1}{2}} \left(\frac{\mu_{\text{diff}} - \mu_0}{\sigma_{\text{diff}}} \right)$$

and

$$\sigma_{\text{diff}} = (\sigma_1^2 + \sigma_2^2 - 2\rho\sigma_1\sigma_2)^{\frac{1}{2}}$$

The test is

$$\text{Reject } H_0 \quad \text{if} \quad \begin{cases} t^2 \geq F_{1-\alpha}(1, N - 1), & \text{two-sided} \\ t \geq t_{1-\alpha}(N - 1), & \text{upper one-sided} \\ t \leq t_{\alpha}(N - 1), & \text{lower one-sided} \end{cases}$$

Exact power computations for t tests are given in O'Brien and Muller (1993, Section 8.2.2):

$$\text{power} = \begin{cases} P(F(1, N - 1, \delta^2) \geq F_{1-\alpha}(1, N - 1)), & \text{two-sided} \\ P(t(N - 1, \delta) \geq t_{1-\alpha}(N - 1)), & \text{upper one-sided} \\ P(t(N - 1, \delta) \leq t_{\alpha}(N - 1)), & \text{lower one-sided} \end{cases}$$

Paired t Test for Mean Ratio with Lognormal Data (TEST=RATIO)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section “Paired t Test (TEST=DIFF)” on page 7374 then apply.

In contrast to the usual t test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means.

The hypotheses for the paired t test with lognormal pairs $\{Y_1, Y_2\}$ are

$$H_0: \frac{\gamma_2}{\gamma_1} = \gamma_0$$

$$H_1: \begin{cases} \frac{\gamma_2}{\gamma_1} \neq \gamma_0, & \text{two-sided} \\ \frac{\gamma_2}{\gamma_1} > \gamma_0, & \text{upper one-sided} \\ \frac{\gamma_2}{\gamma_1} < \gamma_0, & \text{lower one-sided} \end{cases}$$

Let $\mu_1^*, \mu_2^*, \sigma_1^*, \sigma_2^*$, and ρ^* be the (arithmetic) means, standard deviations, and correlation of the bivariate normal distribution of the log-transformed data $\{\log Y_1, \log Y_2\}$. The hypotheses can be rewritten as follows:

$$H_0: \mu_2^* - \mu_1^* = \log(\gamma_0)$$

$$H_1: \begin{cases} \mu_2^* - \mu_1^* \neq \log(\gamma_0), & \text{two-sided} \\ \mu_2^* - \mu_1^* > \log(\gamma_0), & \text{upper one-sided} \\ \mu_2^* - \mu_1^* < \log(\gamma_0), & \text{lower one-sided} \end{cases}$$

where

$$\begin{aligned} \mu_1^* &= \log \gamma_1 \\ \mu_2^* &= \log \gamma_2 \\ \sigma_1^* &= [\log(CV_1^2 + 1)]^{\frac{1}{2}} \\ \sigma_2^* &= [\log(CV_2^2 + 1)]^{\frac{1}{2}} \\ \rho^* &= \frac{\log \{\rho CV_1 CV_2 + 1\}}{\sigma_1^* \sigma_2^*} \end{aligned}$$

and CV_1, CV_2 , and ρ are the coefficients of variation and the correlation of the original untransformed pairs $\{Y_1, Y_2\}$. The conversion from ρ to ρ^* is given by equation (44.36) on page 27 of Kotz, Balakrishnan, and Johnson (2000) and due to Jones and Miller (1966).

The valid range of ρ is restricted to (ρ_L, ρ_U) , where

$$\rho_L = \frac{\exp\left(-[\log(CV_1^2 + 1) \log(CV_2^2 + 1)]^{\frac{1}{2}}\right) - 1}{CV_1 CV_2}$$

$$\rho_U = \frac{\exp\left([\log(CV_1^2 + 1) \log(CV_2^2 + 1)]^{\frac{1}{2}}\right) - 1}{CV_1 CV_2}$$

These bounds are computed from equation (44.36) on page 27 of Kotz, Balakrishnan, and Johnson (2000) by observing that ρ is a monotonically increasing function of ρ^* and plugging in the values $\rho^* = -1$ and $\rho^* = 1$. Note that when the coefficients of variation are equal ($CV_1 = CV_2 = CV$), the bounds simplify to

$$\rho_L = \frac{-1}{CV^2 + 1}$$

$$\rho_U = 1$$

The test assumes lognormally distributed data and requires $N \geq 2$. The power is

$$\text{power} = \begin{cases} P(F(1, N-1, \delta^2) \geq F_{1-\alpha}(1, N-1)), & \text{two-sided} \\ P(t(N-1, \delta) \geq t_{1-\alpha}(N-1)), & \text{upper one-sided} \\ P(t(N-1, \delta) \leq t_{\alpha}(N-1)), & \text{lower one-sided} \end{cases}$$

where

$$\delta = N^{\frac{1}{2}} \left(\frac{\mu_1^* - \mu_2^* - \log(\gamma_0)}{\sigma^*} \right)$$

and

$$\sigma^* = (\sigma_1^{*2} + \sigma_2^{*2} - 2\rho^* \sigma_1^* \sigma_2^*)^{\frac{1}{2}}$$

Additive Equivalence Test for Mean Difference with Normal Data (TEST=EQUIV_DIFF)

The hypotheses for the equivalence test are

$$H_0: \mu_{\text{diff}} < \theta_L \quad \text{or} \quad \mu_{\text{diff}} > \theta_U$$

$$H_1: \theta_L \leq \mu_{\text{diff}} \leq \theta_U$$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987). The test assumes normally distributed data and requires $N \geq 2$. Phillips (1990) derives an expression for the exact power assuming a two-sample balanced design; the results are easily adapted to a paired design:

$$\begin{aligned} \text{power} = & Q_{N-1} \left((-t_{1-\alpha}(N-1)), \frac{\mu_{\text{diff}} - \theta_U}{\sigma_{\text{diff}} N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma_{\text{diff}} N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right) - \\ & Q_{N-1} \left((t_{1-\alpha}(N-1)), \frac{\mu_{\text{diff}} - \theta_L}{\sigma_{\text{diff}} N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma_{\text{diff}} N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right) \end{aligned}$$

where

$$\sigma_{\text{diff}} = (\sigma_1^2 + \sigma_2^2 - 2\rho\sigma_1\sigma_2)^{\frac{1}{2}}$$

and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section "Common Notation" on page 7336.

Multiplicative Equivalence Test for Mean Ratio with Lognormal Data (TEST=EQUIV_RATIO)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section "Additive Equivalence Test for Mean Difference with Normal Data (TEST=EQUIV_DIFF)" on page 7376 then apply.

In contrast to the additive equivalence test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means.

The hypotheses for the equivalence test are

$$H_0: \frac{\gamma_T}{\gamma_R} \leq \theta_L \quad \text{or} \quad \frac{\gamma_T}{\gamma_R} \geq \theta_U$$

$$H_1: \theta_L < \frac{\gamma_T}{\gamma_R} < \theta_U$$

where $0 < \theta_L < \theta_U$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987) on the log-transformed data. The test assumes lognormally distributed data and requires $N \geq 2$. Diletti, Hauschke, and Steinijans (1991) derive an expression for the exact power assuming a crossover design; the results are easily adapted to a paired design:

$$\text{power} = Q_{N-1} \left((-t_{1-\alpha}(N-1)), \frac{\log\left(\frac{\gamma_T}{\gamma_R}\right) - \log(\theta_U)}{\sigma^* N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right) - \\ Q_{N-1} \left((t_{1-\alpha}(N-1)), \frac{\log\left(\frac{\gamma_T}{\gamma_R}\right) - \log(\theta_L)}{\sigma^* N^{-\frac{1}{2}}}; 0, \frac{(N-1)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(t_{1-\alpha}(N-1))} \right)$$

where σ^* is the standard deviation of the differences between the log-transformed pairs (in other words, the standard deviation of $\log(Y_T) - \log(Y_R)$, where Y_T and Y_R are observations from the treatment and reference, respectively), computed as

$$\sigma^* = (\sigma_R^{*2} + \sigma_T^{*2} - 2\rho^* \sigma_R^* \sigma_T^*)^{\frac{1}{2}} \\ \sigma_R^* = [\log(\text{CV}_R^2 + 1)]^{\frac{1}{2}} \\ \sigma_T^* = [\log(\text{CV}_T^2 + 1)]^{\frac{1}{2}} \\ \rho^* = \frac{\log\{\rho \text{CV}_R \text{CV}_T + 1\}}{\sigma_R^* \sigma_T^*}$$

where CV_R , CV_T , and ρ are the coefficients of variation and the correlation of the original untransformed pairs $\{Y_T, Y_R\}$, and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function. The conversion from ρ to ρ^* is given by equation (44.36) on page 27 of Kotz, Balakrishnan, and Johnson (2000) and due to Jones and Miller (1966), and Owen's Q function is defined in the section "[Common Notation](#)" on page 7336.

The valid range of ρ is restricted to (ρ_L, ρ_U) , where

$$\rho_L = \frac{\exp\left(-[\log(\text{CV}_R^2 + 1)\log(\text{CV}_T^2 + 1)]^{\frac{1}{2}}\right) - 1}{\text{CV}_R \text{CV}_T} \\ \rho_U = \frac{\exp\left([\log(\text{CV}_R^2 + 1)\log(\text{CV}_T^2 + 1)]^{\frac{1}{2}}\right) - 1}{\text{CV}_R \text{CV}_T}$$

These bounds are computed from equation (44.36) on page 27 of Kotz, Balakrishnan, and Johnson (2000) by observing that ρ is a monotonically increasing function of ρ^* and plugging in the values $\rho^* = -1$ and $\rho^* = 1$. Note that when the coefficients of variation are equal ($\text{CV}_R = \text{CV}_T = \text{CV}$), the bounds simplify to

$$\rho_L = \frac{-1}{\text{CV}^2 + 1} \\ \rho_U = 1$$

Confidence Interval for Mean Difference (CI=DIFF)

This analysis of precision applies to the standard t -based confidence interval:

$$\begin{cases} \left[\bar{d} - t_{1-\frac{\alpha}{2}}(N-1) \frac{s_d}{\sqrt{N}}, \bar{d} + t_{1-\frac{\alpha}{2}}(N-1) \frac{s_d}{\sqrt{N}} \right], & \text{two-sided} \\ \left[\bar{d} - t_{1-\alpha}(N-1) \frac{s_d}{\sqrt{N}}, \infty \right), & \text{upper one-sided} \\ \left(-\infty, \bar{d} + t_{1-\alpha}(N-1) \frac{s_d}{\sqrt{N}} \right], & \text{lower one-sided} \end{cases}$$

where $\bar{d}$ and s_d are the sample mean and standard deviation of the differences. The “half-width” is defined as the distance from the point estimate $\bar{d}$ to a finite endpoint,

$$\text{half-width} = \begin{cases} t_{1-\frac{\alpha}{2}}(N-1) \frac{s_d}{\sqrt{N}}, & \text{two-sided} \\ t_{1-\alpha}(N-1) \frac{s_d}{\sqrt{N}}, & \text{one-sided} \end{cases}$$

A “valid” confidence interval captures the true mean difference. The exact probability of obtaining at most the target confidence interval half-width h , unconditional or conditional on validity, is given by Beal (1989):

$$\begin{aligned} \Pr(\text{half-width} \leq h) &= \begin{cases} P\left(\chi^2(N-1) \leq \frac{h^2 N(N-1)}{\sigma_{\text{diff}}^2(t_{1-\frac{\alpha}{2}}^2(N-1))}\right), & \text{two-sided} \\ P\left(\chi^2(N-1) \leq \frac{h^2 N(N-1)}{\sigma_{\text{diff}}^2(t_{1-\alpha}^2(N-1))}\right), & \text{one-sided} \end{cases} \\ \Pr(\text{half-width} \leq h | \text{validity}) &= \begin{cases} \left(\frac{1}{1-\alpha}\right) 2 \left[Q_{N-1}\left((t_{1-\frac{\alpha}{2}}(N-1)), 0; 0, b_1\right) - Q_{N-1}(0, 0; 0, b_1) \right], & \text{two-sided} \\ \left(\frac{1}{1-\alpha}\right) Q_{N-1}((t_{1-\alpha}(N-1)), 0; 0, b_1), & \text{one-sided} \end{cases} \end{aligned}$$

where

$$\begin{aligned} \sigma_{\text{diff}} &= (\sigma_1^2 + \sigma_2^2 - 2\rho\sigma_1\sigma_2)^{\frac{1}{2}} \\ b_1 &= \frac{h(N-1)^{\frac{1}{2}}}{\sigma_{\text{diff}}(t_{1-\frac{\alpha}{c}}(N-1))N^{-\frac{1}{2}}} \\ c &= \text{number of sides} \end{aligned}$$

and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen’s Q function, defined in the section “Common Notation” on page 7336.

A “quality” confidence interval is both sufficiently narrow (half-width $\leq h$) and valid:

$$\begin{aligned} \Pr(\text{quality}) &= \Pr(\text{half-width} \leq h \text{ and validity}) \\ &= \Pr(\text{half-width} \leq h | \text{validity})(1 - \alpha) \end{aligned}$$

Analyses in the TWOSAMPLEFREQ Statement**Overview of the 2×2 Table**

Notation:

		Outcome		
		Failure	Success	
Group	1	$n_1 - x_1$	x_1	n_1
	2	$n_2 - x_2$	x_2	n_2
		$N - m$	m	N

$x_1 = \text{\#successes in group 1}$

$x_2 = \text{\#successes in group 2}$

$m = x_1 + x_2 = \text{total \#successes}$

$$\hat{p}_1 = \frac{x_1}{n_1}$$

$$\hat{p}_2 = \frac{x_2}{n_2}$$

$$\hat{p} = \frac{m}{N} = w_1 \hat{p}_1 + w_2 \hat{p}_2$$

The hypotheses are

$$H_0: p_2 - p_1 = p_0$$

$$H_1: \begin{cases} p_2 - p_1 \neq p_0, & \text{two-sided} \\ p_2 - p_1 > p_0, & \text{upper one-sided} \\ p_2 - p_1 < p_0, & \text{lower one-sided} \end{cases}$$

where p_0 is constrained to be 0 for the likelihood ratio and Fisher's exact tests. If $p_0 < 0$ in an upper one-sided test or $p_0 > 0$ in a lower one-sided test, then the test is a noninferiority test. If $p_0 > 0$ in an upper one-sided test or $p_0 < 0$ in a lower one-sided test, then the test is a superiority test. Although p_0 is unconstrained for the Pearson chi-square test, $p_0 \neq 0$ is not recommended for that test. The Farrington-Manning score test is a better choice when $p_0 \neq 0$.

Internal calculations are performed in terms of p_1 , p_2 , and p_0 . An input set consisting of OR, p_1 , and OR_0 is transformed as follows:

$$p_2 = \frac{(OR)p_1}{1 - p_1 + (OR)p_1}$$

$$p_{10} = p_1$$

$$p_{20} = \frac{OR_0 p_{10}}{1 - p_{10} + (OR_0)p_{10}}$$

$$p_0 = p_{20} - p_{10}$$

An input set consisting of RR, p_1 , and RR_0 is transformed as follows:

$$p_2 = (RR)p_1$$

$$p_{10} = p_1$$

$$p_{20} = (RR_0)p_{10}$$

$$p_0 = p_{20} - p_{10}$$

The transformation of either OR_0 or RR_0 to p_0 is not unique. The chosen parameterization fixes the null value p_{10} at the input value of p_1 . Some values of OR_0 or RR_0 might lead to invalid values of p_0 ($p_0 \leq 0$ or $p_0 \geq 1$), in which case an “Invalid input” error occurs.

Farrington-Manning Score Test for Proportion Difference (TEST=FM)

The Farrington-Manning score test for proportion difference is based on equations (1), (2), and (12) in Farrington and Manning (1990). The test statistic, which is assumed to have a null distribution of $N(0, 1)$ under H_0 , is

$$z_{FMD} = \frac{\hat{p}_2 - \hat{p}_1 - p_0}{\left[\frac{\tilde{p}_1(1-\tilde{p}_1)}{n_1} + \frac{\tilde{p}_2(1-\tilde{p}_2)}{n_2} \right]^{\frac{1}{2}}} = [Nw_1w_2]^{\frac{1}{2}} \frac{\hat{p}_2 - \hat{p}_1 - p_0}{[w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)]^{\frac{1}{2}}}$$

where $\tilde{p}_1$ and $\tilde{p}_2$ are the maximum likelihood estimates of the proportions under the restriction $\tilde{p}_2 - \tilde{p}_1 = p_0$.

Sample size for the one-sided cases is given by equations (4) and (12) in Farrington and Manning (1990). One-sided power is computed by inverting the sample size formula. Power for the two-sided case is computed by adding the lower-sided and upper-sided powers, each evaluated at $\alpha/2$. Sample size for the two-sided case is obtained by numerically inverting the power formula,

$$\text{power} = \begin{cases} \Phi \left(\frac{(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\alpha}[w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)]^{\frac{1}{2}}}{[w_2p_1(1-p_1) + w_1p_2(1-p_2)]^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\alpha}[w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)]^{\frac{1}{2}}}{[w_2p_1(1-p_1) + w_1p_2(1-p_2)]^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}[w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)]^{\frac{1}{2}}}{[w_2p_1(1-p_1) + w_1p_2(1-p_2)]^{\frac{1}{2}}} \right) + \\ \Phi \left(\frac{-(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}[w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)]^{\frac{1}{2}}}{[w_2p_1(1-p_1) + w_1p_2(1-p_2)]^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

where

$$\begin{aligned} \tilde{p}_2 &= 2u \cos(w) - b/(3a) \\ \tilde{p}_1 &= \tilde{p}_2 - p_0 \\ w &= (\pi + \cos^{-1}(v/u^3))/3 \\ v &= b^3/(3a)^3 - bc/(6a^2) + d/(2a) \\ u &= \text{sign}(v)\sqrt{b^2/(3a)^2 - c/(3a)} \\ a &= 1 + w_1/w_2 \\ b &= -[1 + w_1/w_2 + p_2 + (w_1/w_2)p_1 + p_0(w_1/w_2 + 2)] \\ c &= p_0^2 + p_0(2p_2 + w_1/w_2 + 1) + p_2 + (w_1/w_2)p_1 \\ d &= -p_2p_0(1 + p_0) \end{aligned}$$

For the one-sided cases, a closed-form inversion of the power equation yields an approximate total sample size of

$$N = \frac{\left[z_{1-\alpha} \{w_2\tilde{p}_1(1-\tilde{p}_1) + w_1\tilde{p}_2(1-\tilde{p}_2)\}^{\frac{1}{2}} + z_{\text{power}} \{w_2p_1(1-p_1) + w_1p_2(1-p_2)\}^{\frac{1}{2}} \right]^2}{w_1w_2(p_2 - p_1 - p_0)^2}$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Farrington-Manning Score Test for Relative Risk (TEST=FM_RR)

The Farrington-Manning score test is based on equations (5), (6), and (13) in Farrington and Manning (1990). The test statistic, which is assumed to have a null distribution of $N(0, 1)$ under H_0 , is

$$z_{\text{FMR}} = \frac{\hat{p}_2 - \text{RR}_0 \hat{p}_1}{\left[\frac{\text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1)}{n_1} + \frac{\tilde{p}_2 (1 - \tilde{p}_2)}{n_2} \right]^{\frac{1}{2}}} = [N w_1 w_2]^{\frac{1}{2}} \frac{\hat{p}_2 - \text{RR}_0 \hat{p}_1}{[w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)]^{\frac{1}{2}}}$$

where $\tilde{p}_1$ and $\tilde{p}_2$ are the maximum likelihood estimates of the proportions under the restriction $\tilde{p}_2 / \tilde{p}_1 = \text{RR}_0$.

Sample size for the one-sided cases is given by equations (8) and (13) in Farrington and Manning (1990). One-sided power is computed by inverting the sample size formula. Power for the two-sided case is computed by adding the lower-sided and upper-sided powers, each evaluated at $\alpha/2$. Sample size for the two-sided case is obtained by numerically inverting the power formula,

$$\text{power} = \begin{cases} \Phi \left(\frac{(\tilde{p}_2 - \text{RR}_0 \tilde{p}_1)(N w_1 w_2)^{\frac{1}{2}} - z_{1-\alpha} [w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)]^{\frac{1}{2}}}{[w_2 \text{RR}_0^2 p_1 (1 - p_1) + w_1 p_2 (1 - p_2)]^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-(\tilde{p}_2 - \text{RR}_0 \tilde{p}_1)(N w_1 w_2)^{\frac{1}{2}} - z_{1-\alpha} [w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)]^{\frac{1}{2}}}{[w_2 \text{RR}_0^2 p_1 (1 - p_1) + w_1 p_2 (1 - p_2)]^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{(\tilde{p}_2 - \text{RR}_0 \tilde{p}_1)(N w_1 w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}} [w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)]^{\frac{1}{2}}}{[w_2 \text{RR}_0^2 p_1 (1 - p_1) + w_1 p_2 (1 - p_2)]^{\frac{1}{2}}} \right) + \\ \Phi \left(\frac{-(\tilde{p}_2 - \text{RR}_0 \tilde{p}_1)(N w_1 w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}} [w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)]^{\frac{1}{2}}}{[w_2 \text{RR}_0^2 p_1 (1 - p_1) + w_1 p_2 (1 - p_2)]^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

where

$$\begin{aligned} \tilde{p}_2 &= \frac{-b - (b^2 - 4ac)^{\frac{1}{2}}}{2a} \\ \tilde{p}_1 &= \tilde{p}_2 / \text{RR}_0 \\ a &= 1 + w_1 / w_2 \\ b &= -[\text{RR}_0 (1 + (w_1 / w_2) p_1) + p_2 + w_1 / w_2] \\ c &= \text{RR}_0 (p_2 + (w_1 / w_2) p_1) \end{aligned}$$

For the one-sided cases, a closed-form inversion of the power equation yields an approximate total sample size of

$$N = \frac{\left[z_{1-\alpha} \{w_2 \text{RR}_0^2 \tilde{p}_1 (1 - \tilde{p}_1) + w_1 \tilde{p}_2 (1 - \tilde{p}_2)\}^{\frac{1}{2}} + z_{\text{power}} \{w_2 \text{RR}_0^2 p_1 (1 - p_1) + w_1 p_2 (1 - p_2)\}^{\frac{1}{2}} \right]^2}{w_1 w_2 (p_2 - \text{RR}_0 p_1)^2}$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Pearson Chi-Square Test for Two Proportions (TEST=PCHI)

The usual Pearson chi-square test is unconditional. The test statistic

$$z_P = \frac{\hat{p}_2 - \hat{p}_1 - p_0}{\left[\hat{p}(1 - \hat{p}) \left(\frac{1}{n_1} + \frac{1}{n_2} \right) \right]^{\frac{1}{2}}} = [Nw_1w_2]^{\frac{1}{2}} \frac{\hat{p}_2 - \hat{p}_1 - p_0}{[\hat{p}(1 - \hat{p})]^{\frac{1}{2}}}$$

is assumed to have a null distribution of $N(0, 1)$.

Sample size for the one-sided cases is given by equation (4) in Fleiss, Tytun, and Ury (1980). One-sided power is computed as suggested by Diegert and Diegert (1981) by inverting the sample size formula. Power for the two-sided case is computed by adding the lower-sided and upper-sided powers each evaluated at $\alpha/2$. Sample size for the two-sided case is obtained by numerically inverting the power formula. A custom null value p_0 for the proportion difference $p_2 - p_1$ is also supported, but it is not recommended. If you are using a nondefault null value, then the Farrington-Manning score test is a better choice.

$$\text{power} = \begin{cases} \Phi \left(\frac{(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\alpha}[(w_1p_1 + w_2p_2)(1 - w_1p_1 - w_2p_2)]^{\frac{1}{2}}}{[w_2p_1(1 - p_1) + w_1p_2(1 - p_2)]^{\frac{1}{2}}} \right), & \text{upper one-sided} \\ \Phi \left(\frac{-(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\alpha}[(w_1p_1 + w_2p_2)(1 - w_1p_1 - w_2p_2)]^{\frac{1}{2}}}{[w_2p_1(1 - p_1) + w_1p_2(1 - p_2)]^{\frac{1}{2}}} \right), & \text{lower one-sided} \\ \Phi \left(\frac{(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}[(w_1p_1 + w_2p_2)(1 - w_1p_1 - w_2p_2)]^{\frac{1}{2}}}{[w_2p_1(1 - p_1) + w_1p_2(1 - p_2)]^{\frac{1}{2}}} \right) + \\ \Phi \left(\frac{-(p_2 - p_1 - p_0)(Nw_1w_2)^{\frac{1}{2}} - z_{1-\frac{\alpha}{2}}[(w_1p_1 + w_2p_2)(1 - w_1p_1 - w_2p_2)]^{\frac{1}{2}}}{[w_2p_1(1 - p_1) + w_1p_2(1 - p_2)]^{\frac{1}{2}}} \right), & \text{two-sided} \end{cases}$$

For the one-sided cases, a closed-form inversion of the power equation yields an approximate total sample size

$$N = \frac{\left[z_{1-\alpha} \{ (w_1p_1 + w_2p_2)(1 - w_1p_1 - w_2p_2) \}^{\frac{1}{2}} + z_{\text{power}} \{ w_2p_1(1 - p_1) + w_1p_2(1 - p_2) \}^{\frac{1}{2}} \right]^2}{w_1w_2(p_2 - p_1 - p_0)^2}$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Likelihood Ratio Chi-Square Test for Two Proportions (TEST=LRCHI)

The usual likelihood ratio chi-square test is unconditional. The test statistic

$$z_{LR} = (-1_{\{p_2 < p_1\}}) \sqrt{2N \sum_{i=1}^2 \left[w_i \hat{p}_i \log \left(\frac{\hat{p}_i}{\hat{p}} \right) + w_i (1 - \hat{p}_i) \log \left(\frac{1 - \hat{p}_i}{1 - \hat{p}} \right) \right]}$$

is assumed to have a null distribution of $N(0, 1)$ and an alternative distribution of $N(\delta, 1)$, where

$$\delta = N^{\frac{1}{2}} (-1_{\{p_2 < p_1\}}) \sqrt{2 \sum_{i=1}^2 \left[w_i p_i \log \left(\frac{p_i}{w_1p_1 + w_2p_2} \right) + w_i (1 - p_i) \log \left(\frac{1 - p_i}{1 - (w_1p_1 + w_2p_2)} \right) \right]}$$

The approximate power is

$$\text{power} = \begin{cases} \Phi(\delta - z_{1-\alpha}), & \text{upper one-sided} \\ \Phi(-\delta - z_{1-\alpha}), & \text{lower one-sided} \\ \Phi\left(\delta - z_{1-\frac{\alpha}{2}}\right) + \Phi\left(-\delta - z_{1-\frac{\alpha}{2}}\right), & \text{two-sided} \end{cases}$$

For the one-sided cases, a closed-form inversion of the power equation yield an approximate total sample size

$$N = \left(\frac{z_{\text{power}} + z_{1-\alpha}}{\delta} \right)^2$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Fisher's Exact Conditional Test for Two Proportions (Test=FISHER)

Fisher's exact test is conditional on the observed total number of successes m . Power and sample size computations are based on a test with similar power properties, the continuity-adjusted arcsine test. The test statistic

$$z_A = (4Nw_1w_2)^{\frac{1}{2}} \left[\arcsin \left(\left[\hat{p}_2 + \frac{1}{2Nw_2} (1_{\{\hat{p}_2 < \hat{p}_1\}} - 1_{\{\hat{p}_2 > \hat{p}_1\}}) \right]^{\frac{1}{2}} \right) - \arcsin \left(\left[\hat{p}_1 + \frac{1}{2Nw_1} (1_{\{\hat{p}_1 < \hat{p}_2\}} - 1_{\{\hat{p}_1 > \hat{p}_2\}}) \right]^{\frac{1}{2}} \right) \right]$$

is assumed to have a null distribution of $N(0, 1)$ and an alternative distribution of $N(\delta, 1)$, where

$$\delta = (4Nw_1w_2)^{\frac{1}{2}} \left[\arcsin \left(\left[p_2 + \frac{1}{2Nw_2} (1_{\{p_2 < p_1\}} - 1_{\{p_2 > p_1\}}) \right]^{\frac{1}{2}} \right) - \arcsin \left(\left[p_1 + \frac{1}{2Nw_1} (1_{\{p_1 < p_2\}} - 1_{\{p_1 > p_2\}}) \right]^{\frac{1}{2}} \right) \right]$$

The approximate power for the one-sided balanced case is given by Walters (1979) and is easily extended to the unbalanced and two-sided cases:

$$\text{power} = \begin{cases} \Phi(\delta - z_{1-\alpha}), & \text{upper one-sided} \\ \Phi(-\delta - z_{1-\alpha}), & \text{lower one-sided} \\ \Phi\left(\delta - z_{1-\frac{\alpha}{2}}\right) + \Phi\left(-\delta - z_{1-\frac{\alpha}{2}}\right), & \text{two-sided} \end{cases}$$

The approximation is valid only for $N \geq 1/(2w_1w_2|p_1 - p_2|)$.

Analyses in the TWOSAMPLEMEANS Statement

Two-Sample t Test Assuming Equal Variances (TEST=DIFF)

The hypotheses for the two-sample t test are

$$H_0: \mu_{\text{diff}} = \mu_0$$

$$H_1: \begin{cases} \mu_{\text{diff}} \neq \mu_0, & \text{two-sided} \\ \mu_{\text{diff}} > \mu_0, & \text{upper one-sided} \\ \mu_{\text{diff}} < \mu_0, & \text{lower one-sided} \end{cases}$$

The test assumes normally distributed data and common standard deviation per group, and it requires $N \geq 3$, $n_1 \geq 1$, and $n_2 \geq 1$. The test statistics are

$$t = N^{\frac{1}{2}}(w_1w_2)^{\frac{1}{2}} \left(\frac{\bar{x}_2 - \bar{x}_1 - \mu_0}{s_p} \right) \sim t(N-2, \delta)$$

$$t^2 \sim F(1, N-2, \delta^2)$$

where $\bar{x}_1$ and $\bar{x}_2$ are the sample means and s_p is the pooled standard deviation, and

$$\delta = N^{\frac{1}{2}}(w_1 w_2)^{\frac{1}{2}} \left(\frac{\mu_{\text{diff}} - \mu_0}{\sigma} \right)$$

The test is

$$\text{Reject } H_0 \text{ if } \begin{cases} t^2 \geq F_{1-\alpha}(1, N-2), & \text{two-sided} \\ t \geq t_{1-\alpha}(N-2), & \text{upper one-sided} \\ t \leq t_{\alpha}(N-2), & \text{lower one-sided} \end{cases}$$

Exact power computations for t tests are given in O'Brien and Muller (1993, Section 8.2.1):

$$\text{power} = \begin{cases} P(F(1, N-2, \delta^2) \geq F_{1-\alpha}(1, N-2)), & \text{two-sided} \\ P(t(N-2, \delta) \geq t_{1-\alpha}(N-2)), & \text{upper one-sided} \\ P(t(N-2, \delta) \leq t_{\alpha}(N-2)), & \text{lower one-sided} \end{cases}$$

Solutions for N , n_1 , n_2 , α , and δ are obtained by numerically inverting the power equation. Closed-form solutions for other parameters, in terms of δ , are as follows:

$$\begin{aligned} \mu_{\text{diff}} &= \delta \sigma (N w_1 w_2)^{-\frac{1}{2}} + \mu_0 \\ \mu_1 &= \delta \sigma (N w_1 w_2)^{-\frac{1}{2}} + \mu_0 - \mu_2 \\ \mu_2 &= \delta \sigma (N w_1 w_2)^{-\frac{1}{2}} + \mu_0 - \mu_1 \\ \sigma &= \begin{cases} \delta^{-1} (N w_1 w_2)^{\frac{1}{2}} (\mu_{\text{diff}} - \mu_0), & |\delta| > 0 \\ \text{undefined}, & \text{otherwise} \end{cases} \\ w_1 &= \begin{cases} \frac{1}{2} \pm \frac{1}{2} \left[1 - \frac{4\delta^2 \sigma^2}{N(\mu_{\text{diff}} - \mu_0)^2} \right]^{\frac{1}{2}}, & 0 < |\delta| \leq \frac{1}{2} N^{\frac{1}{2}} \frac{|\mu_{\text{diff}} - \mu_0|}{\sigma} \\ \text{undefined}, & \text{otherwise} \end{cases} \\ w_2 &= \begin{cases} \frac{1}{2} \pm \frac{1}{2} \left[1 - \frac{4\delta^2 \sigma^2}{N(\mu_{\text{diff}} - \mu_0)^2} \right]^{\frac{1}{2}}, & 0 < |\delta| \leq \frac{1}{2} N^{\frac{1}{2}} \frac{|\mu_{\text{diff}} - \mu_0|}{\sigma} \\ \text{undefined}, & \text{otherwise} \end{cases} \end{aligned}$$

Finally, here is a derivation of the solution for w_1 :

Solve the δ equation for w_1 (which requires the quadratic formula). Then determine the range of δ given w_1 :

$$\begin{aligned} \min(\delta) &= \begin{cases} 0, & \text{when } w_1 = 0 \text{ or } 1, \text{ if } (\mu_{\text{diff}} - \mu_0) \geq 0 \\ \frac{1}{2} N^{\frac{1}{2}} \frac{(\mu_{\text{diff}} - \mu_0)}{\sigma}, & \text{when } w_1 = \frac{1}{2}, \text{ if } (\mu_{\text{diff}} - \mu_0) < 0 \end{cases} \\ \max(\delta) &= \begin{cases} 0, & \text{when } w_1 = 0 \text{ or } 1, \text{ if } (\mu_{\text{diff}} - \mu_0) < 0 \\ \frac{1}{2} N^{\frac{1}{2}} \frac{(\mu_{\text{diff}} - \mu_0)}{\sigma}, & \text{when } w_1 = \frac{1}{2}, \text{ if } (\mu_{\text{diff}} - \mu_0) \geq 0 \end{cases} \end{aligned}$$

This implies

$$|\delta| \leq \frac{1}{2} N^{\frac{1}{2}} \frac{|\mu_{\text{diff}} - \mu_0|}{\sigma}$$

Two-Sample Satterthwaite t Test Assuming Unequal Variances (TEST=DIFF_SATT)

The hypotheses for the two-sample Satterthwaite t test are

$$H_0: \mu_{\text{diff}} = \mu_0$$

$$H_1: \begin{cases} \mu_{\text{diff}} \neq \mu_0, & \text{two-sided} \\ \mu_{\text{diff}} > \mu_0, & \text{upper one-sided} \\ \mu_{\text{diff}} < \mu_0, & \text{lower one-sided} \end{cases}$$

The test assumes normally distributed data and requires $N \geq 3$, $n_1 \geq 1$, and $n_2 \geq 1$. The test statistics are

$$t = \frac{\bar{x}_2 - \bar{x}_1 - \mu_0}{\left[\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right]^{\frac{1}{2}}} = N^{\frac{1}{2}} \frac{\bar{x}_2 - \bar{x}_1 - \mu_0}{\left[\frac{s_1^2}{w_1} + \frac{s_2^2}{w_2} \right]^{\frac{1}{2}}}$$

$$F = t^2$$

where $\bar{x}_1$ and $\bar{x}_2$ are the sample means and s_1 and s_2 are the sample standard deviations.

DiSantostefano and Muller (1995, p. 585) state, the test is based on assuming that under H_0 , F is distributed as $F(1, \nu)$, where ν is given by Satterthwaite's approximation (Satterthwaite 1946),

$$\nu = \frac{\left[\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2} \right]^2}{\frac{\left[\frac{\sigma_1^2}{n_1} \right]^2}{n_1 - 1} + \frac{\left[\frac{\sigma_2^2}{n_2} \right]^2}{n_2 - 1}} = \frac{\left[\frac{\sigma_1^2}{w_1} + \frac{\sigma_2^2}{w_2} \right]^2}{\frac{\left[\frac{\sigma_1^2}{w_1} \right]^2}{Nw_1 - 1} + \frac{\left[\frac{\sigma_2^2}{w_2} \right]^2}{Nw_2 - 1}}$$

Since ν is unknown, in practice it must be replaced by an estimate

$$\hat{\nu} = \frac{\left[\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} \right]^2}{\frac{\left[\frac{s_1^2}{n_1} \right]^2}{n_1 - 1} + \frac{\left[\frac{s_2^2}{n_2} \right]^2}{n_2 - 1}} = \frac{\left[\frac{s_1^2}{w_1} + \frac{s_2^2}{w_2} \right]^2}{\frac{\left[\frac{s_1^2}{w_1} \right]^2}{Nw_1 - 1} + \frac{\left[\frac{s_2^2}{w_2} \right]^2}{Nw_2 - 1}}$$

So the test is

$$\text{Reject } H_0 \quad \text{if} \quad \begin{cases} F \geq F_{1-\alpha}(1, \hat{\nu}), & \text{two-sided} \\ t \geq t_{1-\alpha}(\hat{\nu}), & \text{upper one-sided} \\ t \leq t_{\alpha}(\hat{\nu}), & \text{lower one-sided} \end{cases}$$

Exact solutions for power for the two-sided and upper one-sided cases are given in Moser, Stevens, and Watts (1989). The lower one-sided case follows easily by using symmetry. The equations are as follows:

$$\text{power} = \begin{cases} \int_0^\infty P(F(1, N-2, \lambda) > h(u) F_{1-\alpha}(1, v(u)) | u) f(u) du, & \text{two-sided} \\ \int_0^\infty P\left(t(N-2, \lambda^{\frac{1}{2}}) > [h(u)]^{\frac{1}{2}} t_{1-\alpha}(v(u)) | u\right) f(u) du, & \text{upper one-sided} \\ \int_0^\infty P\left(t(N-2, \lambda^{\frac{1}{2}}) < [h(u)]^{\frac{1}{2}} t_{\alpha}(v(u)) | u\right) f(u) du, & \text{lower one-sided} \end{cases}$$

where

$$h(u) = \frac{\left(\frac{1}{n_1} + \frac{u}{n_2}\right)(n_1 + n_2 - 2)}{\left[(n_1 - 1) + (n_2 - 1)\frac{u\sigma_1^2}{\sigma_2^2}\right]\left(\frac{1}{n_1} + \frac{\sigma_2^2}{\sigma_1^2 n_2}\right)}$$

$$v(u) = \frac{\left(\frac{1}{n_1} + \frac{u}{n_2}\right)^2}{\frac{1}{n_1^2(n_1-1)} + \frac{u^2}{n_2^2(n_2-1)}}$$

$$\lambda = \frac{(\mu_{\text{diff}} - \mu_0)^2}{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$

$$f(u) = \frac{\Gamma\left(\frac{n_1+n_2-2}{2}\right)}{\Gamma\left(\frac{n_1-1}{2}\right)\Gamma\left(\frac{n_2-1}{2}\right)} \left[\frac{\sigma_1^2(n_2-1)}{\sigma_2^2(n_1-1)}\right]^{\frac{n_2-1}{2}} u^{\frac{n_2-3}{2}} \left[1 + \left(\frac{n_2-1}{n_1-1}\right)\frac{u\sigma_1^2}{\sigma_2^2}\right]^{-\left(\frac{n_1+n_2-2}{2}\right)}$$

The density $f(u)$ is obtained from the fact that

$$\frac{u\sigma_1^2}{\sigma_2^2} \sim F(n_2 - 1, n_1 - 1)$$

Because the test is biased, the achieved significance level might differ from the nominal significance level. The actual alpha is computed in the same way as the power, except that the mean difference μ_{diff} is replaced by the null mean difference μ_0 .

Two-Sample Pooled t Test of Mean Ratio with Lognormal Data (TEST=RATIO)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section “Two-Sample t Test Assuming Equal Variances (TEST=DIFF)” on page 7383 then apply.

In contrast to the usual t test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means. The test assumes equal coefficients of variation in the two groups.

The hypotheses for the two-sample t test with lognormal data are

$$H_0: \frac{\gamma_2}{\gamma_1} = \gamma_0$$

$$H_1: \begin{cases} \frac{\gamma_2}{\gamma_1} \neq \gamma_0, & \text{two-sided} \\ \frac{\gamma_2}{\gamma_1} > \gamma_0, & \text{upper one-sided} \\ \frac{\gamma_2}{\gamma_1} < \gamma_0, & \text{lower one-sided} \end{cases}$$

Let μ_1^* , μ_2^* , and σ^* be the (arithmetic) means and common standard deviation of the corresponding normal distributions of the log-transformed data. The hypotheses can be rewritten as follows:

$$H_0: \mu_2^* - \mu_1^* = \log(\gamma_0)$$

$$H_1: \begin{cases} \mu_2^* - \mu_1^* \neq \log(\gamma_0), & \text{two-sided} \\ \mu_2^* - \mu_1^* > \log(\gamma_0), & \text{upper one-sided} \\ \mu_2^* - \mu_1^* < \log(\gamma_0), & \text{lower one-sided} \end{cases}$$

where

$$\mu_1^* = \log \gamma_1$$

$$\mu_2^* = \log \gamma_2$$

The test assumes lognormally distributed data and requires $N \geq 3$, $n_1 \geq 1$, and $n_2 \geq 1$.

The power is

$$\text{power} = \begin{cases} P(F(1, N-2, \delta^2) \geq F_{1-\alpha}(1, N-2)), & \text{two-sided} \\ P(t(N-2, \delta) \geq t_{1-\alpha}(N-2)), & \text{upper one-sided} \\ P(t(N-2, \delta) \leq t_{\alpha}(N-2)), & \text{lower one-sided} \end{cases}$$

where

$$\delta = N^{\frac{1}{2}}(w_1 w_2)^{\frac{1}{2}} \left(\frac{\mu_2^* - \mu_1^* - \log(\gamma_0)}{\sigma^*} \right)$$

$$\sigma^* = [\log(\text{CV}^2 + 1)]^{\frac{1}{2}}$$

Additive Equivalence Test for Mean Difference with Normal Data (TEST=EQUIV_DIFF)

The hypotheses for the equivalence test are

$$H_0: \mu_{\text{diff}} < \theta_L \quad \text{or} \quad \mu_{\text{diff}} > \theta_U$$

$$H_1: \theta_L \leq \mu_{\text{diff}} \leq \theta_U$$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987). The test assumes normally distributed data and requires $N \geq 3$, $n_1 \geq 1$, and $n_2 \geq 1$. Phillips (1990) derives an expression for the exact

power assuming a balanced design; the results are easily adapted to an unbalanced design:

$$\text{power} = Q_{N-2} \left((-t_{1-\alpha}(N-2)), \frac{\mu_{\text{diff}} - \theta_U}{\sigma N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}}; 0, \frac{(N-2)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}(t_{1-\alpha}(N-2))} \right) - \\ Q_{N-2} \left((t_{1-\alpha}(N-2)), \frac{\mu_{\text{diff}} - \theta_L}{\sigma N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}}; 0, \frac{(N-2)^{\frac{1}{2}}(\theta_U - \theta_L)}{2\sigma N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}(t_{1-\alpha}(N-2))} \right)$$

where $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section "Common Notation" on page 7336.

Multiplicative Equivalence Test for Mean Ratio with Lognormal Data (TEST=EQUIV_RATIO)

The lognormal case is handled by reexpressing the analysis equivalently as a normality-based test on the log-transformed data, by using properties of the lognormal distribution as discussed in Johnson, Kotz, and Balakrishnan (1994, Chapter 14). The approaches in the section "Additive Equivalence Test for Mean Difference with Normal Data (TEST=EQUIV_DIFF)" on page 7387 then apply.

In contrast to the additive equivalence test on normal data, the hypotheses with lognormal data are defined in terms of geometric means rather than arithmetic means.

The hypotheses for the equivalence test are

$$H_0: \frac{\gamma_T}{\gamma_R} \leq \theta_L \quad \text{or} \quad \frac{\gamma_T}{\gamma_R} \geq \theta_U$$

$$H_1: \theta_L < \frac{\gamma_T}{\gamma_R} < \theta_U$$

$$\text{where } 0 < \theta_L < \theta_U$$

The analysis is the two one-sided tests (TOST) procedure of Schuirmann (1987) on the log-transformed data. The test assumes lognormally distributed data and requires $N \geq 3$, $n_1 \geq 1$, and $n_2 \geq 1$. Diletti, Hauschke, and Steinijans (1991) derive an expression for the exact power assuming a crossover design; the results are easily adapted to an unbalanced two-sample design:

$$\text{power} = Q_{N-2} \left((-t_{1-\alpha}(N-2)), \frac{\log\left(\frac{\gamma_T}{\gamma_R}\right) - \log(\theta_U)}{\sigma^* N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}}; 0, \frac{(N-2)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}(t_{1-\alpha}(N-2))} \right) - \\ Q_{N-2} \left((t_{1-\alpha}(N-2)), \frac{\log\left(\frac{\gamma_T}{\gamma_R}\right) - \log(\theta_L)}{\sigma^* N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}}; 0, \frac{(N-2)^{\frac{1}{2}}(\log(\theta_U) - \log(\theta_L))}{2\sigma^* N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}(t_{1-\alpha}(N-2))} \right)$$

where

$$\sigma^* = [\log(\text{CV}^2 + 1)]^{\frac{1}{2}}$$

is the (assumed common) standard deviation of the normal distribution of the log-transformed data, and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen's Q function, defined in the section "Common Notation" on page 7336.

Confidence Interval for Mean Difference (CI=DIFF)

This analysis of precision applies to the standard t -based confidence interval:

$$\begin{aligned} & \left[(\bar{x}_2 - \bar{x}_1) - t_{1-\frac{\alpha}{2}}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}}, \right. \\ & \quad \left. (\bar{x}_2 - \bar{x}_1) + t_{1-\frac{\alpha}{2}}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}} \right], \quad \text{two-sided} \\ & \left[(\bar{x}_2 - \bar{x}_1) - t_{1-\alpha}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}}, \infty \right), \quad \text{upper one-sided} \\ & \left(-\infty, (\bar{x}_2 - \bar{x}_1) + t_{1-\alpha}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}} \right], \quad \text{lower one-sided} \end{aligned}$$

where $\bar{x}_1$ and $\bar{x}_2$ are the sample means and s_p is the pooled standard deviation. The “half-width” is defined as the distance from the point estimate $\bar{x}_2 - \bar{x}_1$ to a finite endpoint,

$$\text{half-width} = \begin{cases} t_{1-\frac{\alpha}{2}}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}}, & \text{two-sided} \\ t_{1-\alpha}(N-2) \frac{s_p}{\sqrt{N w_1 w_2}}, & \text{one-sided} \end{cases}$$

A “valid” confidence interval captures the true mean. The exact probability of obtaining at most the target confidence interval half-width h , unconditional or conditional on validity, is given by Beal (1989):

$$\begin{aligned} \Pr(\text{half-width} \leq h) &= \begin{cases} P \left(\chi^2(N-2) \leq \frac{h^2 N(N-2)(w_1 w_2)}{\sigma^2(t_{1-\frac{\alpha}{2}}^2(N-2))} \right), & \text{two-sided} \\ P \left(\chi^2(N-2) \leq \frac{h^2 N(N-2)(w_1 w_2)}{\sigma^2(t_{1-\alpha}^2(N-2))} \right), & \text{one-sided} \end{cases} \\ \Pr(\text{half-width} \leq h | \text{validity}) &= \begin{cases} \left(\frac{1}{1-\alpha} \right) 2 \left[Q_{N-2} \left((t_{1-\frac{\alpha}{2}}(N-2)), 0; \right. \right. \\ \quad \left. \left. 0, b_2 \right) - Q_{N-2}(0, 0; 0, b_2) \right], & \text{two-sided} \\ \left(\frac{1}{1-\alpha} \right) Q_{N-2} \left((t_{1-\alpha}(N-2)), 0; 0, b_2 \right), & \text{one-sided} \end{cases} \end{aligned}$$

where

$$b_2 = \frac{h(N-2)^{\frac{1}{2}}}{\sigma(t_{1-\frac{\alpha}{2}}(N-2))N^{-\frac{1}{2}}(w_1 w_2)^{-\frac{1}{2}}}$$

c = number of sides

and $Q(\cdot, \cdot; \cdot, \cdot)$ is Owen’s Q function, defined in the section “Common Notation” on page 7336.

A “quality” confidence interval is both sufficiently narrow (half-width $\leq h$) and valid:

$$\begin{aligned} \Pr(\text{quality}) &= \Pr(\text{half-width} \leq h \text{ and validity}) \\ &= \Pr(\text{half-width} \leq h | \text{validity})(1 - \alpha) \end{aligned}$$

Analyses in the TWOSAMPLESURVIVAL Statement**Rank Tests for Two Survival Curves (TEST=LOGRANK, TEST=GEHAN, TEST=TARONEWARE)**

The method is from Lakatos (1988) and Cantor (1997, pp. 83–92).

Define the following notation:

- $X_j(i)$ = i th input time point on survival curve for group j
- $S_j(i)$ = input survivor function value that corresponds to $X_j(i)$
- $h_j(t)$ = hazard rate for group j at time t
- $\Psi_j(t)$ = loss hazard rate for group j at time t
- λ_j = exponential hazard rate for group j
- R = hazard ratio of group 2 to group 1 $\equiv$ (assumed constant) value of $\frac{h_2(t)}{h_1(t)}$
- m_j = median survival time for group j
- b = number of subintervals per time unit
- T = accrual time
- τ = follow-up time after accrual
- L_j = exponential loss rate for group j
- XL_j = input time point on loss curve for group j
- SL_j = input survivor function value that corresponds to XL_j
- mL_j = median survival time for group j
- r_i = rank for i th time point

Each survival curve can be specified in one of several ways.

- For exponential curves:
 - a single point $(X_j(1), S_j(1))$ on the curve
 - median survival time
 - hazard rate
 - hazard ratio (for curve 2, with respect to curve 1)
- For piecewise linear curves with proportional hazards:
 - a set of points $\{(X_1(1), S_1(1)), (X_1(2), S_1(2)), \dots\}$ (for curve 1)
 - hazard ratio (for curve 2, with respect to curve 1)
- For arbitrary piecewise linear curves:
 - a set of points $\{(X_j(1), S_j(1)), (X_j(2), S_j(2)), \dots\}$

A total of $M + 1$ evenly spaced time points $\{t_0 = 0, t_1, t_2, \dots, t_M = T + \tau\}$ are used in calculations, where

$$M = \text{floor}((T + \tau)b)$$

The hazard function is calculated for each survival curve at each time point. For an exponential curve, the (constant) hazard is given by one of the following, depending on the input parameterization:

$$h_j(t) = \begin{cases} \lambda_j \\ \lambda_1 R \\ \frac{-\log(\frac{1}{2})}{m_j} \\ \frac{-\log(S_j(1))}{\bar{X}_j(1)} \\ \frac{-\log(S_1(1))}{\bar{X}_1(1)} R \end{cases}$$

For a piecewise linear curve, define the following additional notation:

t_i^- = largest input time X such that $X \leq t_i$

t_i^+ = smallest input time X such that $X > t_i$

The hazard is computed by using linear interpolation as follows:

$$h_j(t_i) = \frac{S_j(t_i^-) - S_j(t_i^+)}{[S_j(t_i^+) - S_j(t_i^-)][t_i - t_i^-] + S_j(t_i^-)[t_i^+ - t_i^-]}$$

With proportional hazards, the hazard rate of group 2's curve in terms of the hazard rate of group 1's curve is

$$h_2(t) = h_1(t)R$$

Hazard function values $\{\Psi_j(t_i)\}$ for the loss curves are computed in an analogous way from $\{L_j, XL_j, SL_j, mL_j\}$.

The expected number at risk $N_j(i)$ at time i in group j is calculated for each group and time points 0 through $M - 1$, as follows:

$$\begin{aligned} N_j(0) &= Nw_j \\ N_j(i+1) &= N_j(i) \left[1 - h_j(t_i) \left(\frac{1}{b} \right) - \Psi_j(t_i) \left(\frac{1}{b} \right) - \left(\frac{1}{b(T + \tau - t_i)} \right) 1_{\{t_i > \tau\}} \right] \end{aligned}$$

Define θ_i as the ratio of hazards and ϕ_i as the ratio of expected numbers at risk for time t_i :

$$\begin{aligned} \theta_i &= \frac{h_2(t_i)}{h_1(t_i)} \\ \phi_i &= \frac{N_2(i)}{N_1(i)} \end{aligned}$$

The expected number of deaths in each subinterval is calculated as follows:

$$D_i = [h_1(t_i)N_1(i) + h_2(t_i)N_2(i)] \left(\frac{1}{b} \right)$$

The rank values are calculated as follows according to which test statistic is used:

$$r_i = \begin{cases} 1, & \text{log-rank} \\ N_1(i) + N_2(i), & \text{Gehan} \\ \sqrt{N_1(i) + N_2(i)}, & \text{Tarone-Ware} \end{cases}$$

The distribution of the test statistic is approximated by $N(E, 1)$ where

$$E = \frac{\sum_{i=0}^{M-1} D_i r_i \left[\frac{\phi_i \theta_i}{1 + \phi_i \theta_i} - \frac{\phi_i}{1 + \phi_i} \right]}{\sqrt{\sum_{i=0}^{M-1} D_i r_i^2 \frac{\phi_i}{(1 + \phi_i)^2}}}$$

Note that $N^{\frac{1}{2}}$ can be factored out of the mean E , and so it can be expressed equivalently as

$$E = N^{\frac{1}{2}} E^* = N^{\frac{1}{2}} \left[\frac{\sum_{i=0}^{M-1} D_i^* r_i^* \left[\frac{\phi_i \theta_i}{1 + \phi_i \theta_i} - \frac{\phi_i}{1 + \phi_i} \right]}{\sqrt{\sum_{i=0}^{M-1} D_i^* r_i^{*2} \frac{\phi_i}{(1 + \phi_i)^2}}} \right]$$

where E^* is free of N and

$$\begin{aligned} D_i^* &= [h_1(t_i) N_1^*(i) + h_2(t_i) N_2^*(i)] \left(\frac{1}{b} \right) \\ r_i^* &= \begin{cases} 1, & \text{log-rank} \\ \frac{N_1^*(i) + N_2^*(i)}{\sqrt{N_1^*(i) + N_2^*(i)}}, & \text{Gehan} \\ \sqrt{N_1^*(i) + N_2^*(i)}, & \text{Tarone-Ware} \end{cases} \\ N_j^*(0) &= w_j \\ N_j^*(i+1) &= N_j^*(i) \left[1 - h_j(t_i) \left(\frac{1}{b} \right) - \Psi_j(t_i) \left(\frac{1}{b} \right) - \left(\frac{1}{b(T + \tau - t_i)} \right) 1_{\{t_i > \tau\}} \right] \end{aligned}$$

The approximate power is

$$\text{power} = \begin{cases} \Phi \left(-N^{\frac{1}{2}} E^* - z_{1-\alpha} \right), & \text{upper one-sided} \\ \Phi \left(N^{\frac{1}{2}} E^* - z_{1-\alpha} \right), & \text{lower one-sided} \\ \Phi \left(-N^{\frac{1}{2}} E^* - z_{1-\frac{\alpha}{2}} \right) + \Phi \left(N^{\frac{1}{2}} E^* - z_{1-\frac{\alpha}{2}} \right), & \text{two-sided} \end{cases}$$

Note that the upper and lower one-sided cases are expressed differently than in other analyses. This is because $E^* > 0$ corresponds to a higher survival curve in group 1 and thus, by the convention used in PROC power for two-group analyses, the lower side.

For the one-sided cases, a closed-form inversion of the power equation yield an approximate total sample size

$$N = \left(\frac{z_{\text{power}} + z_{1-\alpha}}{E^*} \right)^2$$

For the two-sided case, the solution for N is obtained by numerically inverting the power equation.

Accrual rates are converted to and from sample sizes according to the equation $a_j = n_j/T$, where a_j is the accrual rate for group j .

Expected numbers of events—that is, deaths, whether observed or censored—are converted to and from sample sizes according to the equation

$$e_j = \begin{cases} n_j [1 - S_j(\tau)], & T = 0 \\ n_j \left[1 - \frac{1}{T} \int_0^T S_j(T + \tau - t) dt \right], & T > 0 \end{cases}$$

where e_j is the expected number of events in group j . For an exponential curve, the equation simplifies to

$$e_j = \begin{cases} n_j [1 - \exp(-\lambda_j \tau)], & T = 0 \\ n_j \left[1 - \frac{1}{\lambda_j T} (\exp(-\lambda_j \tau) - \exp(-\lambda_j (T + \tau))) \right], & T > 0 \end{cases}$$

For a piecewise linear curve, first define K_j as the number of time points in the following collection: τ , $T + \tau$, and input time points for group j strictly between τ and $T + \tau$. Denote the ordered set of these points as $\{u_{j1}, \dots, u_{jK_j}\}$. The survival function values $S_j(\tau)$ and $S_j(T + \tau)$ are calculated by linear interpolation between adjacent input time points if they do not coincide with any input time points. Then the equation for a piecewise linear curve simplifies to

$$e_j = \begin{cases} n_j [1 - S_j(\tau)], & T = 0 \\ n_j \left[1 - \frac{1}{2T} \sum_{i=1}^{K_j-1} (u_{j,i+1} - u_{ji}) (S_j(u_{ji}) + S_j(u_{j,i+1})) \right], & T > 0 \end{cases}$$

Analyses in the TWOSAMPLEWILCOXON Statement

Wilcoxon-Mann-Whitney Test for Comparing Two Distributions (TEST=WMW)

The power approximation in this section is applicable to the Wilcoxon-Mann-Whitney (WMW) test as invoked with the WILCOXON option in the PROC NPAR1WAY statement of the NPAR1WAY procedure. The approximation is based on O'Brien and Casteloe (2006) and an estimator called $\widehat{WMW}_{\text{odds}}$. See O'Brien and Casteloe (2006) for a definition of $\widehat{WMW}_{\text{odds}}$, which need not be derived in detail here for purposes of explaining the power formula.

Let Y_1 and Y_2 be independent observations from any two distributions that you want to compare using the WMW test. For purposes of deriving the asymptotic distribution of $\widehat{WMW}_{\text{odds}}$ (and consequently the power computation as well), these distributions must be formulated as ordered categorical (“ordinal”) distributions.

If a distribution is continuous, it can be discretized using a large number of categories with negligible loss of accuracy. Each nonordinal distribution is divided into b categories, where b is the value of the NBINS parameter, with breakpoints evenly spaced on the probability scale. That is, each bin contains an equal probability $1/b$ for that distribution. Then the breakpoints across both distributions are pooled to form a collection of C bins (heretofore called “categories”), and the probabilities of bin membership for each distribution are recalculated. The motivation for this method of binning is to avoid degenerate representations of the distributions—that is, small handfuls of large probabilities among mostly empty bins—as can be caused by something like an evenly spaced grid across raw values rather than probabilities.

After the discretization process just mentioned, there are now two ordinal distributions, each with a set of probabilities across a common set of C ordered categories. For simplicity of notation, assume (without loss of generality) the response values to be $1, \dots, C$. Represent the conditional probabilities as

$$\tilde{p}_{ij} = \text{Prob}(Y_i = j \mid \text{group} = i), i \in \{1, 2\} \quad \text{and} \quad j \in \{1, \dots, C\}$$

and the group allocation weights as

$$w_i = \frac{n_i}{N} = \text{Prob}(\text{group} = i), \quad i \in \{1, 2\}$$

The joint probabilities can then be calculated simply as

$$p_{ij} = \text{Prob}(\text{group} = i, Y_i = j) = w_i \tilde{p}_{ij}, i \in \{1, 2\} \quad \text{and} \quad j \in \{1, \dots, C\}$$

The next step in the power computation is to compute the probabilities that a randomly chosen pair of observations from the two groups is concordant, discordant, or tied. It is useful to define these probabilities as functions of the terms Rs_{ij} and Rd_{ij} , defined as follows, where Y is a random observation drawn from the joint distribution across groups and categories:

$$\begin{aligned} Rs_{ij} &= \text{Prob}(Y \text{ is concordant with cell}(i, j)) + \frac{1}{2}\text{Prob}(Y \text{ is tied with cell}(i, j)) \\ &= \text{Prob}((\text{group} < i \text{ and } Y < j) \text{ or } (\text{group} > i \text{ and } Y > j)) + \\ &\quad \frac{1}{2}\text{Prob}(\text{group} \neq i \text{ and } Y = j) \\ &= \sum_{g=1}^2 \sum_{c=1}^C w_g \tilde{p}_{gc} \left[I_{(g-i)(c-j)>0} + \frac{1}{2}I_{g \neq i, c=j} \right] \end{aligned}$$

and

$$\begin{aligned} Rd_{ij} &= \text{Prob}(Y \text{ is discordant with cell}(i, j)) + \frac{1}{2}\text{Prob}(Y \text{ is tied with cell}(i, j)) \\ &= \text{Prob}((\text{group} < i \text{ and } Y > j) \text{ or } (\text{group} > i \text{ and } Y < j)) + \\ &\quad \frac{1}{2}\text{Prob}(\text{group} \neq i \text{ and } Y = j) \\ &= \sum_{g=1}^2 \sum_{c=1}^C w_g \tilde{p}_{gc} \left[I_{(g-i)(c-j)<0} + \frac{1}{2}I_{g \neq i, c=j} \right] \end{aligned}$$

For an independent random draw Y_1, Y_2 from the two distributions,

$$\begin{aligned} P_c &= \text{Prob}(Y_1, Y_2 \text{ concordant}) + \frac{1}{2}\text{Prob}(Y_1, Y_2 \text{ tied}) \\ &= \sum_{i=1}^2 \sum_{j=1}^C w_i \tilde{p}_{ij} Rs_{ij} \end{aligned}$$

and

$$\begin{aligned} P_d &= \text{Prob}(Y_1, Y_2 \text{ discordant}) + \frac{1}{2}\text{Prob}(Y_1, Y_2 \text{ tied}) \\ &= \sum_{i=1}^2 \sum_{j=1}^C w_i \tilde{p}_{ij} Rd_{ij} \end{aligned}$$

Then

$$\text{WMW}_{\text{odds}} = \frac{P_c}{P_d}$$

Proceeding to compute the theoretical standard error associated with WMW_{odds} (that is, the population analogue to the sample standard error),

$$\text{SE}(\text{WMW}_{\text{odds}}) = \frac{2}{P_d} \left[\sum_{i=1}^2 \sum_{j=1}^C w_i \tilde{p}_{ij} (\text{WMW}_{\text{odds}} Rd_{ij} - Rs_{ij})^2 / N \right]^{\frac{1}{2}}$$

Converting to the natural log scale and using the delta method,

$$SE(\log(WMW_{\text{odds}})) = \frac{SE(WMW_{\text{odds}})}{WMW_{\text{odds}}}$$

The next step is to produce a “smoothed” version of the $2 \times C$ cell probabilities that conforms to the null hypothesis of the Wilcoxon-Mann-Whitney test (in other words, independence in the $2 \times C$ contingency table of probabilities). Let $SE_{H_0}(\log(WMW_{\text{odds}}))$ denote the theoretical standard error of $\log(WMW_{\text{odds}})$ assuming H_0 .

Finally, compute the power using the noncentral chi-square and normal distributions:

$$\text{power} = \begin{cases} P\left(Z \geq \frac{SE_{H_0}(\log(WMW_{\text{odds}}))}{SE(\log(WMW_{\text{odds}}))} z_{1-\alpha} - \delta^* N^{\frac{1}{2}}\right), & \text{upper one-sided} \\ P\left(Z \leq \frac{SE_{H_0}(\log(WMW_{\text{odds}}))}{SE(\log(WMW_{\text{odds}}))} z_{\alpha} - \delta^* N^{\frac{1}{2}}\right), & \text{lower one-sided} \\ P\left(\chi^2(1, (\delta^*)^2 N) \geq \left[\frac{SE_{H_0}(\log(WMW_{\text{odds}}))}{SE(\log(WMW_{\text{odds}}))}\right]^2 \chi^2_{1-\alpha}(1)\right), & \text{two-sided} \end{cases}$$

where

$$\delta^* = \frac{\log(WMW_{\text{odds}})}{N^{\frac{1}{2}} SE(\log(WMW_{\text{odds}}))}$$

is the primary noncentrality—that is, the “effect size” that quantifies how much the two conjectured distributions differ. Z is a standard normal random variable, $\chi^2(df, nc)$ is a noncentral χ^2 random variable with degrees of freedom df and noncentrality nc , and N is the total sample size.

ODS Graphics

Statistical procedures use ODS Graphics to create graphs as part of their output. ODS Graphics is described in detail in Chapter 21, “[Statistical Graphics Using ODS](#).”

Before you create graphs, ODS Graphics must be enabled (for example, by specifying the ODS GRAPHICS ON statement). For more information about enabling and disabling ODS Graphics, see the section “[Enabling and Disabling ODS Graphics](#)” on page 607 in Chapter 21, “[Statistical Graphics Using ODS](#).”

The overall appearance of graphs is controlled by ODS styles. Styles and other aspects of using ODS Graphics are discussed in the section “[A Primer on ODS Statistical Graphics](#)” on page 606 in Chapter 21, “[Statistical Graphics Using ODS](#).”

If ODS Graphics is not enabled, then PROC POWER creates traditional graphics.

You can reference every graph produced through ODS Graphics with a name. The names of the graphs that PROC POWER generates are listed in [Table 90.37](#), along with the required statements and options.

Table 90.37 Graphs Produced by PROC POWER

ODS Graph Name	Plot Description	Option
PowerPlot	Plot with two of the following three parameters on the X and Y axes: power, sample size, and effect size	PLOT
PowerAbort	Empty plot that shows an error message when a plot could not be produced	PLOT

Examples: POWER Procedure

Example 90.1: One-Way ANOVA

This example deals with the same situation as in [Example 49.1](#) of Chapter 49, “The GLMPower Procedure.”

Hocking (1985, p. 109) describes a study of the effectiveness of electrolytes in reducing lactic acid buildup for long-distance runners. You are planning a similar study in which you will allocate five different fluids to runners on a 10-mile course and measure lactic acid buildup immediately after the run. The fluids consist of water and two commercial electrolyte drinks, EZDure and LactoZap, each prepared at two concentrations, low (EZD1 and LZ1) and high (EZD2 and LZ2).

You conjecture that the standard deviation of lactic acid measurements given any particular fluid is about 3.75, and that the expected lactic acid values will correspond roughly to those in [Table 90.38](#). You are least familiar with the LZ1 drink and hence decide to consider a range of reasonable values for that mean.

Table 90.38 Mean Lactic Acid Buildup by Fluid

Water	EZD1	EZD2	LZ1	LZ2
35.6	33.7	30.2	29 or 28	25.9

You are interested in four different comparisons, shown in [Table 90.39](#) with appropriate contrast coefficients.

Table 90.39 Planned Comparisons

Comparison	Contrast Coefficients				
	Water	EZD1	EZD2	LZ1	LZ2
Water versus electrolytes	4	−1	−1	−1	−1
EZD versus LZ	0	1	1	−1	−1
EZD1 versus EZD2	0	1	−1	0	0
LZ1 versus LZ2	0	0	0	1	−1

For each of these contrasts you want to determine the sample size required to achieve a power of 0.9 for detecting an effect with magnitude in accord with [Table 90.38](#). You are not yet attempting to choose a single sample size for the study, but rather checking the range of sample sizes needed for individual contrasts. You plan to test each contrast at $\alpha = 0.025$. In the interests of reducing costs, you will provide twice as many runners with water as with any of the electrolytes; in other words, you will use a sample size weighting scheme of 2:1:1:1:1. Use the `ONEWAYANOVA` statement in the `POWER` procedure to compute the sample sizes.

The statements required to perform this analysis are as follows:

```
proc power;
  onewayanova
    groupmeans = 35.6 | 33.7 | 30.2 | 29 28 | 25.9
    stddev = 3.75
    groupweights = (2 1 1 1 1)
    alpha = 0.025
    ntotal = .
    power = 0.9
    contrast = (4 -1 -1 -1 -1) (0 1 1 -1 -1)
               (0 1 -1 0 0) (0 0 0 1 -1);
run;
```

The `NTOTAL=` option with a missing value (.) indicates total sample size as the result parameter. The `GROUPMEANS=` option with values from [Table 90.38](#) specifies your conjectures for the means. With only one mean varying (the LZ1 mean), the “crossed” notation is simpler, showing scenarios for each group mean, separated by vertical bars (|). For more information about crossed and matched notations for grouped values, see the section “[Specifying Value Lists in Analysis Statements](#)” on page 7329. The contrasts in [Table 90.39](#) are specified with the `CONTRAST=` option, by using the “matched” notation with each contrast enclosed in parentheses. The `STDDEV=`, `ALPHA=`, and `POWER=` options specify the error standard deviation, significance level, and power. The `GROUPWEIGHTS=` option specifies the weighting schemes. Default values for the `NULLCONTRAST=` and `SIDES=` options specify a two-sided t test of the contrast equal to 0. See [Output 90.1.1](#) for the results.

Output 90.1.1 Sample Sizes for One-Way ANOVA Contrasts

The POWER Procedure
Single DF Contrast in One-Way ANOVA

Fixed Scenario Elements	
Method	Exact
Alpha	0.025
Standard Deviation	3.75
Group Weights	2 1 1 1 1
Nominal Power	0.9
Number of Sides	2
Null Contrast Value	0

Output 90.1.1 continued

Computed N Total											
Index	Contrast					Means					Actual Power Total
1	4	-1	-1	-1	-1	35.6	33.7	30.2	29	25.9	0.947 30
2	4	-1	-1	-1	-1	35.6	33.7	30.2	28	25.9	0.901 24
3	0	1	1	-1	-1	35.6	33.7	30.2	29	25.9	0.929 60
4	0	1	1	-1	-1	35.6	33.7	30.2	28	25.9	0.922 48
5	0	1	-1	0	0	35.6	33.7	30.2	29	25.9	0.901 174
6	0	1	-1	0	0	35.6	33.7	30.2	28	25.9	0.901 174
7	0	0	0	1	-1	35.6	33.7	30.2	29	25.9	0.902 222
8	0	0	0	1	-1	35.6	33.7	30.2	28	25.9	0.902 480

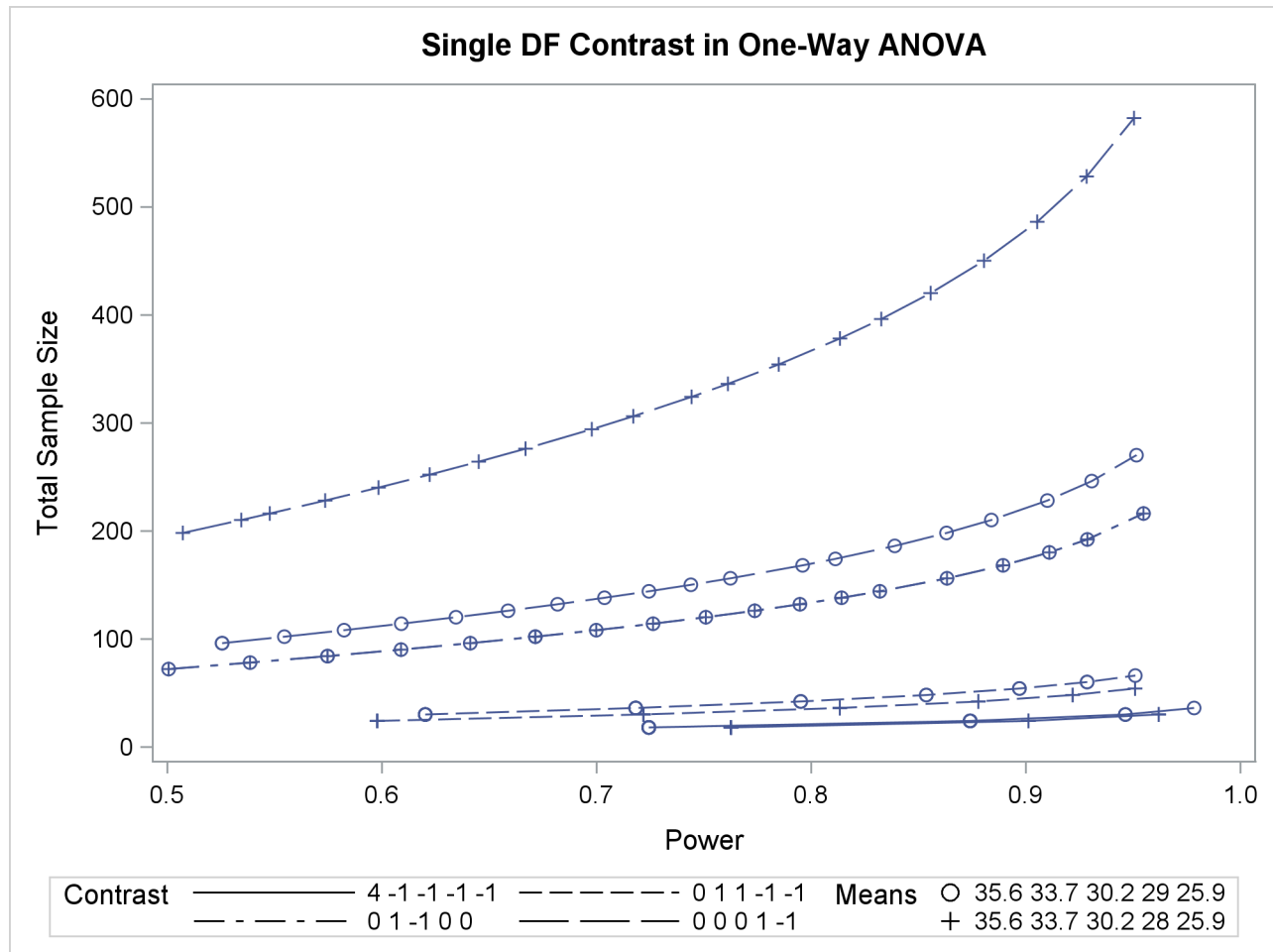
The sample sizes in [Output 90.1.1](#) range from 24 for the comparison of water versus electrolytes to 480 for the comparison of LZ1 versus LZ2, both assuming the smaller LZ1 mean. The sample size for the latter comparison is relatively large because the small mean difference of $28 - 25.9 = 2.1$ is hard to detect.

The Nominal Power of 0.9 in the “Fixed Scenario Elements” table in [Output 90.1.1](#) represents the input target power, and the Actual Power column in the “Computed N Total” table is the power at the sample size (N Total) adjusted to achieve the specified sample weighting. Note that all of the sample sizes are rounded up to multiples of 6 to preserve integer group sizes (since the group weights add up to 6). You can use the [NFRACTIONAL](#) option in the [ONEWAYANOVA](#) statement to compute raw fractional sample sizes.

Suppose you want to plot the required sample size for the range of power values from 0.5 to 0.95. First, define the analysis by specifying the same statements as before, but add the [PLOTONLY](#) option to the [PROC POWER](#) statement to disable the nongraphical results. Next, specify the [PLOT](#) statement with [X=POWER](#) to request a plot with power on the X axis. (The result parameter, here sample size, is always plotted on the other axis.) Use the [MIN=](#) and [MAX=](#) options in the [PLOT](#) statement to specify the power range. The following statements produce the plot shown in [Output 90.1.2](#).

```
ods graphics on;

proc power plotonly;
  onewayanova
    groupmeans = 35.6 | 33.7 | 30.2 | 29 28 | 25.9
    stddev = 3.75
    groupweights = (2 1 1 1 1)
    alpha = 0.025
    ntotal = .
    power = 0.9
    contrast = (4 -1 -1 -1 -1) (0 1 1 -1 -1)
              (0 1 -1 0 0) (0 0 0 1 -1);
  plot x=power min=.5 max=.95;
run;
```

Output 90.1.2 Plot of Sample Size versus Power for One-Way ANOVA Contrasts

In [Output 90.1.2](#), the line style identifies the contrast, and the plotting symbol identifies the group means scenario. The plot shows that the required sample size is highest for the (0 0 0 1 -1) contrast, which corresponds to the test of LZ1 versus LZ2 that was previously found to require the most resources, in either cell means scenario.

Note that some of the plotted points in [Output 90.1.2](#) are unevenly spaced. This is because the plotted points are the *rounded* sample size results at their corresponding *actual* power levels. The range specified with the **MIN=** and **MAX=** values in the **PLOT** statement corresponds to *nominal* power levels. In some cases, actual power is substantially higher than nominal power. To obtain plots with evenly spaced points (but with *fractional* sample sizes at the computed points), you can use the **NFRACTIONAL** option in the analysis statement preceding the **PLOT** statement.

Finally, suppose you want to plot the power for the range of sample sizes you will likely consider for the study (the range of 24 to 480 that achieves 0.9 power for different comparisons). In the **ONEWAYANOVA** statement, identify power as the result (**POWER=.**), and specify **NTOTAL=24**. The following statements produce the plot:

```
proc power plotonly;
  onewayanova
    groupmeans = 35.6 | 33.7 | 30.2 | 29 28 | 25.9
    stddev = 3.75
    groupweights = (2 1 1 1 1)
    alpha = 0.025
    ntotal = 24
    power = .
    contrast = (4 -1 -1 -1 -1) (0 1 1 -1 -1)
              (0 1 -1 0 0) (0 0 0 1 -1);
  plot x=n min=24 max=480;
run;

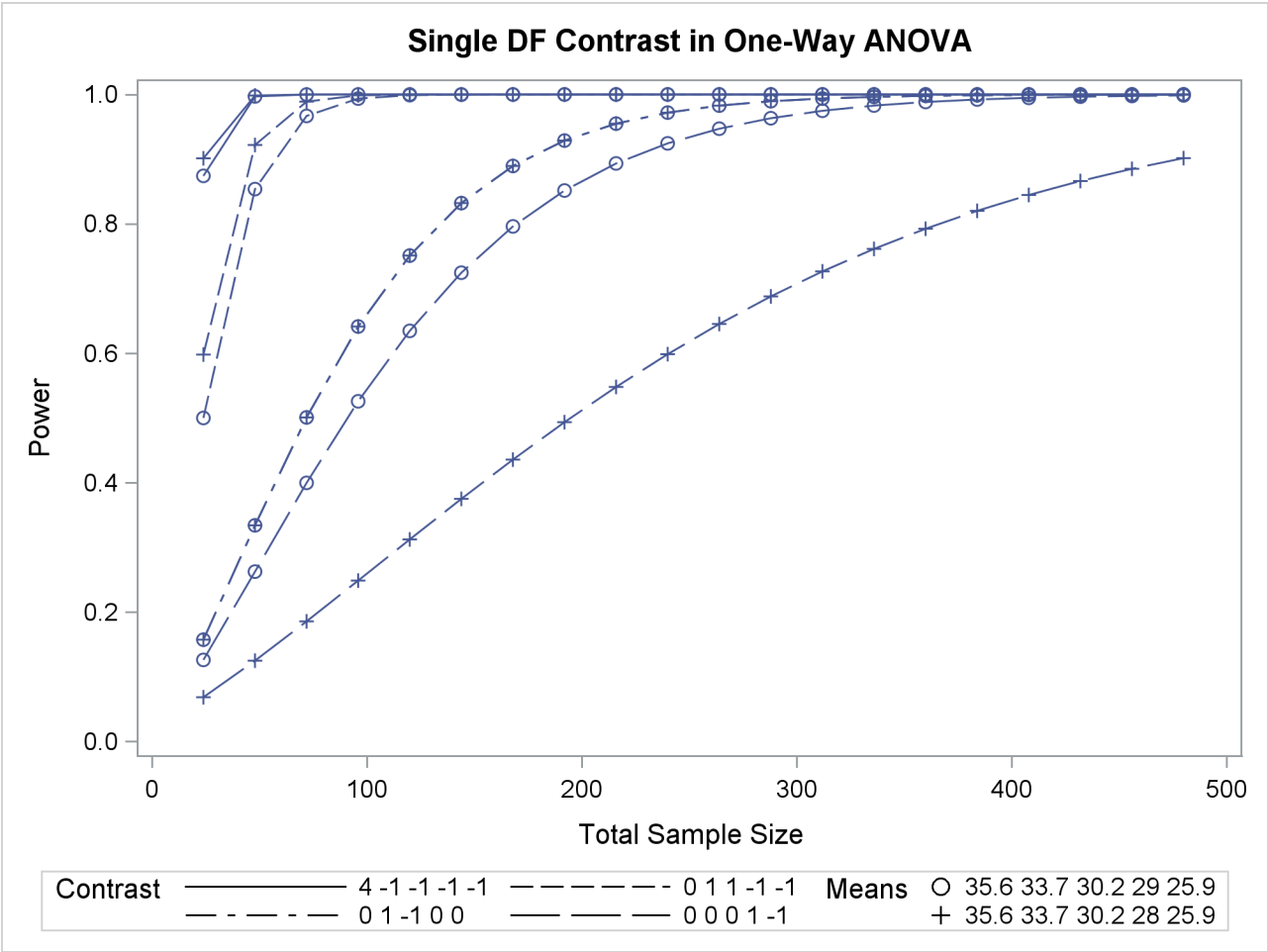
ods graphics off;
```

The **X=N** option in the **PLOT** statement requests a plot with sample size on the X axis.

Note that the value specified with the **NTOTAL=24** option is not used. It is overridden in the plot by the **MIN=** and **MAX=** options in the **PLOT** statement, and the **PLOTONLY** option in the **PROC POWER** statement disables nongraphical results. But the **NTOTAL=** option (along with a value) is still needed in the **ONEWAYANOVA** statement as a placeholder, to identify the desired parameterization for sample size.

Output 90.1.3 shows the resulting plot.

Output 90.1.3 Plot of Power versus Sample Size for One-Way ANOVA Contrasts



Although [Output 90.1.2](#) and [Output 90.1.3](#) surface essentially the same computations for practical power ranges, they each provide a different quick visual assessment. [Output 90.1.2](#) reveals the range of required sample sizes for powers of interest, and [Output 90.1.3](#) reveals the range of achieved powers for sample sizes of interest.

Example 90.2: The Sawtooth Power Function in Proportion Analyses

For many common statistical analyses, the power curve is monotonically increasing: the more samples you take, the more power you achieve. However, in statistical analyses of discrete data, such as tests of proportions, the power curve is often nonmonotonic. A small increase in sample size can result in a *decrease* in power, a decrease that is sometimes substantial. The explanation is that the actual significance level (in other words, the achieved Type I error rate) for discrete tests strays below the target level and varies with sample size. The power loss from a decrease in the Type I error rate can outweigh the power gain from an increase in sample size. The example discussed here demonstrates this “sawtooth” phenomenon. For additional discussion on the topic, see Chernick and Liu (2002).

Suppose you have a new scheduling system for an airline, and you want to determine how many flights you must observe to have at least an 80% chance of establishing an improvement in the proportion of late arrivals

on a specific travel route. You will use a one-sided exact binomial proportion test with a null proportion of 30%, the frequency of late arrivals under the previous scheduling system, and a nominal significance level of $\alpha = 0.05$. Well-supported predictions estimate the new late arrival rate to be about 20%, and you will base your sample size determination on this assumption.

The POWER procedure does not currently compute exact sample size directly for the exact binomial test. But you can get an initial estimate by computing the approximate sample size required for a z test. Use the `ONESAMPLEFREQ` statement in the POWER procedure with `TEST=Z` and `METHOD=NORMAL` to compute the approximate sample size to achieve a power of 0.8 by using the z test. The following statements perform the analysis:

```
proc power;
  onesamplefreq test=z method=normal
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.3
    proportion     = 0.2
    ntotal         = .
    power          = 0.8;
run;
```

The `NTOTAL=` option with a missing value (.) indicates sample size as the result parameter. The `SIDES=1` option specifies a one-sided test. The `ALPHA=`, `NULLPROPORTION=`, and `POWER=` options specify the significance level of 0.05, null value of 0.3, and target power of 0.8, respectively. The `PROPORTION=` option specifies your conjecture of 0.3 for the true proportion.

Output 90.2.1 Approximate Sample Size for z Test of a Proportion

The POWER Procedure		
Z Test for Binomial Proportion		
Fixed Scenario Elements		
Method	Normal approximation	
Number of Sides	1	
Null Proportion	0.3	
Alpha	0.05	
Binomial Proportion	0.2	
Nominal Power	0.8	
Variance Estimate	Null Variance	
	Computed N	
	Total	
	Actual	N
	Power	Total
	0.800	119

The results, shown in [Output 90.2.1](#), indicate that you need to observe about $N = 119$ flights to have an 80% chance of rejecting the hypothesis of a late arrival proportion of 30% or higher, if the true proportion is 20%, by using the z test. A similar analysis ([Output 90.2.2](#)) reveals an approximate sample size of $N = 129$ for the z test with continuity correction, which is performed by using `TEST=ADJZ`:

```

proc power;
  onesamplefreq test=adjz method=normal
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.3
    proportion     = 0.2
    ntotal         = .
    power          = 0.8;
run;

```

Output 90.2.2 Approximate Sample Size for z Test with Continuity Correction

The POWER Procedure
Z Test for Binomial Proportion with Continuity Adjustment

Fixed Scenario Elements		
Method	Normal approximation	
Number of Sides		1
Null Proportion		0.3
Alpha		0.05
Binomial Proportion		0.2
Nominal Power		0.8
Variance Estimate	Null Variance	

Computed N		
Total		
Actual	N	
Power	Total	
0.801	129	

Based on the approximate sample size results, you decide to explore the power of the exact binomial test for sample sizes between 110 and 140. The following statements produce the plot:

```

ods graphics on;

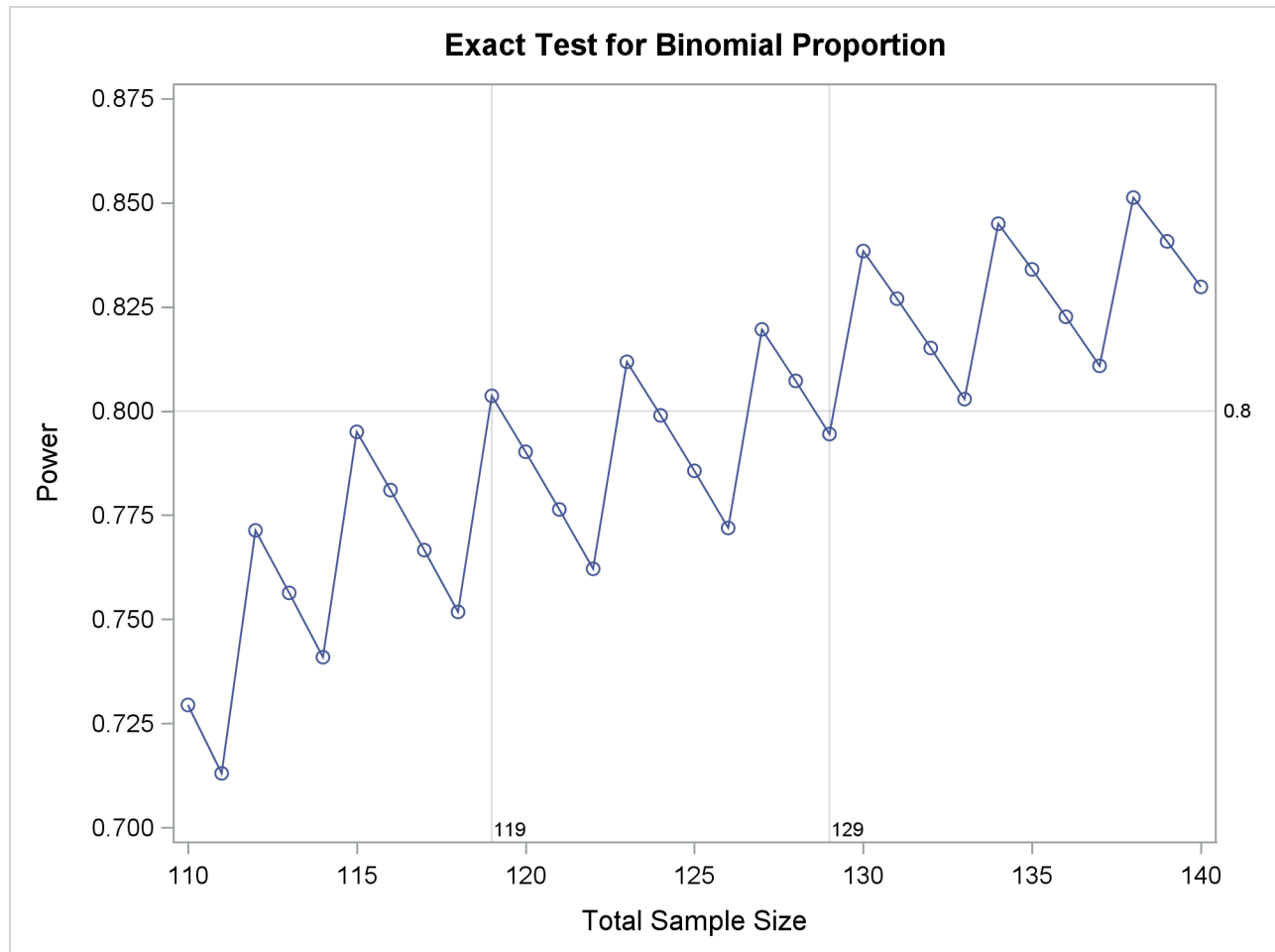
proc power plotonly;
  onesamplefreq test=exact
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.3
    proportion     = 0.2
    ntotal         = 119
    power          = .;
  plot x=n min=110 max=140 step=1
    yopts=(ref=.8) xopts=(ref=119 129);
run;

```

The **TEST=EXACT** option in the **ONESAMPLEFREQ** statement specifies the exact binomial test, and the missing value (.) for the **POWER=** option indicates power as the result parameter. The **PLOTONLY** option in the **PROC POWER** statement disables nongraphical output. The **PLOT** statement with **X=N** requests a plot with sample size on the X axis. The **MIN=** and **MAX=** options in the **PLOT** statement specify the sample size range. The **YOPTS=(REF=)** and **XOPTS=(REF=)** options add reference lines to highlight the approximate sample size results. The **STEP=1** option produces a point at each integer sample size. The sample size value

specified with the `NTOTAL=` option in the `ONESAMPLEFREQ` statement is overridden by the `MIN=` and `MAX=` options in the `PLOT` statement. Output 90.2.3 shows the resulting plot.

Output 90.2.3 Plot of Power versus Sample Size for Exact Binomial Test



Note the sawtooth pattern in Output 90.2.3. Although the power surpasses the target level of 0.8 at $N = 119$, it decreases to 0.79 with $N = 120$ and further to 0.76 with $N = 122$ before rising again to 0.81 with $N = 123$. Not until $N = 130$ does the power stay above the 0.8 target. Thus, a more conservative sample size recommendation of 130 might be appropriate, depending on the precise goals of the sample size determination.

In addition to considering alternative sample sizes, you might also want to assess the sensitivity of the power to inaccuracies in assumptions about the true proportion. The following statements produce a plot including true proportion values of 0.18 and 0.22. They are identical to the previous statements except for the additional true proportion values specified with the `PROPORTION=` option in the `ONESAMPLEFREQ` statement.

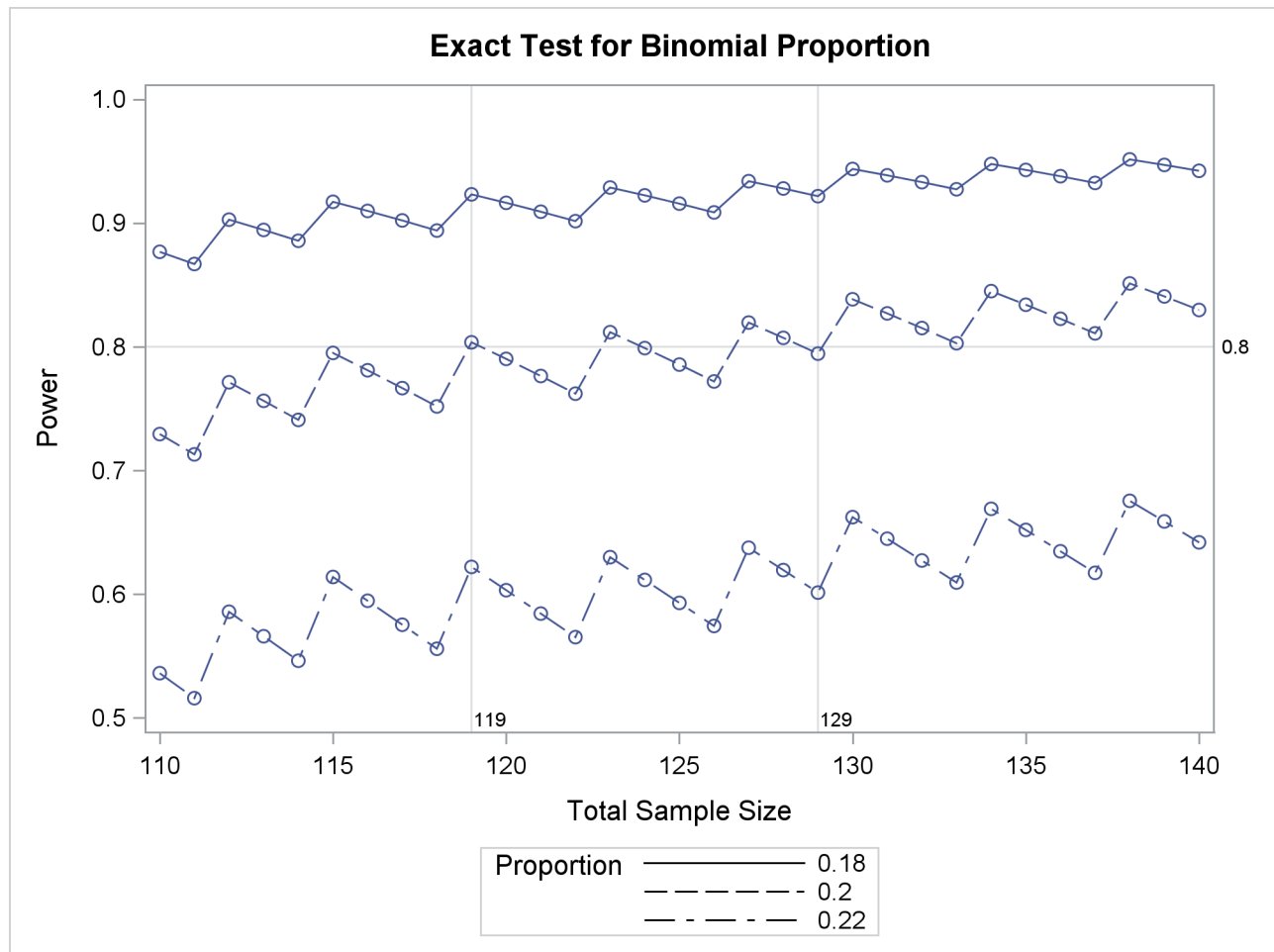
```

proc power plotonly;
  onesamplefreq test=exact
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.3
    proportion      = 0.18 0.2 0.22
    ntotal         = 119
    power          = .;
  plot x=n min=110 max=140 step=1
    yopts=(ref=.8) xopts=(ref=119 129);
run;

```

Output 90.2.4 shows the resulting plot.

Output 90.2.4 Plot for Assessing Sensitivity to True Proportion Value



The plot reveals a dramatic sensitivity to the true proportion value. For $N=119$, the power is about 0.92 if the true proportion is 0.18, and as low as 0.62 if the proportion is 0.22. Note also that the power jumps occur at the same sample sizes in all three curves; the curves are only shifted and stretched vertically. This is because spikes and valleys in power curves are invariant to the true proportion value; they are due to changes in the critical value of the test.

A closer look at some ancillary output from the analysis sheds light on this property of the sawtooth pattern. You can add an ODS OUTPUT statement to save the plot content that corresponds to [Output 90.2.3](#) to a data set:

```
proc power plotonly;
  ods output plotcontent=PlotData;
  onesamplefreq test=exact
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.3
    proportion      = 0.2
    ntotal         = 119
    power          = .;
  plot x=n min=110 max=140 step=1
       yopts=(ref=.8) xopts=(ref=119 129);
run;
```

The PlotData data set contains parameter values for each point in the plot. The parameters include underlying characteristics of the putative test. The following statements print the critical value and actual significance level along with sample size and power:

```
proc print data=PlotData;
  var NTotal LowerCritVal Alpha Power;
run;
```

[Output 90.2.5](#) shows the plot data.

Output 90.2.5 Numerical Content of Plot

Obs	NTotal	LowerCritVal	Alpha	Power
1	110	24	0.0356	0.729
2	111	24	0.0313	0.713
3	112	25	0.0446	0.771
4	113	25	0.0395	0.756
5	114	25	0.0349	0.741
6	115	26	0.0490	0.795
7	116	26	0.0435	0.781
8	117	26	0.0386	0.767
9	118	26	0.0341	0.752
10	119	27	0.0478	0.804
11	120	27	0.0425	0.790
12	121	27	0.0377	0.776
13	122	27	0.0334	0.762
14	123	28	0.0465	0.812
15	124	28	0.0414	0.799
16	125	28	0.0368	0.786
17	126	28	0.0327	0.772
18	127	29	0.0453	0.820
19	128	29	0.0404	0.807
20	129	29	0.0359	0.794
21	130	30	0.0493	0.838
22	131	30	0.0441	0.827
23	132	30	0.0394	0.815
24	133	30	0.0351	0.803
25	134	31	0.0480	0.845
26	135	31	0.0429	0.834
27	136	31	0.0384	0.823
28	137	31	0.0342	0.811
29	138	32	0.0466	0.851
30	139	32	0.0418	0.841
31	140	32	0.0374	0.830

Note that whenever the critical value changes, the actual α jumps up to a value close to the nominal $\alpha = 0.05$, and the power also jumps up. Then while the critical value stays constant, the actual α and power slowly decrease. The critical value is independent of the true proportion value. So you can achieve a locally maximal power by choosing a sample size corresponding to a spike on the sawtooth curve, and this choice is locally optimal *regardless* of the unknown value of the true proportion. Locally optimal sample sizes in this case include 115, 119, 123, 127, 130, and 134.

As a point of interest, the power does not always jump sharply and decrease gradually. The shape of the sawtooth depends on the direction of the test and the location of the null proportion relative to 0.5. For example, if the direction of the hypothesis in this example is reversed (by switching true and null proportion values) so that the rejection region is in the upper tail, then the power curve exhibits sharp decreases and gradual increases. The following statements are similar to those producing the plot in [Output 90.2.3](#) but with values of the `PROPORTION=` and `NULLPROPORTION=` options switched:

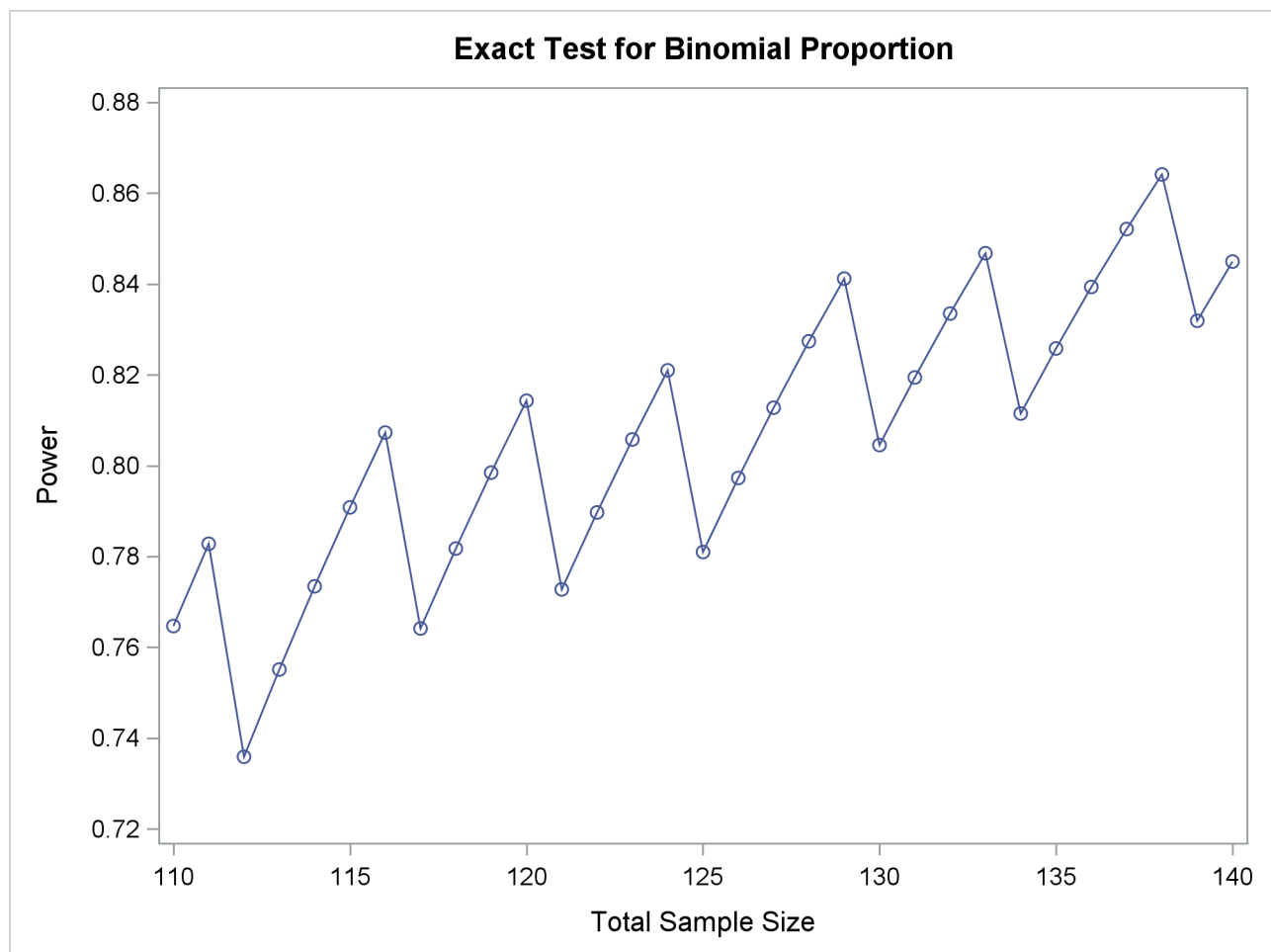
```

proc power plotonly;
  onesamplefreq test=exact
    sides          = 1
    alpha          = 0.05
    nullproportion = 0.2
    proportion     = 0.3
    ntotal         = 119
    power          = .;
  plot x=n min=110 max=140 step=1;
run;

```

The resulting plot is shown in [Output 90.2.6](#).

Output 90.2.6 Plot of Power versus Sample Size for Another One-sided Test



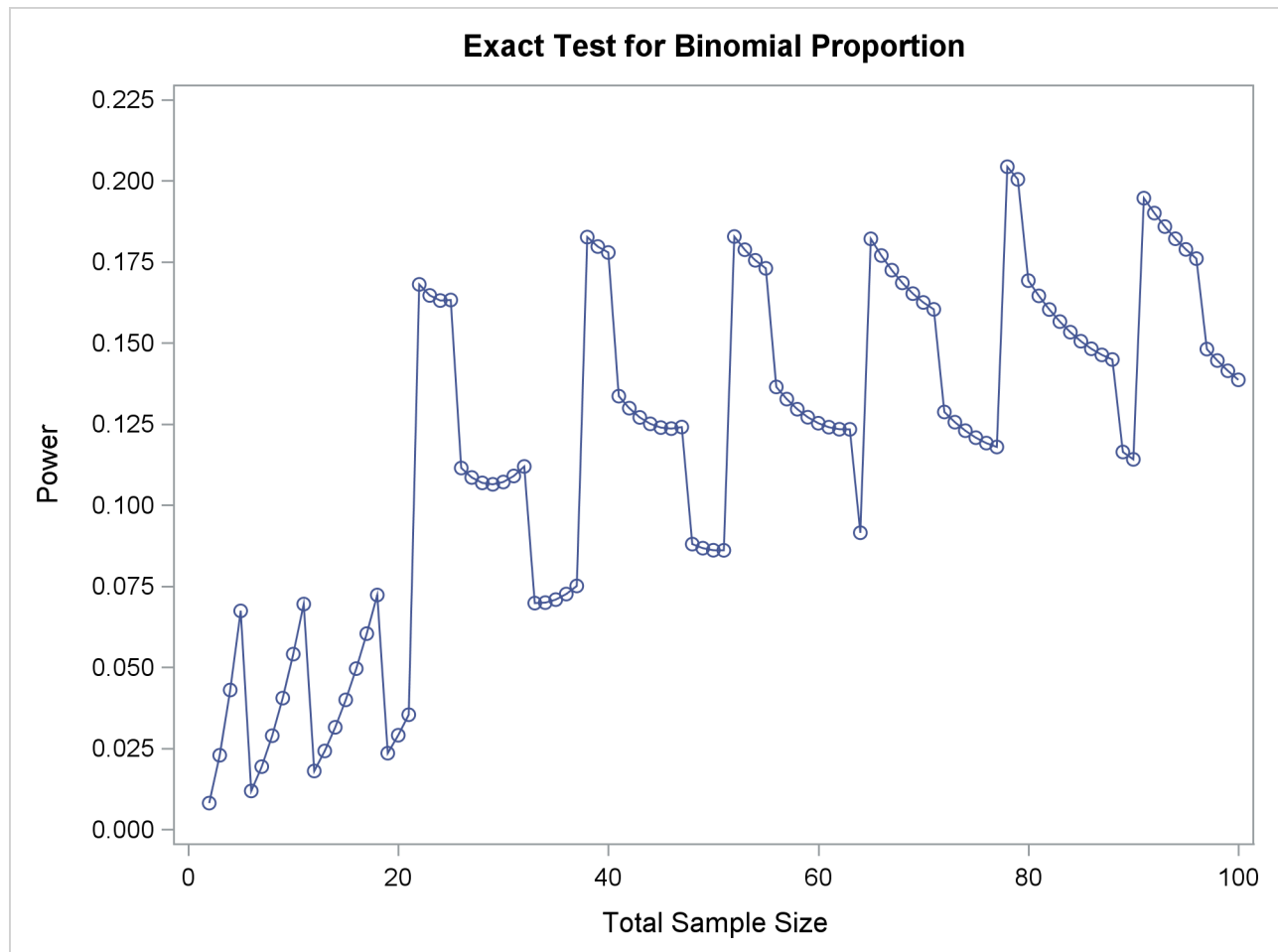
Finally, two-sided tests can lead to even more irregular power curve shapes, since changes in lower and upper critical values affect the power in different ways. The following statements produce a plot of power versus sample size for the scenario of a two-sided test with high alpha and a true proportion close to the null value:

```
proc power plotonly;
  onesamplefreq test=exact
    sides          = 2
    alpha          = 0.2
    nullproportion = 0.1
    proportion      = 0.09
    ntotal         = 10
    power          = .;
  plot x=n min=2 max=100 step=1;
run;

ods graphics off;
```

Output 90.2.7 shows the resulting plot.

Output 90.2.7 Plot of Power versus Sample Size for a Two-Sided Test



Due to the irregular shapes of power curves for proportion tests, the question “Which sample size should I use?” is often insufficient. A sample size solution produced directly in PROC POWER reveals the smallest possible sample size to achieve your target power. But as the examples in this section demonstrate, it is helpful to consult graphs for answers to questions such as the following:

- Which sample size will guarantee that all higher sample sizes also achieve my target power?
- Given a candidate sample size, can I increase it slightly to achieve locally maximal power, or perhaps even decrease it and get higher power?

Example 90.3: Simple AB/BA Crossover Designs

Crossover trials are experiments in which each subject is given a sequence of different treatments. They are especially common in clinical trials for medical studies. The reduction in variability from taking multiple measurements on a subject allows for more precise treatment comparisons. The simplest such design is the AB/BA crossover, in which each subject receives each of two treatments in a randomized order.

Under certain simplifying assumptions, you can test the treatment difference in an AB/BA crossover trial by using either a paired or two-sample t test (or equivalence test, depending on the hypothesis). This example will demonstrate when and how you can use the **PAIREDMEANS** statement in PROC POWER to perform power analyses for AB/BA crossover designs.

Senn (1993, Chapter 3) discusses a study comparing the effects of two bronchodilator medications in treatment of asthma, by using an AB/BA crossover design. Suppose you want to plan a similar study comparing two new medications, “Xilodol” and “Brantium.” Half of the patients would be assigned to sequence AB, getting a dose of Xilodol in the first treatment period, a wash-out period of one week, and then a dose of Brantium in the second treatment period. The other half would be assigned to sequence BA, following the same schedule but with the drugs reversed. In each treatment period you would administer the drugs in the morning and then measure peak expiratory flow (PEF) at the end of the day, with higher PEF representing better lung function.

You conjecture that the mean and standard deviation of PEF are about $\mu_A = 330$ and $\sigma_A = 40$ for Xilodol and $\mu_B = 310$ and $\sigma_B = 55$ for Brantium, and that each pair of measurements on the same subject will have a correlation of about 0.3. You want to compute the power of both one-sided and two-sided tests of mean difference, with a significance level of $\alpha = 0.01$, for a sample size of 100 patients and also plot the power for a range of 50 to 200 patients. Note that the allocation ratio of patients to the two sequences is irrelevant in this analysis.

The choice of statistical test depends on which assumptions are reasonable. One possibility is a t test. A paired or two-sample t test is valid when there is no carryover effect and no interactions between patients, treatments, and periods. See Senn (1993, Chapter 3) for more details. The choice between a paired or a two-sample test depends on what you assume about the period effect. If you assume no period effect, then a paired t test is the appropriate analysis for the design, with the first member of each pair being the Xilodol measurement (regardless of which sequence the patient belongs to). Otherwise, the two-sample t test approach is called for, since this analysis adjusts for the period effect by using an extra degree of freedom.

Suppose you assume no period effect. Then you can use the **PAIREDMEANS** statement in PROC POWER with the **TEST=DIFF** option to perform a sample size analysis for the paired t test. Indicate power as the result parameter by specifying the **POWER=** option with a missing value (.). Specify the conjectured means and standard deviations for each drug by using the **PAIREDMEANS=** and **PAIREDSTDDEVS=** options and the correlation by using the **CORR=** option. Specify both one- and two-sided tests by using the **SIDES=** option, the significance level by using the **ALPHA=** option, and the sample size (in terms of number of pairs) by using the **NPAIRS=** option. Generate a plot of power versus sample size by specifying the **PLOT** statement with **X=N** to request a plot with sample size on the X axis. (The result parameter, here power, is always plotted on the other axis.) Use the **MIN=** and **MAX=** options in the **PLOT** statement to specify the sample size range (as numbers of pairs).

The following statements perform the sample size analysis:

```
ods graphics on;

proc power;
  pairedmeans test=diff
    pairedmeans    = (330 310)
    pairedstddevs  = (40 55)
    corr           = 0.3
    sides          = 1 2
    alpha          = 0.01
    npairs         = 100
    power          = .;
  plot x=n min=50 max=200;
run;

ods graphics off;
```

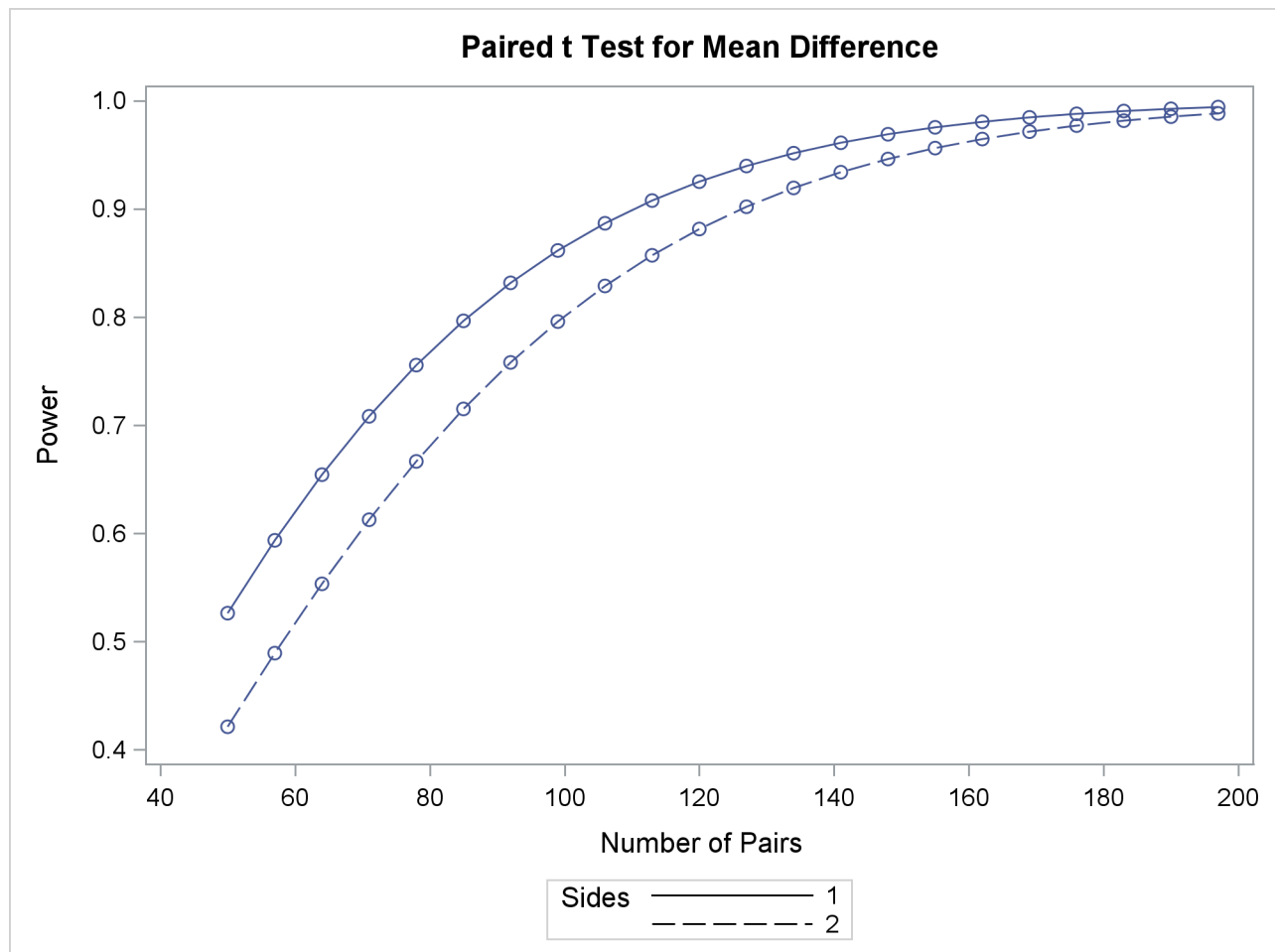
Default values for the **NULLDIFF=** and **DIST=** options specify a null mean difference of 0 and the assumption of normally distributed data. The output is shown in [Output 90.3.1](#) and [Output 90.3.2](#).

Output 90.3.1 Power for Paired *t* Analysis of Crossover Design

The POWER Procedure
Paired t Test for Mean Difference

Fixed Scenario Elements		
Distribution	Normal	
Method	Exact	
Alpha	0.01	
Mean 1	330	
Mean 2	310	
Standard Deviation 1	40	
Standard Deviation 2	55	
Correlation	0.3	
Number of Pairs	100	
Null Difference	0	

Computed Power		
Index	Sides	Power
1	1	0.865
2	2	0.801

Output 90.3.2 Plot of Power versus Sample Size for Paired t Analysis of Crossover Design

The “Computed Power” table in [Output 90.3.1](#) shows that the power with 100 patients is about 0.8 for the two-sided test and 0.87 for the one-sided test with the alternative of larger Brantium mean. In [Output 90.3.2](#), the line style identifies the number of sides of the test. The plotting symbols identify locations of actual computed powers; the curves are linear interpolations of these points. The plot demonstrates how much higher the power is in the one-sided test than in the two-sided test for the range of sample sizes.

Suppose now that instead of detecting a difference between Xilodol and Brantium, you want to establish that they are similar—in particular, that the absolute mean PEF difference is at most 35. You might consider this goal if, for example, one of the drugs has fewer side effects and if a difference of no more than 35 is considered clinically small. Instead of a standard t test, you would conduct an *equivalence test* of the treatment mean difference for the two drugs. You would test the hypothesis that the true difference is less than -35 or more than 35 against the alternative that the mean difference is between -35 and 35 , by using an additive model and a two one-sided tests (“TOST”) analysis.

Assuming no period effect, you can use the **PAIREDMEANS** statement with the **TEST=EQUIV_DIFF** option to perform a sample size analysis for the paired equivalence test. Indicate power as the result parameter by specifying the **POWER=** option with a missing value (.). Use the **LOWER=** and **UPPER=** options to specify the equivalence bounds of -35 and 35 . Use the **PAIREDMEANS=**, **PAIREDSTDDEVS=**, **CORR=**, and **ALPHA=** options in the same way as in the t test at the beginning of this example to specify the remaining parameters.

The following statements perform the sample size analysis:

```
proc power;
  pairedmeans test=equiv_add
    lower      = -35
    upper      = 35
    pairedmeans = (330 310)
    pairedstddevs = (40 55)
    corr       = 0.3
    alpha      = 0.01
    npairs     = 100
    power      = .;
run;
```

The default option **DIST=NORMAL** specifies an assumption of normally distributed data. The output is shown in [Output 90.3.3](#).

Output 90.3.3 Power for Paired Equivalence Test for Crossover Design

The POWER Procedure
Equivalence Test for Paired Mean Difference

Fixed Scenario Elements	
Distribution	Normal
Method	Exact
Lower Equivalence Bound	-35
Upper Equivalence Bound	35
Alpha	0.01
Reference Mean	330
Treatment Mean	310
Standard Deviation 1	40
Standard Deviation 2	55
Correlation	0.3
Number of Pairs	100
Computed Power	
Power	
0.598	

The power for the paired equivalence test with 100 patients is about 0.6.

Example 90.4: Noninferiority Test with Lognormal Data

The typical goal in noninferiority testing is to conclude that a new treatment or process or product is not appreciably worse than some standard. This is accomplished by convincingly rejecting a one-sided null hypothesis that the new treatment is appreciably worse than the standard. When designing such studies, investigators must define precisely what constitutes “appreciably worse.”

You can use the POWER procedure for sample size analyses for a variety of noninferiority tests, by specifying custom, one-sided null hypotheses for common tests. This example illustrates the strategy (often called

Blackwelder's scheme; Blackwelder 1982) by comparing the means of two independent lognormal samples. The logic applies to one-sample, two-sample, and paired-sample problems involving normally distributed measures and proportions.

Suppose you are designing a study hoping to show that a new (less expensive) manufacturing process does not produce appreciably more pollution than the current process. Quantifying “appreciably worse” as 10%, you seek to show that the mean pollutant level from the new process is less than 110% of that from the current process. In standard hypothesis testing notation, you seek to reject

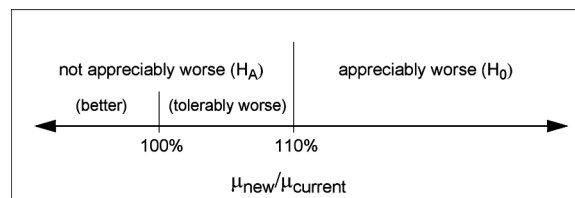
$$H_0: \frac{\mu_{\text{new}}}{\mu_{\text{current}}} \geq 1.10$$

in favor of

$$H_A: \frac{\mu_{\text{new}}}{\mu_{\text{current}}} < 1.10$$

This is described graphically in Figure 90.8. Mean ratios below 100% are better levels for the new process; a ratio of 100% indicates absolute equivalence; ratios of 100–110% are “tolerably” worse; and ratios exceeding 110% are appreciably worse.

Figure 90.8 Hypotheses for the Pollutant Study



An appropriate test for this situation is the common two-group *t* test on log-transformed data. The hypotheses become

$$H_0: \log(\mu_{\text{new}}) - \log(\mu_{\text{current}}) \geq \log(1.10)$$

$$H_A: \log(\mu_{\text{new}}) - \log(\mu_{\text{current}}) < \log(1.10)$$

Measurements of the pollutant level will be taken by using laboratory models of the two processes and will be treated as independent lognormal observations with a coefficient of variation (σ/μ) between 0.5 and 0.6 for both processes. You will end up with 300 measurements for the current process and 180 for the new one. It is important to avoid a Type I error here, so you set the Type I error rate to 0.01. Your theoretical work suggests that the new process will actually reduce the pollutant by about 10% (to 90% of current), but you need to compute and graph the power of the study if the new levels are actually between 70% and 120% of current levels.

Implement the sample size analysis by using the **TWOSAMPLEMEANS** statement in PROC POWER with the **TEST=RATIO** option. Indicate power as the result parameter by specifying the **POWER=** option with a missing value (.). Specify a series of scenarios for the mean ratio between 0.7 and 1.2 by using the **MEANRATIO=** option. Use the **NULLRATIO=** option to specify the null mean ratio of 1.10. Specify **SIDES=L** to indicate a one-sided test with the alternative hypothesis stating that the mean ratio is *lower* than the null value. Specify the significance level, scenarios for the coefficient of variation, and the group sample sizes by using the **ALPHA=**, **CV=**, and **GROUPNS=** options. Generate a plot of power versus mean ratio by specifying the **PLOT** statement with the **X=EFFECT** option to request a plot with mean ratio on the X axis.

(The result parameter, here power, is always plotted on the other axis.) Use the **STEP=** option in the **PLOT** statement to specify an interval of 0.05 between computed points in the plot.

The following statements perform the desired analysis:

```
ods graphics on;

proc power;
  twosamplemeans test=ratio
    meanratio = 0.7 to 1.2 by 0.1
    nullratio = 1.10
    sides      = L
    alpha      = 0.01
    cv         = 0.5 0.6
    groupns    = (300 180)
    power      = .;
  plot x=effect step=0.05;
run;

ods graphics off;
```

Note the use of **SIDES=L**, which forces computations for cases that need a rejection region that is opposite to the one providing the most one-tailed power; in this case, it is the lower tail. Such cases will show power that is less than the prescribed Type I error rate. The default option **DIST=LOGNORMAL** specifies the assumption of lognormally distributed data. The default **MIN=** and **MAX=** options in the plot statement specify an X axis range identical to the effect size range in the **TWOSAMPLEMEANS** statement (mean ratios between 0.7 and 1.2).

Output 90.4.1 and Output 90.4.2 show the results.

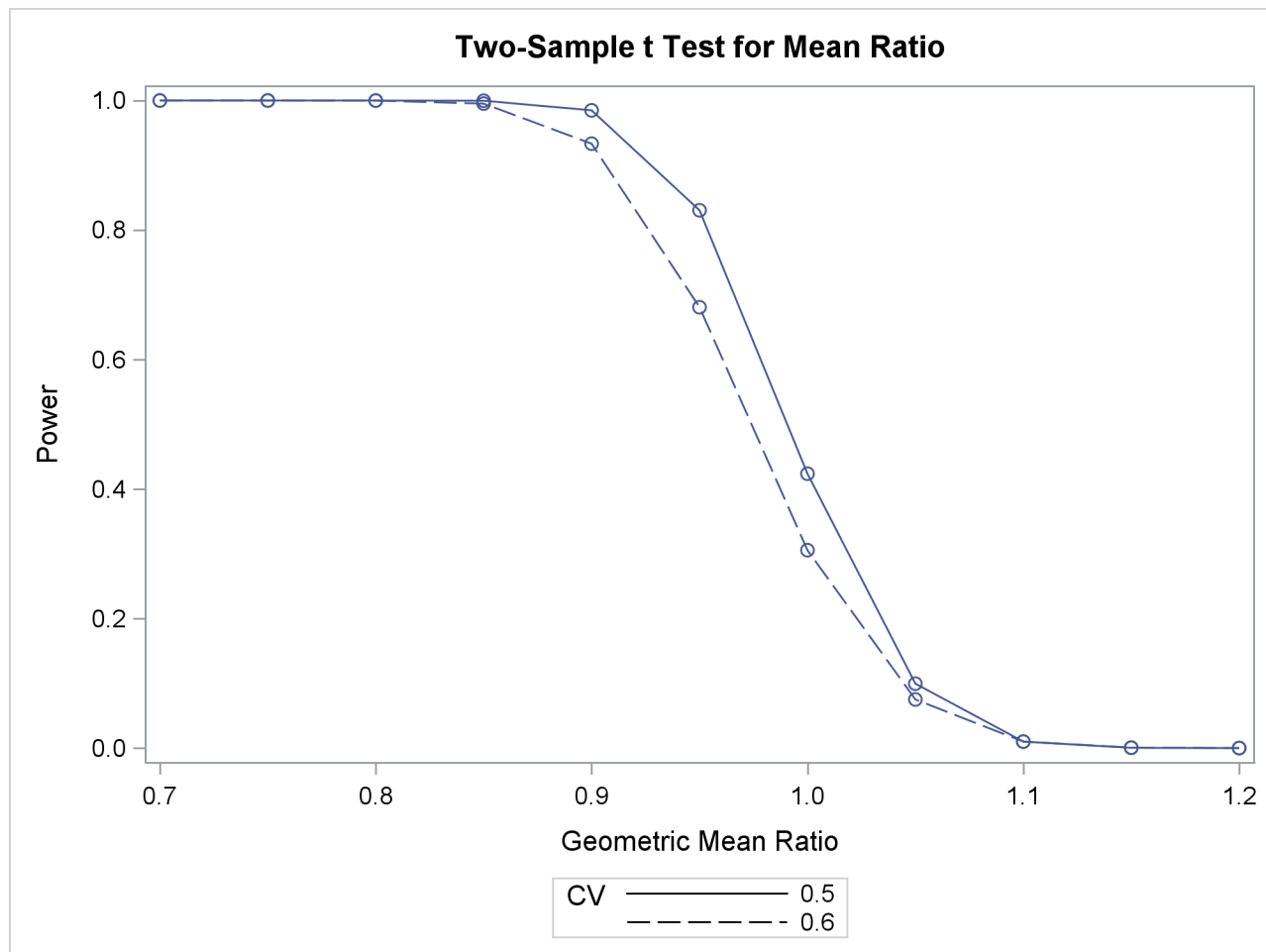
Output 90.4.1 Power for Noninferiority Test of Ratio

The POWER Procedure Two-Sample t Test for Mean Ratio

Fixed Scenario Elements	
Distribution	Lognormal
Method	Exact
Number of Sides	L
Null Geometric Mean Ratio	1.1
Alpha	0.01
Group 1 Sample Size	300
Group 2 Sample Size	180

Output 90.4.1 continued

Computed Power			
Geo			
Mean			
Index	Ratio	CV	Power
1	0.7	0.5	>.999
2	0.7	0.6	>.999
3	0.8	0.5	>.999
4	0.8	0.6	>.999
5	0.9	0.5	0.985
6	0.9	0.6	0.933
7	1.0	0.5	0.424
8	1.0	0.6	0.306
9	1.1	0.5	0.010
10	1.1	0.6	0.010
11	1.2	0.5	<.001
12	1.2	0.6	<.001

Output 90.4.2 Plot of Power versus Mean Ratio for Noninferiority Test

The “Computed Power” table in [Output 90.4.1](#) shows that power exceeds 0.90 if the true mean ratio is 90% or less, as surmised. But power is unacceptably low (0.31–0.42) if the processes happen to be truly equivalent. Note that the power is identical to the alpha level (0.01) if the true mean ratio is 1.10 and below 0.01 if the true mean ratio is appreciably worse (>110%). In [Output 90.4.2](#), the line style identifies the coefficient of variation. The plotting symbols identify locations of actual computed powers; the curves are linear interpolations of these points.

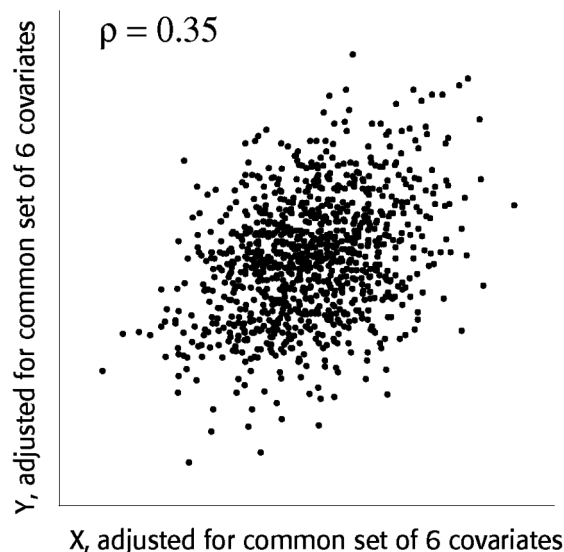
Example 90.5: Multiple Regression and Correlation

You are working with a team of preventive cardiologists investigating whether elevated serum homocysteine levels are linked to atherosclerosis (plaque buildup) in coronary arteries. The planned analysis is an ordinary least squares regression to assess the relationship between total homocysteine level (tHcy) and a plaque burden index (PBI), adjusting for six other variables: age, gender, plasma levels of folate, vitamin B₆, vitamin B₁₂, and a serum cholesterol index. You will regress PBI on tHcy and the six other predictors (plus the intercept) and use a Type III *F* test to assess whether tHcy is a significant predictor after adjusting for the others. You wonder whether 100 subjects will provide adequate statistical power.

This is a correlational study at a single time. Subjects will be screened so that about half will have had a heart problem. All eight variables will be measured during one visit. Most clinicians are familiar with simple correlations between two variables, so you decide to pose the statistical problem in terms of estimating and testing the partial correlation between $X_1 = \text{tHcy}$ and $Y = \text{PBI}$, controlling for the six other predictor variables ($R_{YX_1|X_{-1}}$). This greatly simplifies matters, especially the elicitation of the conjectured effect.

You use partial regression plots like that shown in [Figure 90.9](#) to teach the team that the partial correlation between PBI and tHcy is the correlation of two sets of residuals obtained from ordinary regression models, one from regressing PBI on the six covariates and the other from regressing tHcy on the same covariates. Thus each subject has “expected” tHcy and PBI values based on the six covariates. The cardiologists believe that subjects whose tHcy is relatively higher than expected will also have a PBI that is relatively higher than expected. The partial correlation quantifies that adjusted association just as a standard simple correlation does with the unadjusted linear association between two variables.

Figure 90.9 Partial Regression Plot



Based on previously published studies of various coronary risk factors and after viewing a set of scatterplots showing various correlations, the team surmises that the true partial correlation is likely to be at least 0.35.

You want to compute the statistical power for a sample size of $N = 100$ by using $\alpha = 0.05$. You also want to plot power for sample sizes between 50 and 150. Use the **MULTREG** statement to compute the power and the **PLOT** statement to produce the graph. Since the predictors are observed rather than fixed in advanced, and a joint multivariate normal assumption seems tenable, use **MODEL=**RANDOM. The following statements perform the power analysis:

```
ods graphics on;

proc power;
  multreg
    model = random
    nfullpredictors = 7
    ntestpredictors = 1
    partialcorr = 0.35
    ntotal = 100
    power = .;
  plot x=n min=50 max=150;
run;

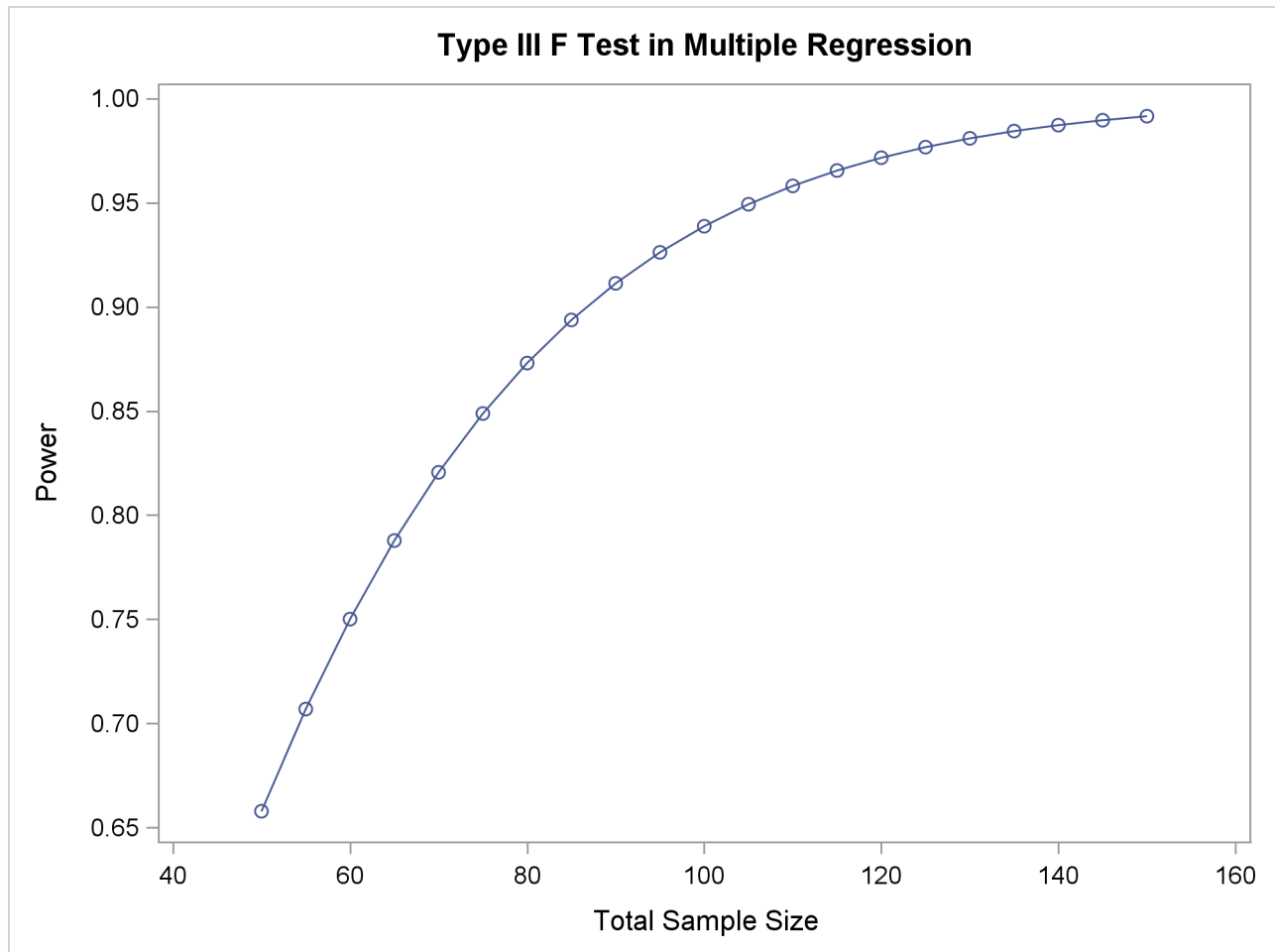
ods graphics off;
```

The **POWER=.** option identifies power as the parameter to compute. The **NFULLPREDICTORS=** option specifies seven total predictors (not including the intercept), and the **NTESTPREDICTORS=** option indicates that one of those predictors is being tested. The **PARTIALCORR=** and **NTOTAL=** options specify the partial correlation and sample size, respectively. The default value for the **ALPHA=** option sets the significance level to 0.05. The **X=N** option in the plot statement requests a plot of sample size on the X axis, and the **MIN=** and **MAX=** options specify the sample size range.

Output 90.5.1 shows the output, and Output 90.5.2 shows the plot.

Output 90.5.1 Power Analysis for Multiple Regression

The POWER Procedure	
Type III F Test in Multiple Regression	
Fixed Scenario Elements	
Method	Exact
Model	Random X
Number of Predictors in Full Model	7
Number of Test Predictors	1
Partial Correlation	0.35
Total Sample Size	100
Alpha	0.05
Computed Power	
Power	
0.939	

Output 90.5.2 Plot of Power versus Sample Size for Multiple Regression

For the sample size $N = 100$, the study is almost balanced with respect to Type I and Type II error rates, with $\alpha = 0.05$ and $\beta = 1 - 0.937 = 0.063$. The study thus seems well designed at this sample size.

Now suppose that in a follow-up meeting with the cardiologists, you discover that their specific intent is to demonstrate that the (partial) correlation between PBI and tHcy is greater than 0.2. You suggest changing the planned data analysis to a one-sided Fisher's z test with a null correlation of 0.2. The following statements perform a power analysis for this test:

```
proc power;
  onecorr dist=fisherz
    npvars = 6
    corr = 0.35
    nullcorr = 0.2
    sides = 1
    ntotal = 100
    power = .;
run;
```

The **DIST=FISHERZ** option in the **ONECORR** statement specifies Fisher's z test. The **NPARTIALVARS=** option specifies that six additional variables are adjusted for in the partial correlation. The **CORR=** option

specifies the conjectured correlation of 0.35, and the **NULLCORR=** option indicates the null value of 0.2. The **SIDES=** option specifies a one-sided test.

Output 90.5.3 shows the output.

Output 90.5.3 Power Analysis for Fisher's z Test

The POWER Procedure
Fisher's z Test for Pearson Correlation

Fixed Scenario Elements	
Distribution	Fisher's z transformation of r
Method	Normal approximation
Number of Sides	1
Null Correlation	0.2
Number of Variables Partialled Out	6
Correlation	0.35
Total Sample Size	100
Nominal Alpha	0.05

Computed Power	
Actual Alpha	Power
0.05	0.466

The power for Fisher's z test is less than 50%, the decrease being mostly due to the smaller effect size (relative to the null value). When asked for a recommendation for a new sample size goal, you compute the required sample size to achieve a power of 0.95 (to balance Type I and Type II errors) and 0.85 (a threshold deemed to be minimally acceptable to the team). The following statements perform the sample size determination:

```
proc power;
  onecorr dist=fisherz
    npvars = 6
    corr = 0.35
    nullcorr = 0.2
    sides = 1
    ntotal = .
    power = 0.85 0.95;
run;
```

The **NTOTAL=** option identifies sample size as the parameter to compute, and the **POWER=** option specifies the target powers.

Output 90.5.4 Sample Size Determination for Fisher's z Test**The POWER Procedure
Fisher's z Test for Pearson Correlation**

Fixed Scenario Elements				
Distribution	Fisher's z transformation of r			
Method	Normal approximation			
Number of Sides				1
Null Correlation				0.2
Number of Variables Partialled Out				6
Correlation				0.35
Nominal Alpha				0.05

Computed N Total				
Index	Nominal Power	Actual Alpha	Actual Power	N Total
1	0.85	0.05	0.850	280
2	0.95	0.05	0.950	417

The results in [Output 90.5.4](#) reveal a required sample size of 417 to achieve a power of 0.95 and a required sample size of 280 to achieve a power of 0.85.

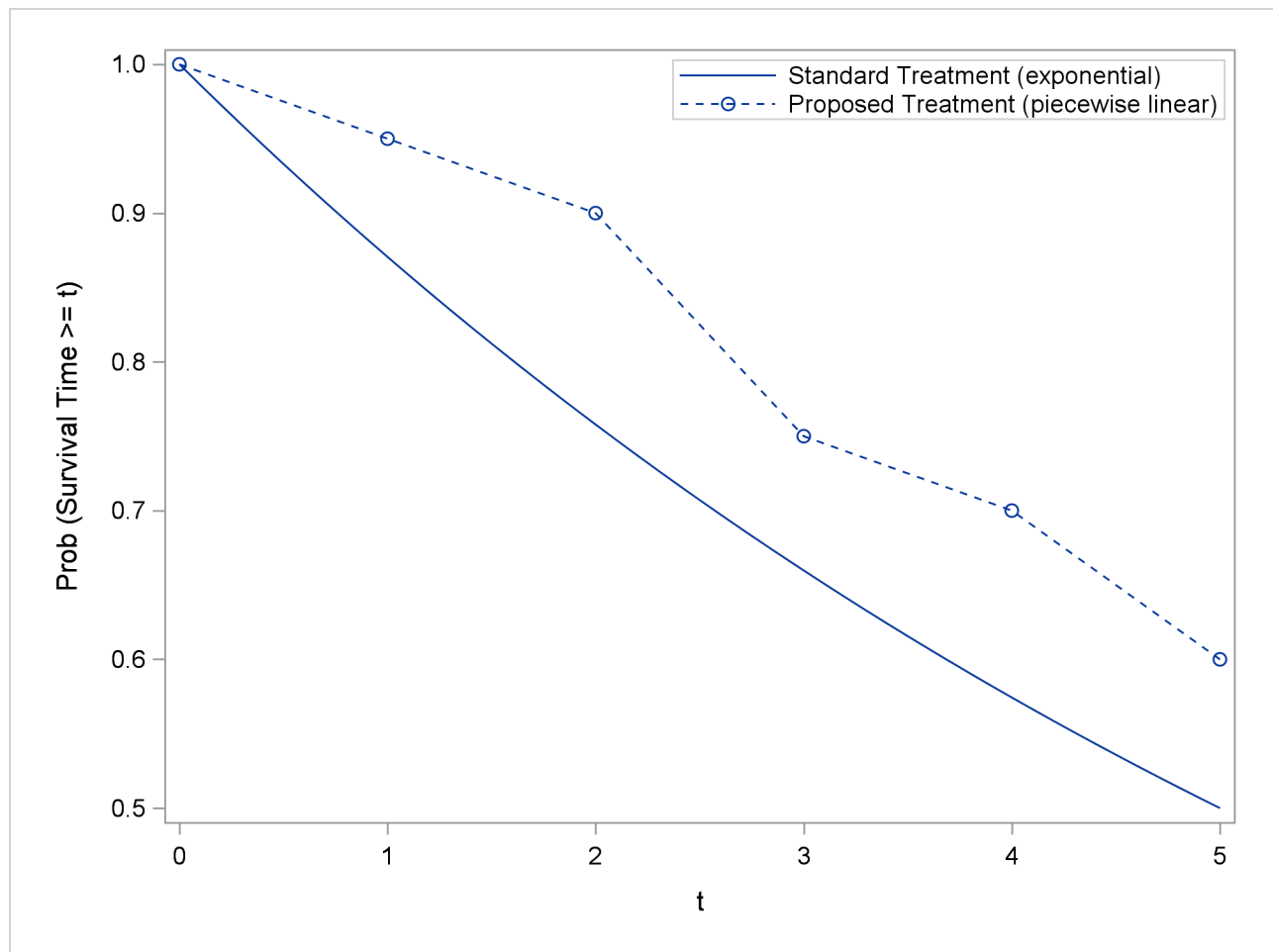
Example 90.6: Comparing Two Survival Curves

You are consulting for a clinical research group planning a trial to compare survival rates for proposed and standard cancer treatments. The planned data analysis is a log-rank test to nonparametrically compare the overall survival curves for the two treatments. Your goal is to determine an appropriate sample size to achieve a power of 0.8 for a two-sided test with $\alpha = 0.05$ by using a balanced design.

The survival curve for patients on the standard treatment is well known to be approximately exponential with a median survival time of five years. The research group conjectures that the new proposed treatment will yield a (nonexponential) survival curve similar to the dashed line in [Figure 90.6.1](#).

Patients will be accrued uniformly over two years and then followed for an additional three years past the accrual period. Some loss to follow-up is expected, with roughly exponential rates that would result in about 50% loss with the standard treatment within 10 years. The loss to follow-up with the proposed treatment is more difficult to predict, but 50% loss would be expected to occur sometime between years 5 and 20.

Output 90.6.1 Survival Curves



Use the **TWOSAMPLESURVIVAL** statement with the **TEST=LOGRANK** option to compute the required sample size for the log-rank test. The following statements perform the analysis:

```
proc power;
  twosamplesurvival test=logrank
    curve("Standard") = 5 : 0.5
    curve("Proposed") = (1 to 5 by 1):(0.95 0.9 0.75 0.7 0.6)
    groupsurvival = "Standard" | "Proposed"
    accrualtime = 2
    followuptime = 3
    groupmedlosstimes = 10 | 20 5
    power = 0.8
    npergroup = .;
run;
```

The **CURVE=** option defines the two survival curves. The “Standard” curve has only one point, specifying an exponential form with a survival probability of 0.5 at year 5. The “Proposed” curve is a piecewise linear curve defined by the five points shown in Figure 90.6.1. The **GROUPSURVIVAL=** option assigns the survival curves to the two groups, and the **ACCRUALTIME=** and **FOLLOWUPTIME=** options specify the accrual and follow-up times. The **GROUPMEDLOSSTIMES=** option specifies the years at which 50% loss is expected

to occur. The **POWER=** option specifies the target power, and the **NPERGROUP=.** option identifies sample size per group as the parameter to compute. Default values for the **SIDES=** and **ALPHA=** options specify a two-sided test with $\alpha = 0.05$.

Output 90.6.2 shows the results.

Output 90.6.2 Sample Size Determination for Log-Rank Test

The POWER Procedure
Log-Rank Test for Two Survival Curves

Fixed Scenario Elements			
Method	Lakatos normal approximation		
Accrual Time	2		
Follow-up Time	3		
Group 1 Survival Curve	Standard		
Form of Survival Curve 1	Exponential		
Group 2 Survival Curve	Proposed		
Form of Survival Curve 2	Piecewise Linear		
Group 1 Median Loss Time	10		
Nominal Power	0.8		
Number of Sides	2		
Number of Time Sub-Intervals	12		
Alpha	0.05		

Computed N per Group			
Index	Median		N per Group
	Loss Time	Actual Power	
1	20	0.800	228
2	5	0.801	234

The required sample size per group to achieve a power of 0.8 is 228 if the median loss time is 20 years for the proposed treatment. Only six more patients are required in each group if the median loss time is as short as five years.

Example 90.7: Confidence Interval Precision

An investment firm has hired you to help plan a study to estimate the success of a new investment strategy called IntuiVest. The study involves complex simulations of market conditions over time, and it tracks the balance of a hypothetical brokerage account starting with \$50,000. Each simulation is very expensive in terms of computing time. You are asked to determine an appropriate number of simulations to estimate the average change in the account balance at the end of three years. The goal is to have a 95% chance of obtaining a 90% confidence interval whose half-width is at most \$1,000. That is, the firm wants to have a 95% chance of being able to correctly claim at the end of the study that “Our research shows with 90% confidence that IntuiVest yields a profit of \$X +/- \$1,000 at the end of three years on an initial investment of \$50,000 (under simulated market conditions).”

The probability of achieving the desired precision (that is, a small interval width) can be calculated either unconditionally or conditionally given that the true mean is captured by the interval. You decide to use the conditional form, considering two of its advantages:

- The conditional probability is usually lower than the unconditional probability for the same sample size, meaning that the conditional form is generally conservative.
- The overall probability of achieving the desired precision *and* capturing the true mean is easily computed as the product of the half-width probability and the confidence level. In this case, the overall probability is $0.95 \times 0.9 = 0.855$.

Based on some initial simulations, you expect a standard deviation between \$25,000 and \$45,000 for the ending account balance. You will consider both of these values in the sample size analysis.

As mentioned in the section “[Overview of Power Concepts](#)” on page 7326, an analysis of confidence interval precision is analogous to a traditional power analysis, with “CI Half-Width” taking the place of effect size and “Prob(Width)” taking the place of power. In this example, the target CI Half-Width is 1000, and the desired Prob(Width) is 0.95.

In addition to computing sample sizes for a half-width of \$1,000, you are asked to plot the required number of simulations for a range of half-widths between \$500 and \$2,000. Use the [ONESAMPLEMEANS](#) statement with the [CI=T](#) option to implement the sample size determination. The following statements perform the analysis:

```
ods graphics on;

proc power;
  onesamplemeans ci=t
    alpha = 0.1
    halfwidth = 1000
    stddev = 25000 45000
    probwidth = 0.95
    ntotal = .;
  plot x=effect min=500 max=2000;
run;

ods graphics off;
```

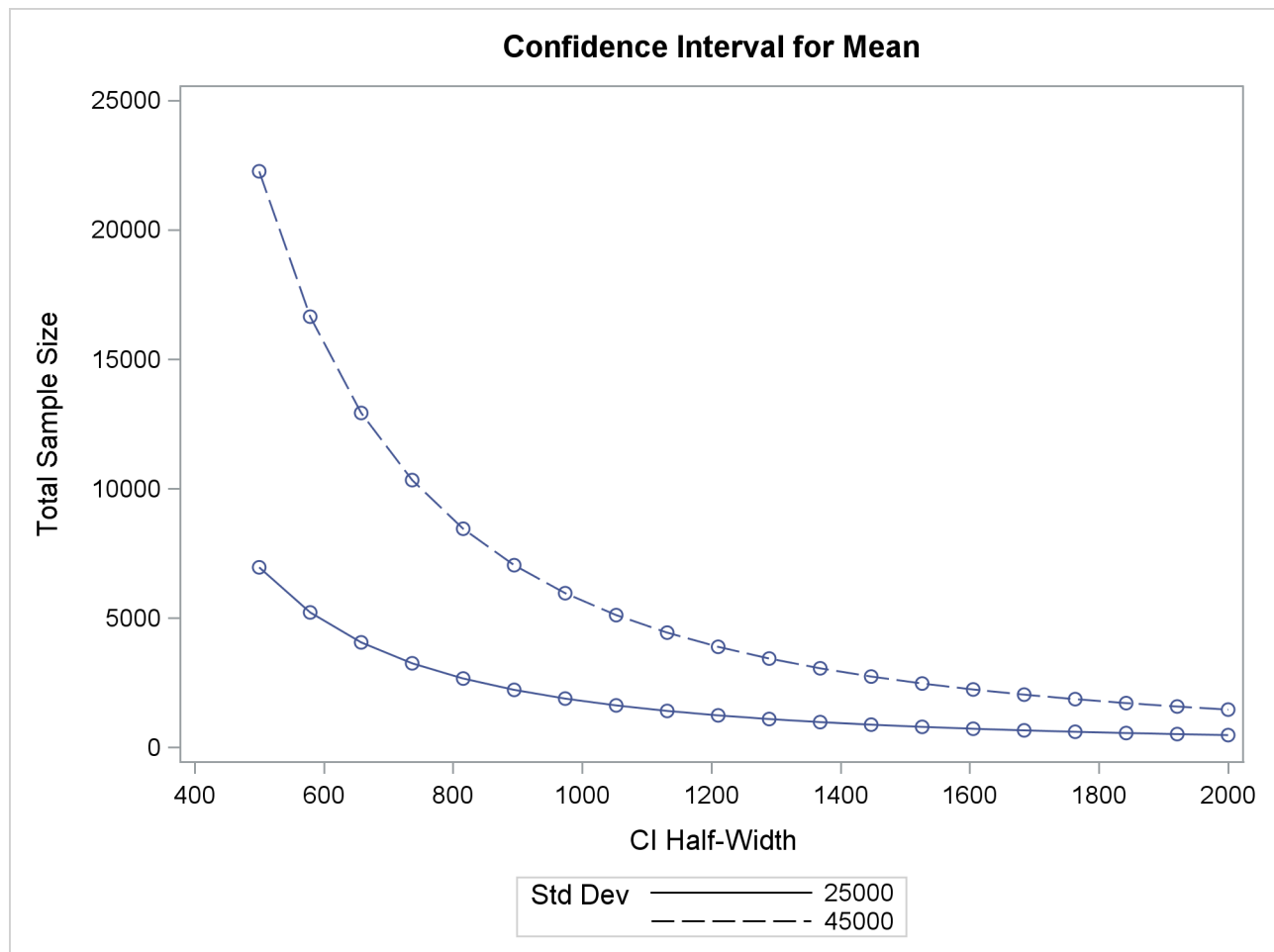
The **NTOTAL=** option identifies sample size as the parameter to compute. The **ALPHA=0.1** option specifies a confidence level of $1 - \alpha = 0.9$. The **HALFWIDTH=** option specifies the target half-width, and the **STDDEV=** option specifies the conjectured standard deviation values. The **PROBWIDTH=** option specifies the desired probability of achieving the target precision. The default value **PROBTYPE=CONDITIONAL** specifies that this probability is conditional on the true mean being captured by the interval. The default of **SIDES=2** indicates a two-sided interval.

Output 90.7.1 shows the output, and Output 90.7.2 shows the plot.

Output 90.7.1 Sample Size Determination for Confidence Interval Precision

The POWER Procedure
Confidence Interval for Mean

Fixed Scenario Elements			
Distribution		Normal	
Method		Exact	
Alpha		0.1	
CI Half-Width		1000	
Nominal Prob(Width)		0.95	
Number of Sides		2	
Prob Type		Conditional	
Computed N Total			
Index	Std Dev	Actual Prob(Width)	N Total
1	25000	0.951	1788
2	45000	0.950	5652

Output 90.7.2 Plot of Sample Size versus Confidence Interval Half-Width

The number of simulations required in order to have a 95% chance of obtaining a half-width of at most 1000 is between 1788 and 5652, depending on the standard deviation. The plot reveals that more than 20,000 simulations would be required for a half-width of 500, assuming the higher standard deviation.

Example 90.8: Customizing Plots

This example demonstrates various ways you can modify and enhance plots:

- assigning analysis parameters to axes
- fine-tuning a sample size axis
- adding reference lines
- linking plot features to analysis parameters
- choosing key (legend) styles
- modifying symbol locations

The example plots are all based on a sample size analysis for a two-sample t test of group mean difference. You start by computing the sample size required to achieve a power of 0.9 by using a two-sided test with $\alpha = 0.05$, assuming the first mean is 12, the second mean is either 15 or 18, and the standard deviation is either 7 or 9.

Use the **TWOSAMPLEMEANS** statement with the **TEST=DIFF** option to compute the required sample sizes. Indicate total sample size as the result parameter by supplying a missing value (.) with the **NTOTAL=** option. Use the **GROUPMEANS=**, **STDDEV=**, and **POWER=** options to specify values of the other parameters. The following statements perform the sample size computations:

```
proc power;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = 0.9
    ntotal        = .;
run;
```

Default values for the **NULLDIFF=**, **SIDES=**, **GROUPWEIGHTS=**, and **DIST=** options specify a null mean difference of 0, two-sided test, balanced design, and assumption of normally distributed data, respectively.

Output 90.8.1 shows that the required sample size ranges from 60 to 382, depending on the unknown standard deviation and second mean.

Output 90.8.1 Computed Sample Sizes

The POWER Procedure Two-Sample t Test for Mean Difference

Fixed Scenario Elements				
Distribution	Normal			
Method	Exact			
Group 1 Mean	12			
Nominal Power	0.9			
Number of Sides	2			
Null Difference	0			
Alpha	0.05			
Group 1 Weight	1			
Group 2 Weight	1			

Computed N Total				
Index	Mean2	Std Dev	Actual Power	N Total
1	15	7	0.902	232
2	15	9	0.901	382
3	18	7	0.904	60
4	18	9	0.904	98

Assigning Analysis Parameters to Axes

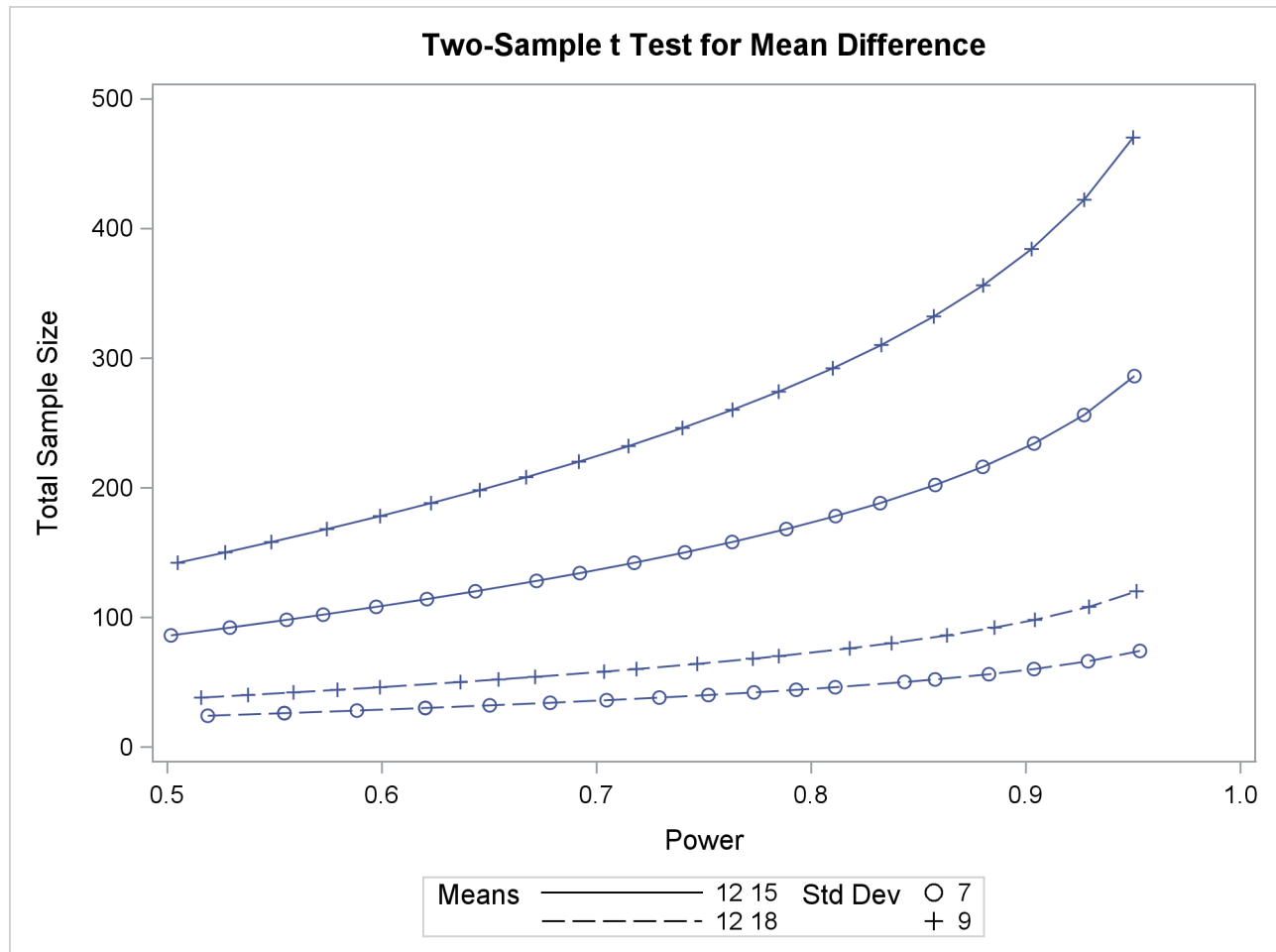
Use the **PLOT** statement to produce plots for all power and sample size analyses in PROC POWER. For the sample size analysis described at the beginning of this example, suppose you want to plot the required sample size on the Y axis against a range of powers between 0.5 and 0.95 on the X axis. The **X=** and **Y=** options specify which parameter to plot against the result and which axis to assign to this parameter. You can use either the **X=** or the **Y=** option, but not both. Use the **X=POWER** option in the **PLOT** statement to request a plot with power on the X axis. The result parameter, here total sample size, is always plotted on the other axis. Use the **MIN=** and **MAX=** options to specify the range of the axis indicated with either the **X=** or the **Y=** option. Here, specify **MIN=0.5** and **MAX=0.95** to specify the power range. The following statements produce the plot:

```
ods graphics on;

proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = 0.9
    ntotal        = .;
  plot x=power min=0.5 max=0.95;
run;
```

Note that the value (0.9) of the **POWER=** option in the **TWOSAMPLEMEANS** statement is only a placeholder when the **PLOTONLY** option is used and both the **MIN=** and **MAX=** options are used, because the values of the **MIN=** and **MAX=** options override the value of 0.9. But the **POWER=** option itself is still required in the **TWOSAMPLEMEANS** statement, to provide a complete specification of the sample size analysis.

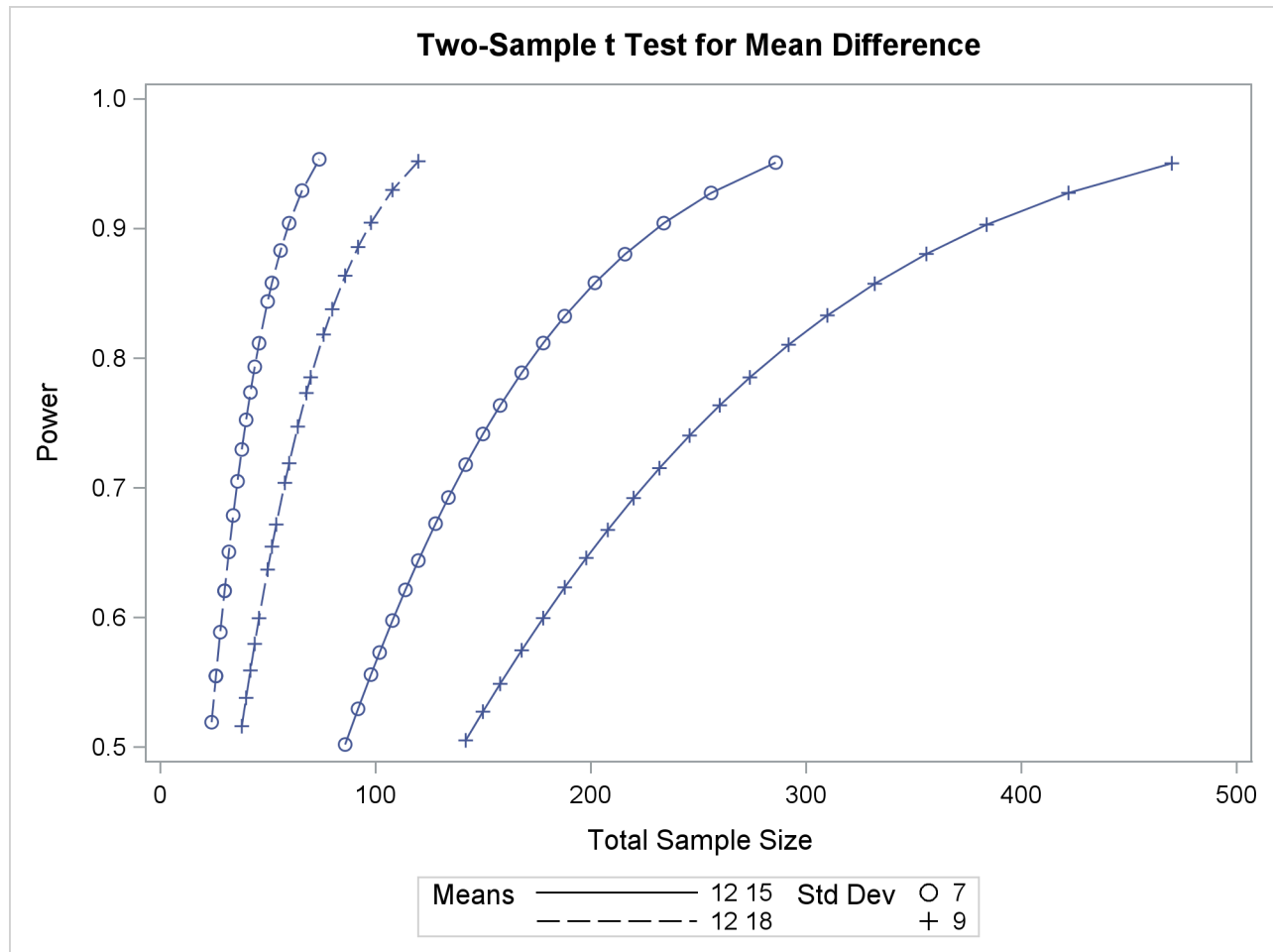
The resulting plot is shown in [Output 90.8.2](#).

Output 90.8.2 Plot of Sample Size versus Power

The line style identifies the group means scenario, and the plotting symbol identifies the standard deviation scenario. The locations of plotting symbols indicate computed sample sizes; the curves are linear interpolations of these points. By default, each curve consists of approximately 20 computed points (sometimes slightly more or less, depending on the analysis).

If you would rather plot power on the Y axis versus sample size on the X axis, you have two general strategies to choose from. One strategy is to use the **Y=** option instead of the **X=** option in the **PLOT** statement:

```
plot y=power min=0.5 max=0.95;
```

Output 90.8.3 Plot of Power versus Sample Size using First Strategy

Note that the resulting plot ([Output 90.8.3](#)) is essentially a mirror image of [Output 90.8.2](#). The axis ranges are set such that each curve in [Output 90.8.3](#) contains similar values of Y instead of X. Each plotted point represents the computed value of the X axis at the input value of the Y axis.

A second strategy for plotting power versus sample size (when originally solving for sample size) is to invert the analysis and base the plot on computed power for a given range of sample sizes. This strategy works well for monotonic power curves (as is the case for the t test and most other continuous analyses). It is advantageous in the sense of preserving the traditional role of the Y axis as the computed parameter. A common way to implement this strategy is as follows:

- Determine the range of sample sizes sufficient to cover at the desired power range for all curves (where each “curve” represents a scenario for standard deviation and second group mean).
- Use this range for the X axis of a plot.

To determine the required sample sizes for target powers of 0.5 and 0.95, change the values in the **POWER=** option as follows to reflect this range:

```
proc power;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = 0.5 0.95
    ntotal        = .;
run;
```

Output 90.8.4 reveals that a sample size range of 24 to 470 is approximately sufficient to cover the desired power range of 0.5 to 0.95 for all curves (“approximately” because the actual power at the rounded sample size of 24 is slightly higher than the nominal power of 0.5).

Output 90.8.4 Computed Sample Sizes

The POWER Procedure Two-Sample t Test for Mean Difference

Fixed Scenario Elements					
Distribution	Normal				
Method	Exact				
Group 1 Mean	12				
Number of Sides	2				
Null Difference	0				
Alpha	0.05				
Group 1 Weight	1				
Group 2 Weight	1				

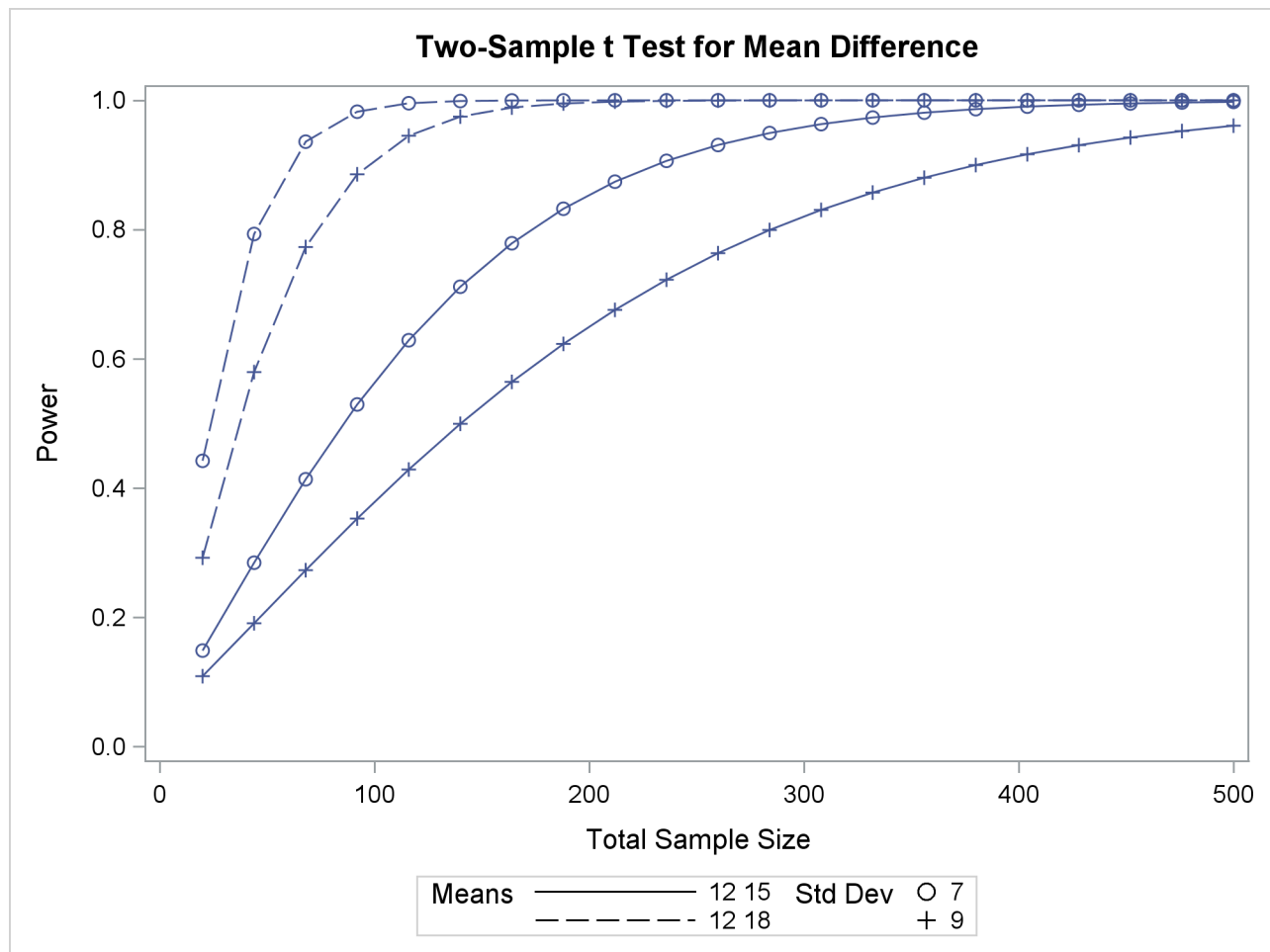
Computed N Total					
Index	Mean2	Std Dev	Nominal Power	Actual Power	N Total
1	15	7	0.50	0.502	86
2	15	7	0.95	0.951	286
3	15	9	0.50	0.505	142
4	15	9	0.95	0.950	470
5	18	7	0.50	0.519	24
6	18	7	0.95	0.953	74
7	18	9	0.50	0.516	38
8	18	9	0.95	0.952	120

To plot power on the Y axis for sample sizes between 20 and 500, use the **X=N** option in the **PLOT** statement with **MIN=20** and **MAX=500**:

```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = .
    ntotal        = 200;
  plot x=n min=20 max=500;
run;
```

Each curve in the resulting plot in [Output 90.8.5](#) covers at least a power range of 0.5 to 0.95.

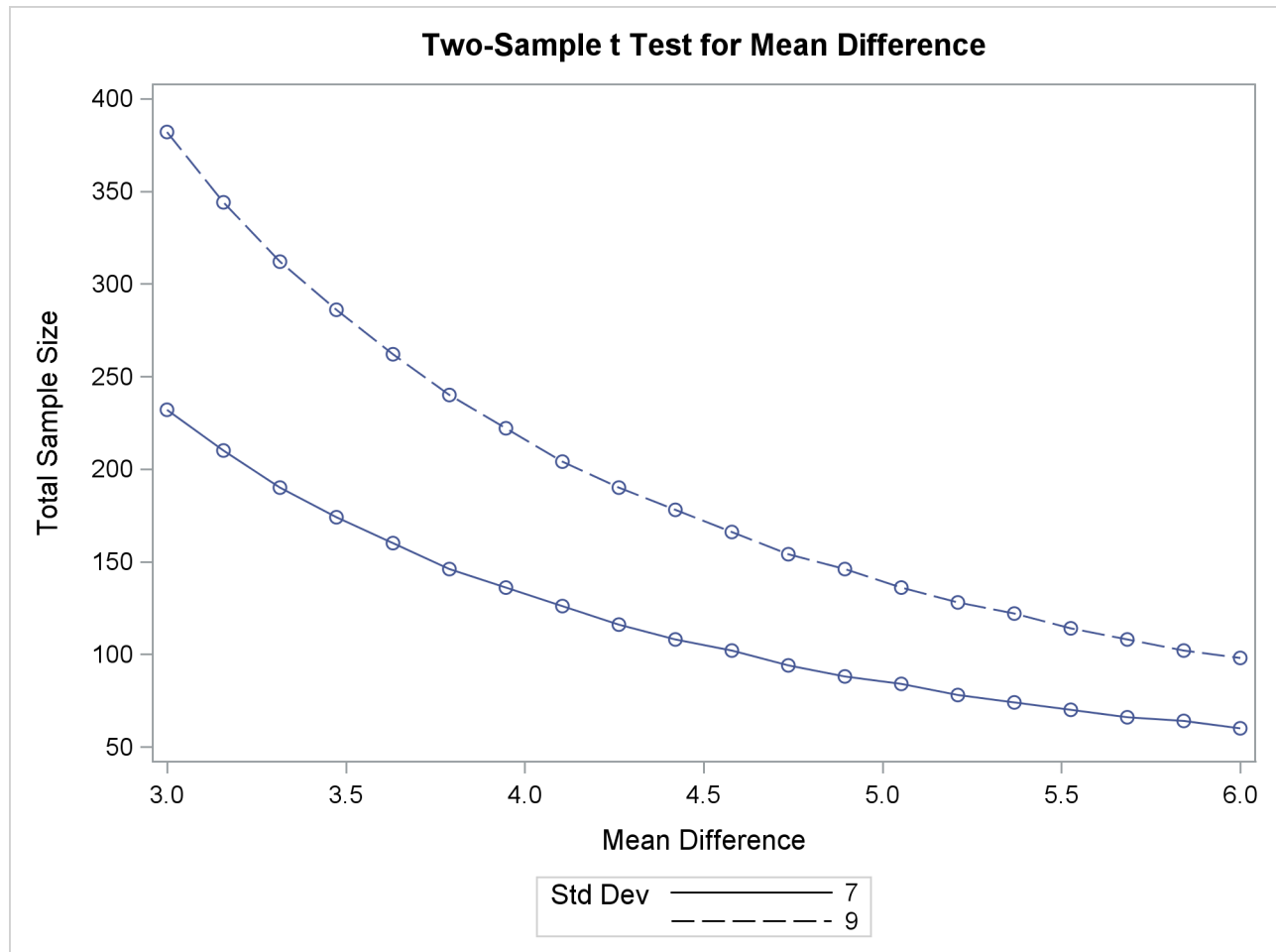
Output 90.8.5 Plot of Power versus Sample Size Using Second Strategy



Finally, suppose you want to produce a plot of sample size versus effect size for a power of 0.9. In this case, the “effect size” is defined to be the mean difference. You need to reparameterize the analysis by using the **MEANDIFF=** option instead of the **GROUPMEANS=** option to produce a plot, since each plot axis must be represented by a scalar parameter. Use the **X=EFFECT** option in the **PLOT** statement to assign the mean difference to the X axis. The following statements produce a plot of required sample size to detect mean differences between 3 and 6:

```
proc power plotonly;
  twosamplemeans test=diff
    meandiff      = 3 6
    stddev        = 7 9
    power         = 0.9
    ntotal        = .;
  plot x=effect min=3 max=6;
run;
```

The resulting plot [Output 90.8.6](#) shows how the required sample size decreases with increasing mean difference.

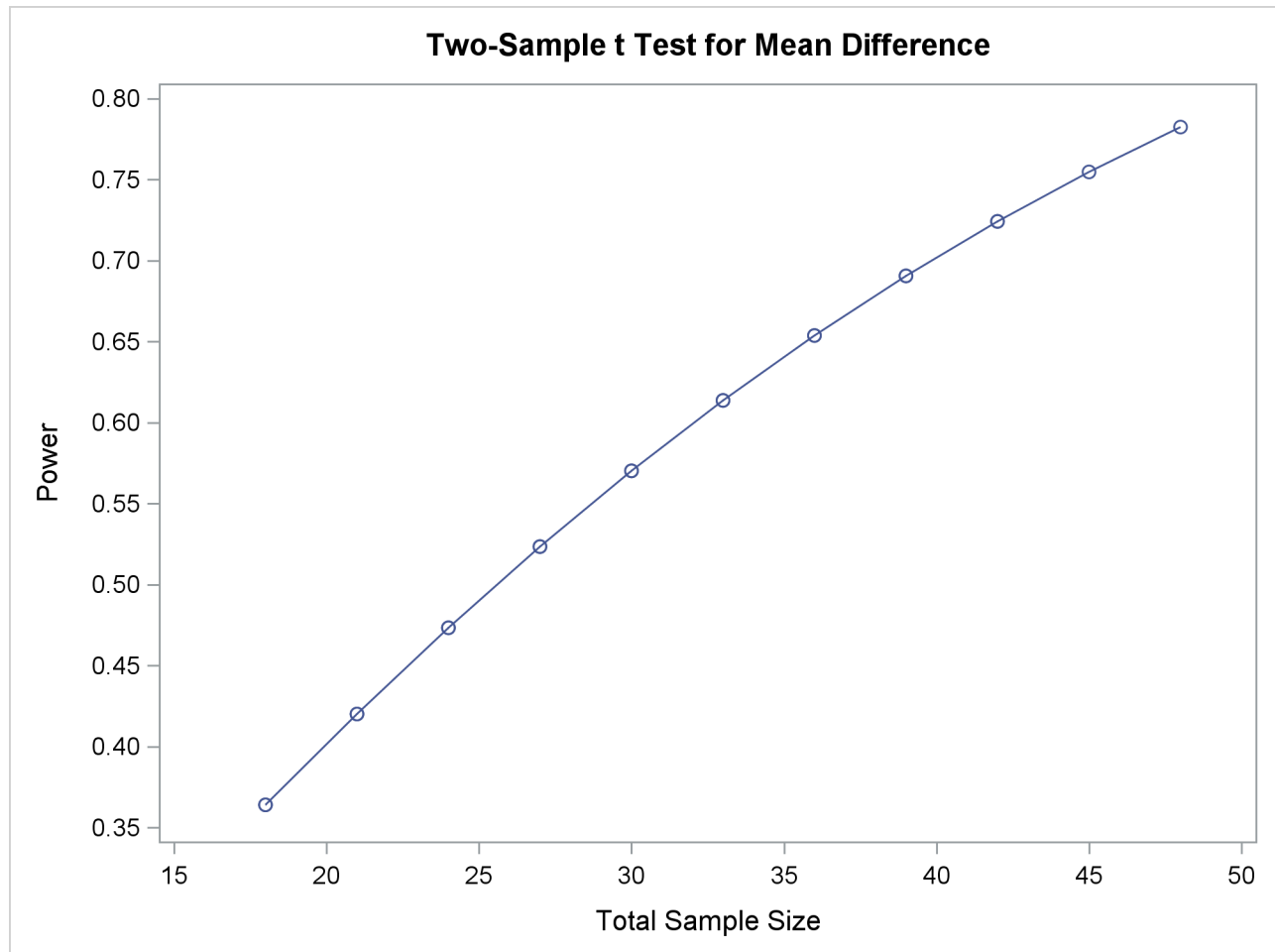
Output 90.8.6 Plot of Sample Size versus Mean Difference

Fine-Tuning a Sample Size Axis

Consider the following plot request for a sample size analysis similar to the one in [Output 90.8.1](#) but with only a single scenario, and with unbalanced sample size allocation of 2:1:

```
proc power plotonly;
  ods output plotcontent=PlotData;
  twosamplemeans test=diff
    groupmeans    = 12 | 18
    stddev        = 7
    groupweights  = 2 | 1
    power         = .
    ntotal        = 20;
  plot x=n min=20 max=50 npoints=20;
run;
```

The `MIN=`, `MAX=`, and `NPOINTS=` options in the `PLOT` statement request a plot with 20 points between 20 and 50. But the resulting plot ([Output 90.8.7](#)) appears to have only 11 points, and they range from 18 to 48.

Output 90.8.7 Plot with Overlapping Points

The reason that this plot has fewer points than usual is due to the rounding of sample sizes. If you do not use the **NFRACTIONAL** option in the analysis statement (here, the **TWOSAMPLEMEANS** statement), then the set of sample size points determined by the **MIN=**, **MAX=**, **NPOINTS=**, and **STEP=** options in the **PLOT** statement can be rounded to satisfy the allocation weights. In this case, they are rounded down to the nearest multiples of 3 (the sum of the weights), and many of the points overlap. To see the overlap, you can print the **NominalNTotal** (unadjusted) and **NTotal** (rounded) variables in the **PlotContent** ODS object (here saved to a data set called **PlotData**):

```
proc print data=PlotData;
  var NominalNTotal NTotal;
run;
```

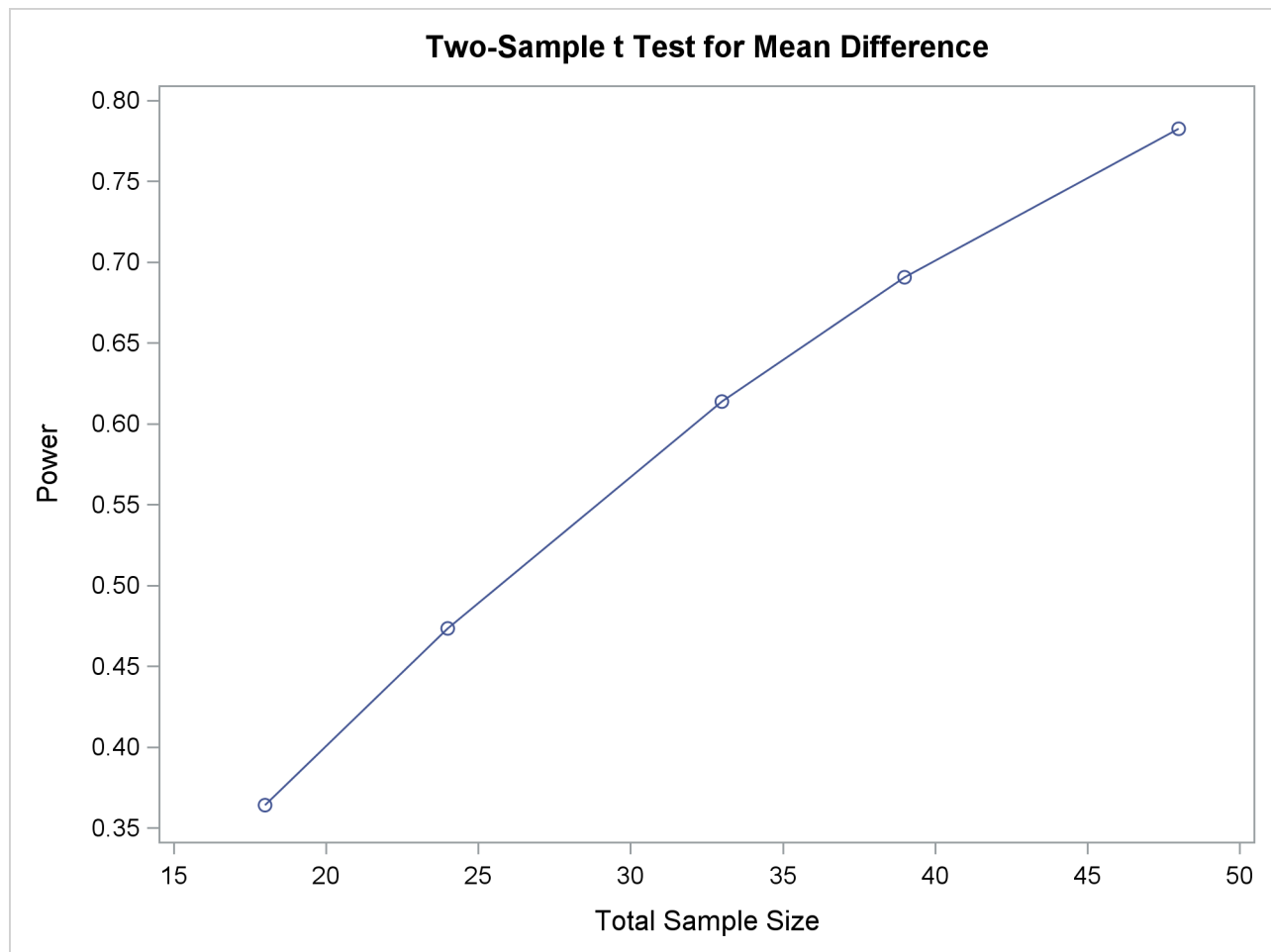
The output is shown in [Output 90.8.8](#).

Output 90.8.8 Sample Sizes

Obs	Nominal	NTotal	NTotal
1		18.0	18
2		19.6	18
3		21.2	21
4		22.7	21
5		24.3	24
6		25.9	24
7		27.5	27
8		29.1	27
9		30.6	30
10		32.2	30
11		33.8	33
12		35.4	33
13		36.9	36
14		38.5	36
15		40.1	39
16		41.7	39
17		43.3	42
18		44.8	42
19		46.4	45
20		48.0	48

Besides overlapping of sample size points, another peculiarity that might occur without the **NFRACTIONAL** option is unequal spacing—for example, in the plot in [Output 90.8.9](#), created with the following statements:

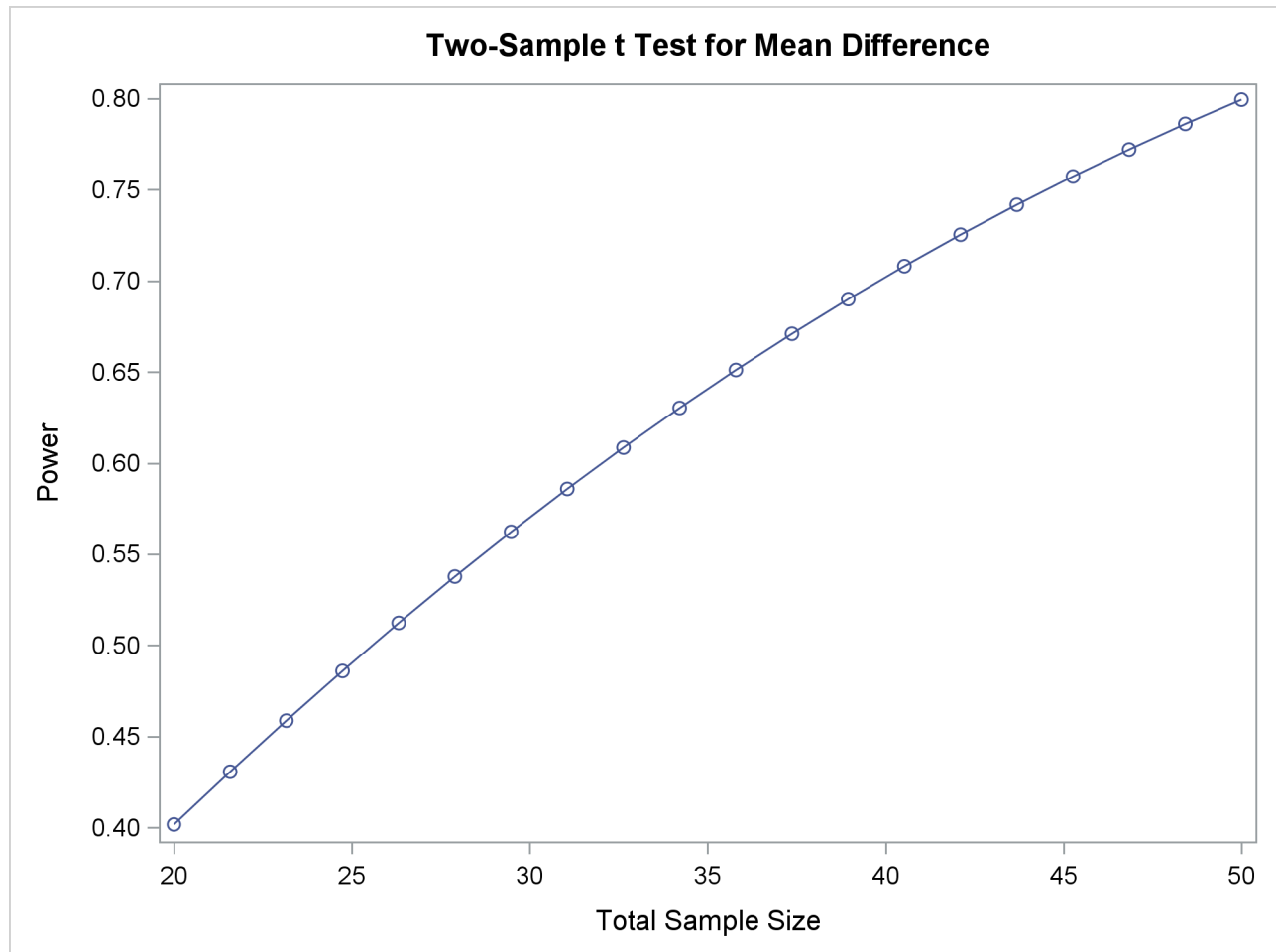
```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 18
    stddev        = 7
    groupweights  = 2 | 1
    power         = .
    ntotal        = 20;
  plot x=n min=20 max=50 npoints=5;
run;
```

Output 90.8.9 Plot with Unequally Spaced Points

If you want to guarantee evenly spaced, nonoverlapping sample size points in your plots, you can either (1) use the **NFRACTIONAL** option in the analysis statement preceding the **PLOT** statement or (2) use the **STEP=** option and provide values for the **MIN=**, **MAX=**, and **STEP=** options in the **PLOT** statement that are multiples of the sum of the allocation weights. Note that this sum is simply 1 for one-sample and paired designs and 2 for balanced two-sample designs. So any integer step value works well for one-sample and paired designs, and any even step value works well for balanced two-sample designs. Both of these strategies will avoid rounding adjustments.

The following statements implement the first strategy to create the plot in [Output 90.8.10](#), by using the **NFRACTIONAL** option in the **TWOSAMPLEMEANS** statement:

```
proc power plotonly;
  twosamplemeans test=diff
    nfractional
    groupmeans    = 12 | 18
    stddev        = 7
    groupweights  = 2 | 1
    power         = .
    ntotal        = 20;
  plot x=n min=20 max=50 npoints=20;
run;
```

Output 90.8.10 Plot with Fractional Sample Sizes

To implement the second strategy, use multiples of 3 for the **STEP=**, **MIN=**, and **MAX=** options in the **PLOT** statement (because the sum of the allocation weights is $2 + 1 = 3$). The following statements use **STEP=3**, **MIN=18**, and **MAX=48** to create a plot that looks identical to the plot in [Output 90.8.7](#) but suffers no overlapping of points:

```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 18
    stddev        = 7
    groupweights  = 2 | 1
    power         = .
    ntotal        = 20;
  plot x=n min=18 max=48 step=3;
run;
```

Adding Reference Lines

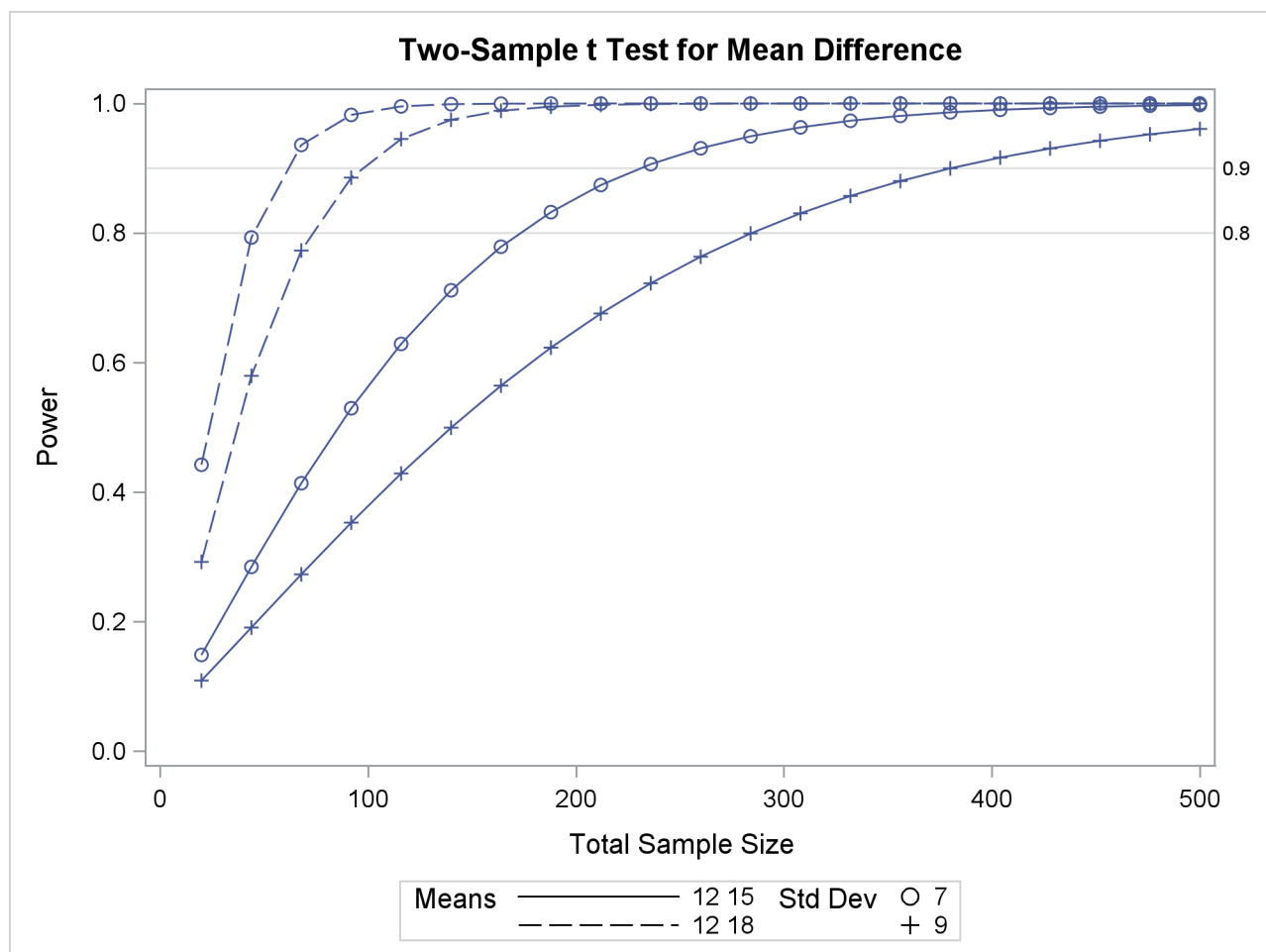
Suppose you want to add reference lines to highlight $\text{power}=0.8$ and $\text{power}=0.9$ on the plot in [Output 90.8.5](#). You can add simple reference lines by using the **YOPTS=** option and **REF=** suboption in the **PLOT** statement to produce [Output 90.8.11](#), with the following statements:

```

proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = .
    ntotal        = 100;
  plot x=n min=20 max=500
    yopts=(ref=0.8 0.9);
run;

```

Output 90.8.11 Plot with Simple Reference Lines on Y Axis



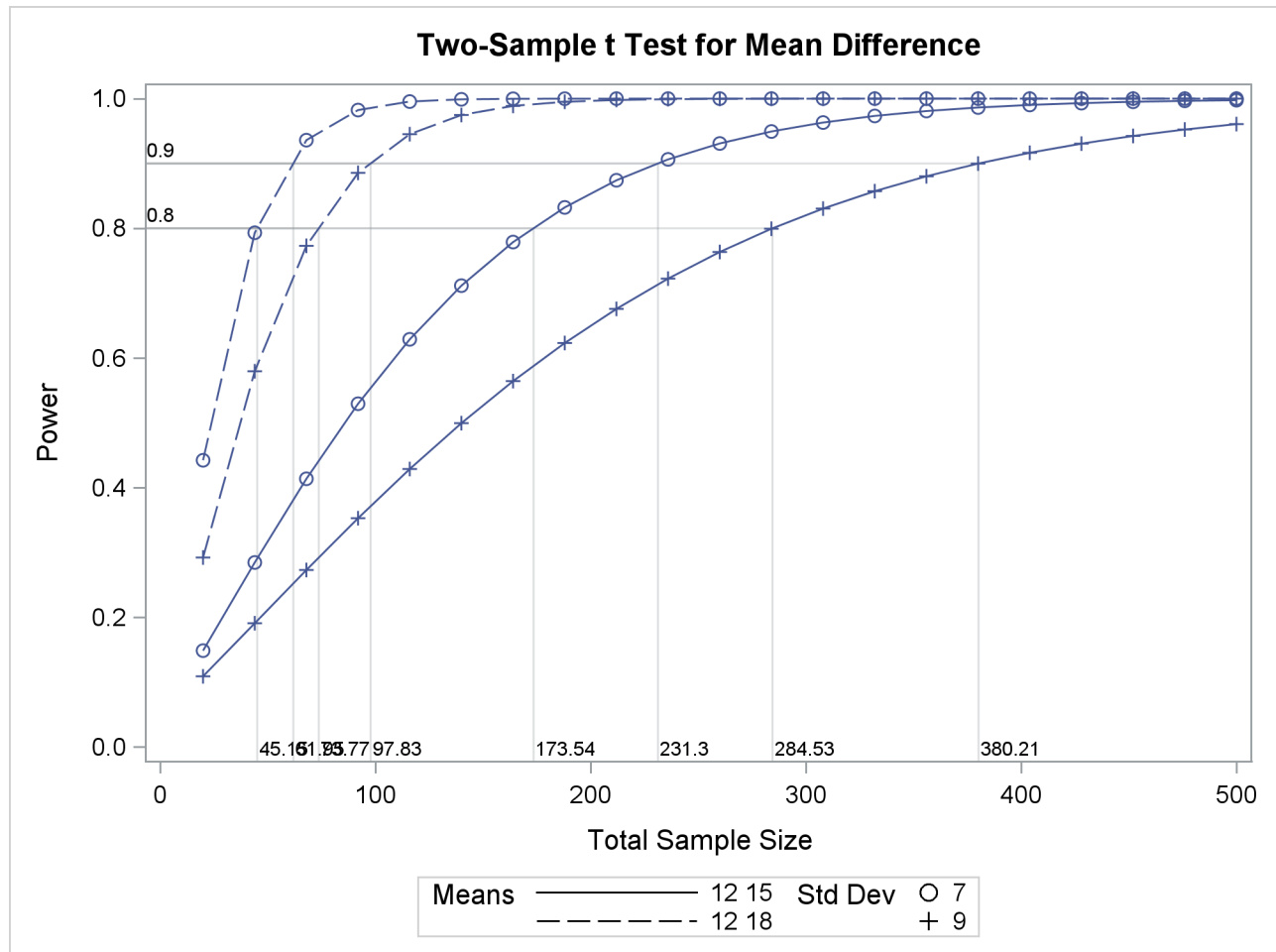
Or you can specify **CROSSREF=YES** to add reference lines that intersect each curve and cross over to the other axis:

```

plot x=n min=20 max=500
  yopts=(ref=0.8 0.9 crossref=yes);

```

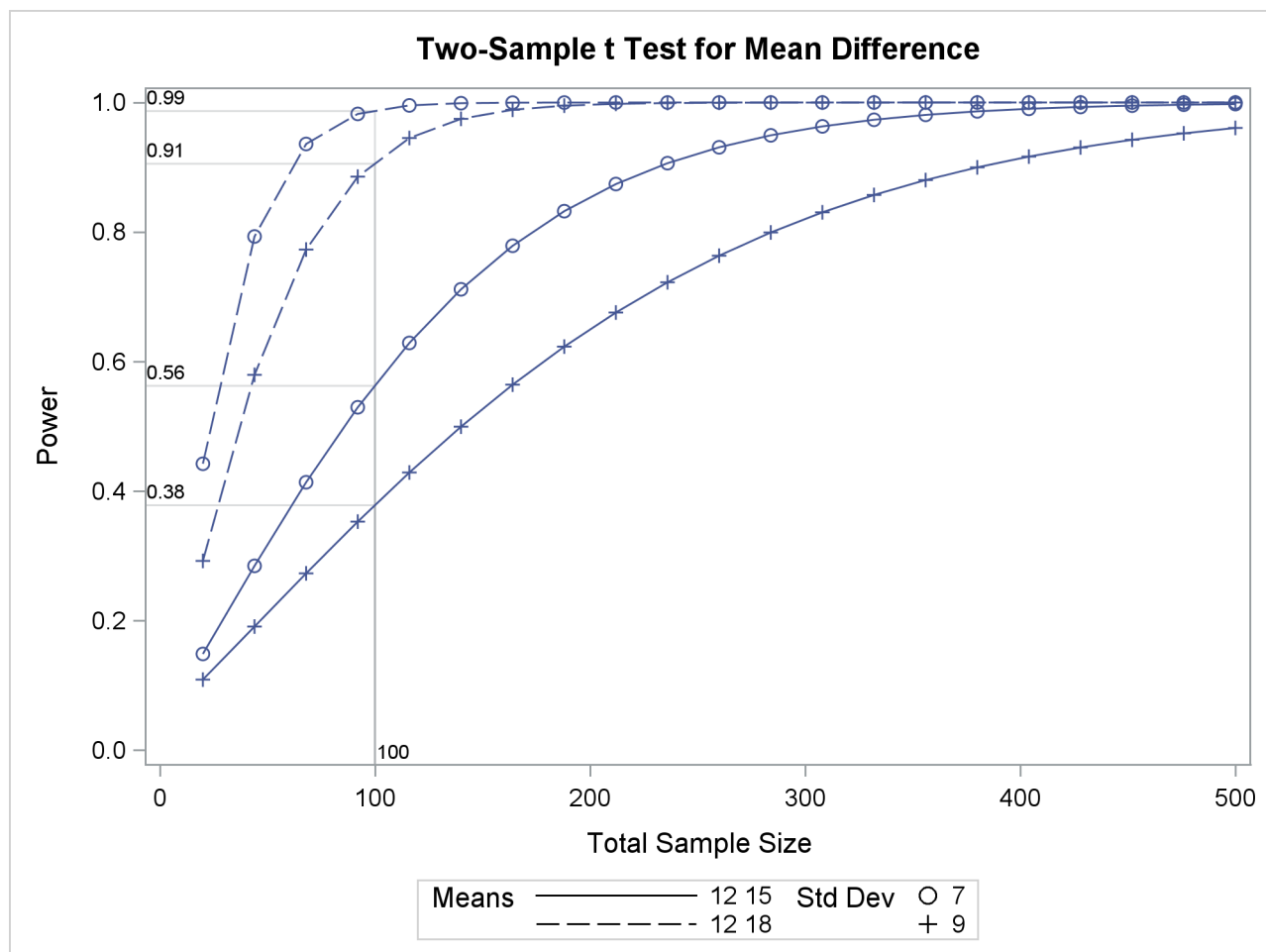
The resulting plot is shown in [Output 90.8.12](#).

Output 90.8.12 Plot with CROSSREF=YES Style Reference Lines from Y Axis

You can also add reference lines for the X axis by using the **XOPTS=** option instead of the **YOPTS=** option. For example, the following **PLOT** statement produces [Output 90.8.13](#), which has crossing reference lines highlighting the sample size of 100:

```
plot x=n min=20 max=500
     xopts=(ref=100 crossref=yes);
```

Note that the values that label the reference lines at the X axis in [Output 90.8.12](#) and at the Y axis in [Output 90.8.13](#) are linearly interpolated from two neighboring points on the curves. Thus they might not exactly match corresponding values that are computed directly from the methods in the section “[Computational Methods and Formulas](#)” on page 7335—that is, computed by PROC POWER in the absence of a PLOT statement. The two ways of computing these values generally differ by a negligible amount.

Output 90.8.13 Plot with CROSSREF=YES Style Reference Lines from X Axis

Linking Plot Features to Analysis Parameters

You can use the **VARY** option in the **PLOT** statement to specify which of the following features you want to associate with analysis parameters.

- line style
- plotting symbol
- color
- panel

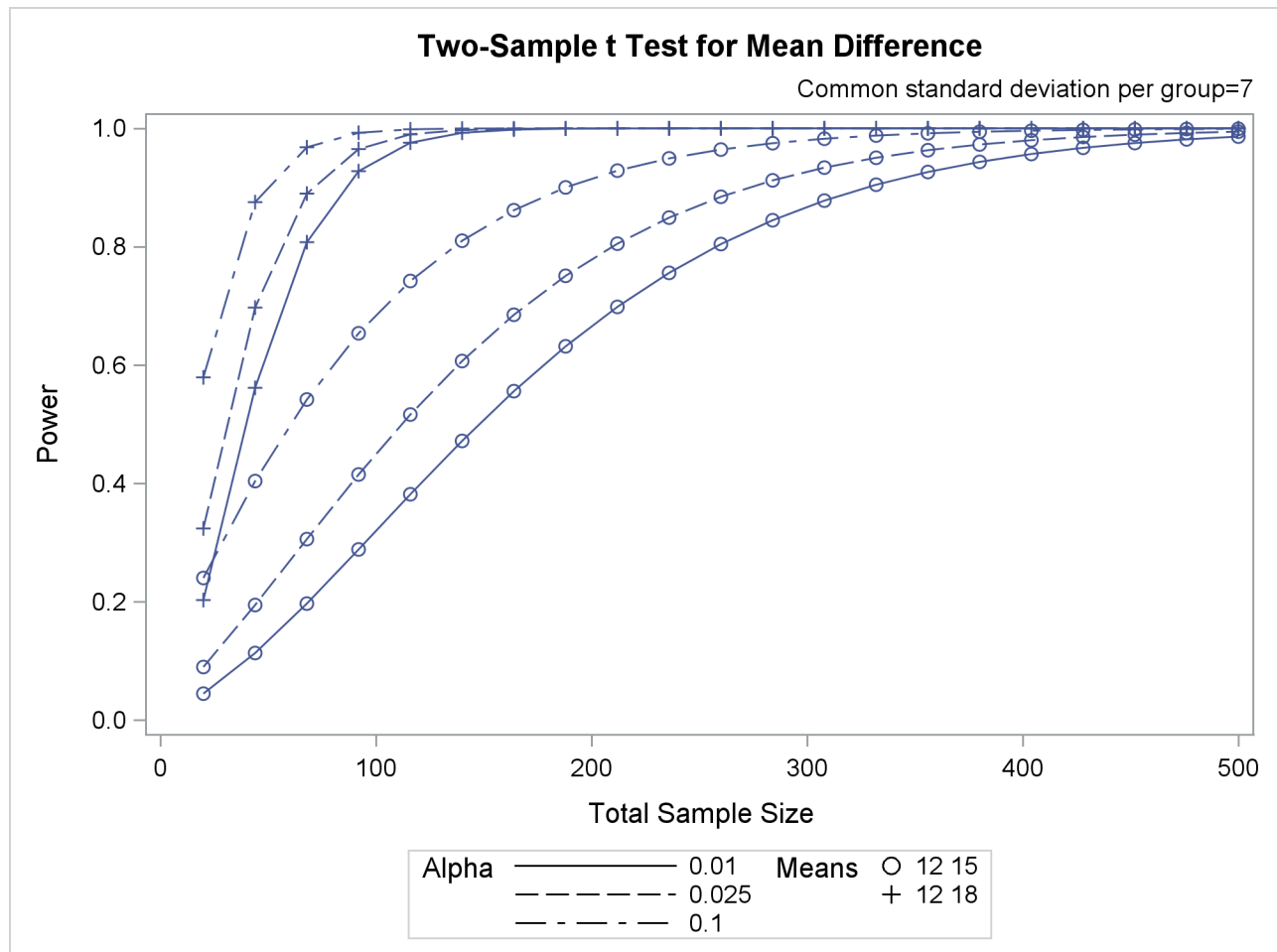
You can specify mappings between each of these features and one or more analysis parameters, or you can simply choose a subset of these features to use (and rely on default settings to associate these features with multiple-valued analysis parameters).

Suppose you supplement the sample size analysis in [Output 90.8.5](#) to include three values of alpha, by using the following statements:

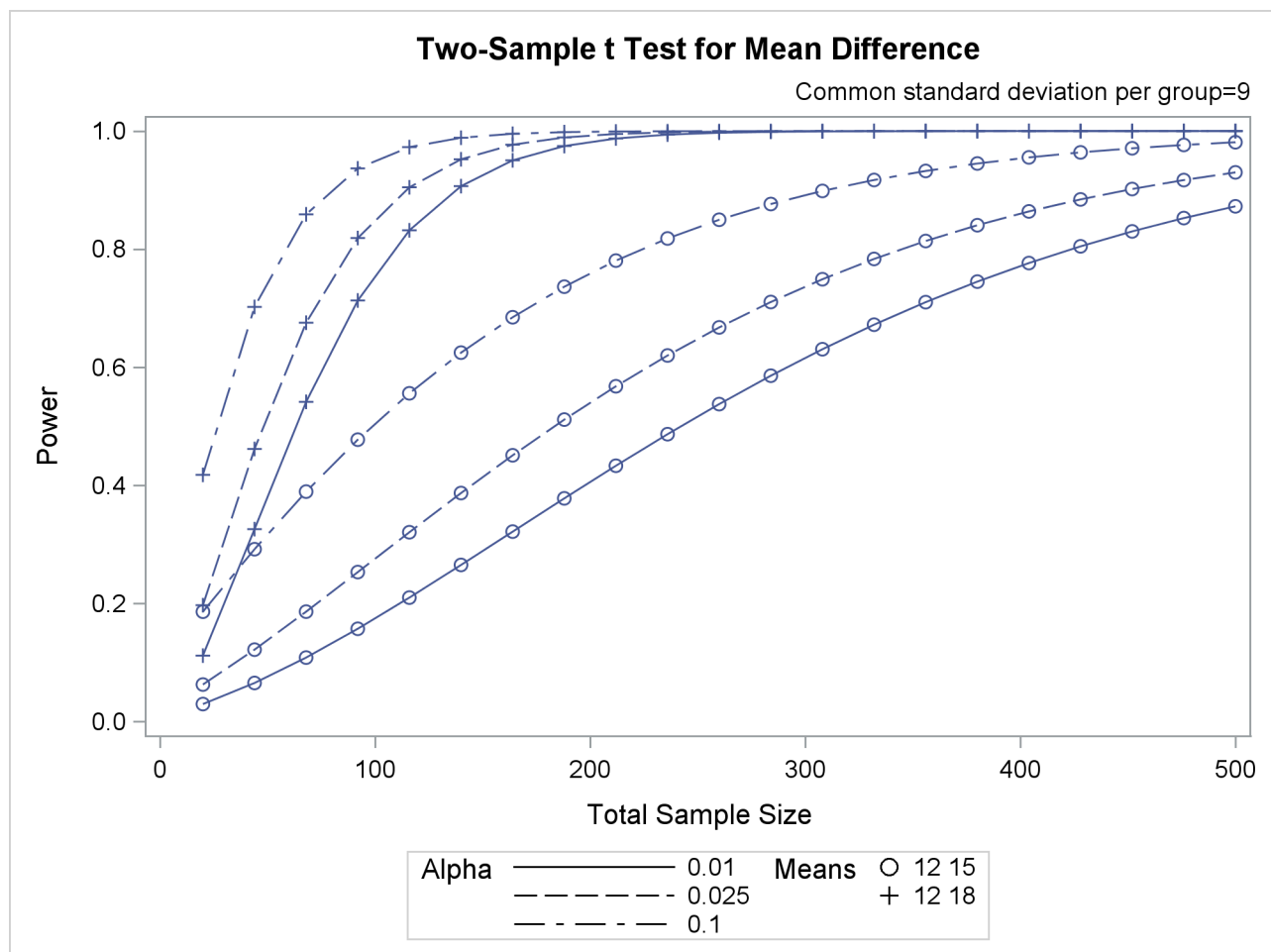
```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    alpha         = 0.01 0.025 0.1
    power         = .
    ntotal        = 100;
  plot x=n min=20 max=500;
run;
```

The defaults for the **VARY** option in the **PLOT** statement specify line style varying by the **ALPHA=** parameter, plotting symbol varying by the **GROUPMEANS=** parameter, panel varying by the **STDDEV=** parameter, and color remaining constant. The resulting plot, consisting of two panels, is shown in [Output 90.8.14](#).

Output 90.8.14 Plot with Default VARY Settings

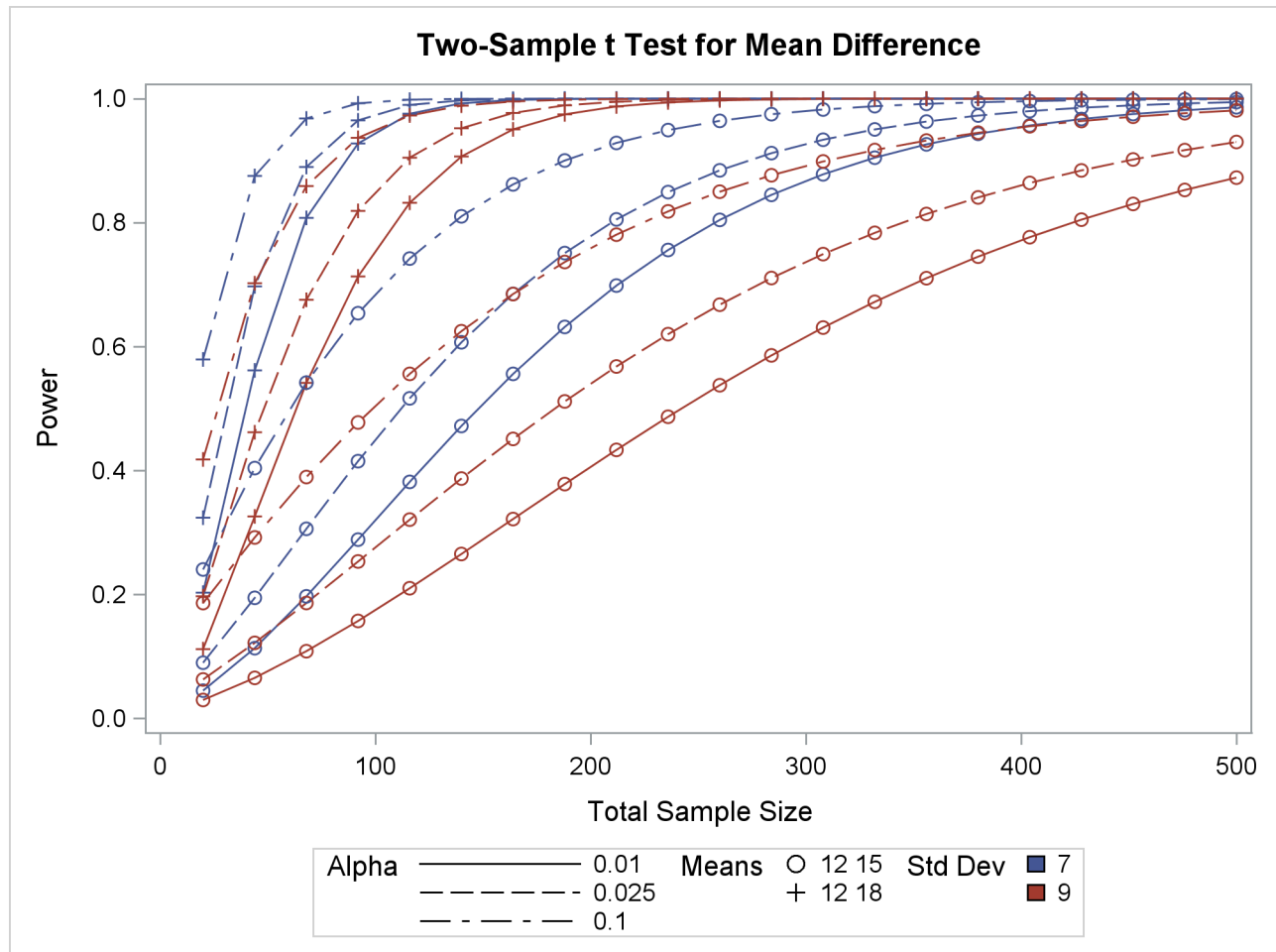


Output 90.8.14 continued



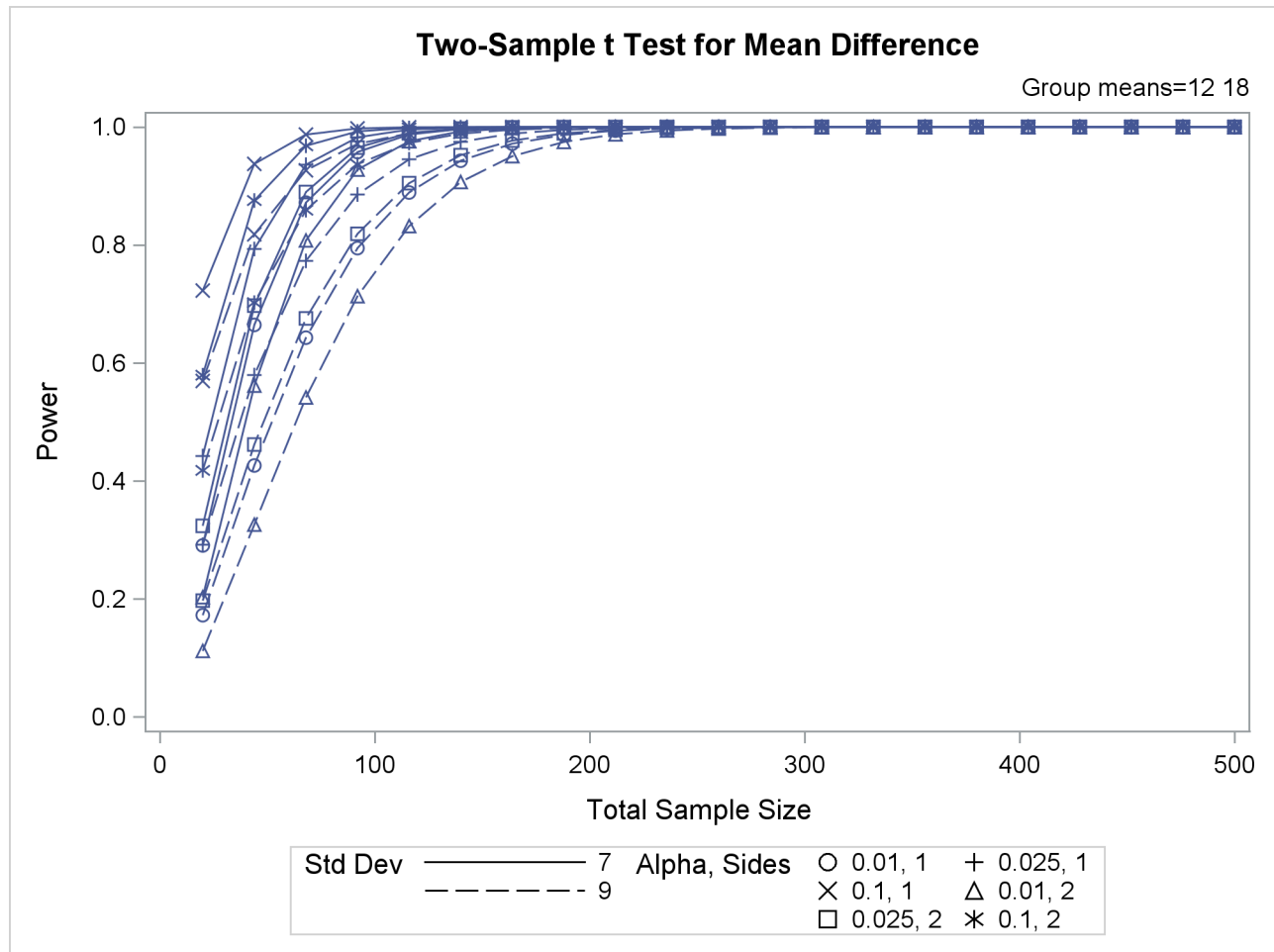
Suppose you want to produce a plot with only one panel that varies color in addition to line style and plotting symbol. Include the **LINESTYLE**, **SYMBOL**, and **COLOR** keywords in the **VARY** option in the **PLOT** statement, as follows, to produce the plot in [Output 90.8.15](#):

```
plot x=n min=20 max=500
    vary (linestyle, symbol, color);
```

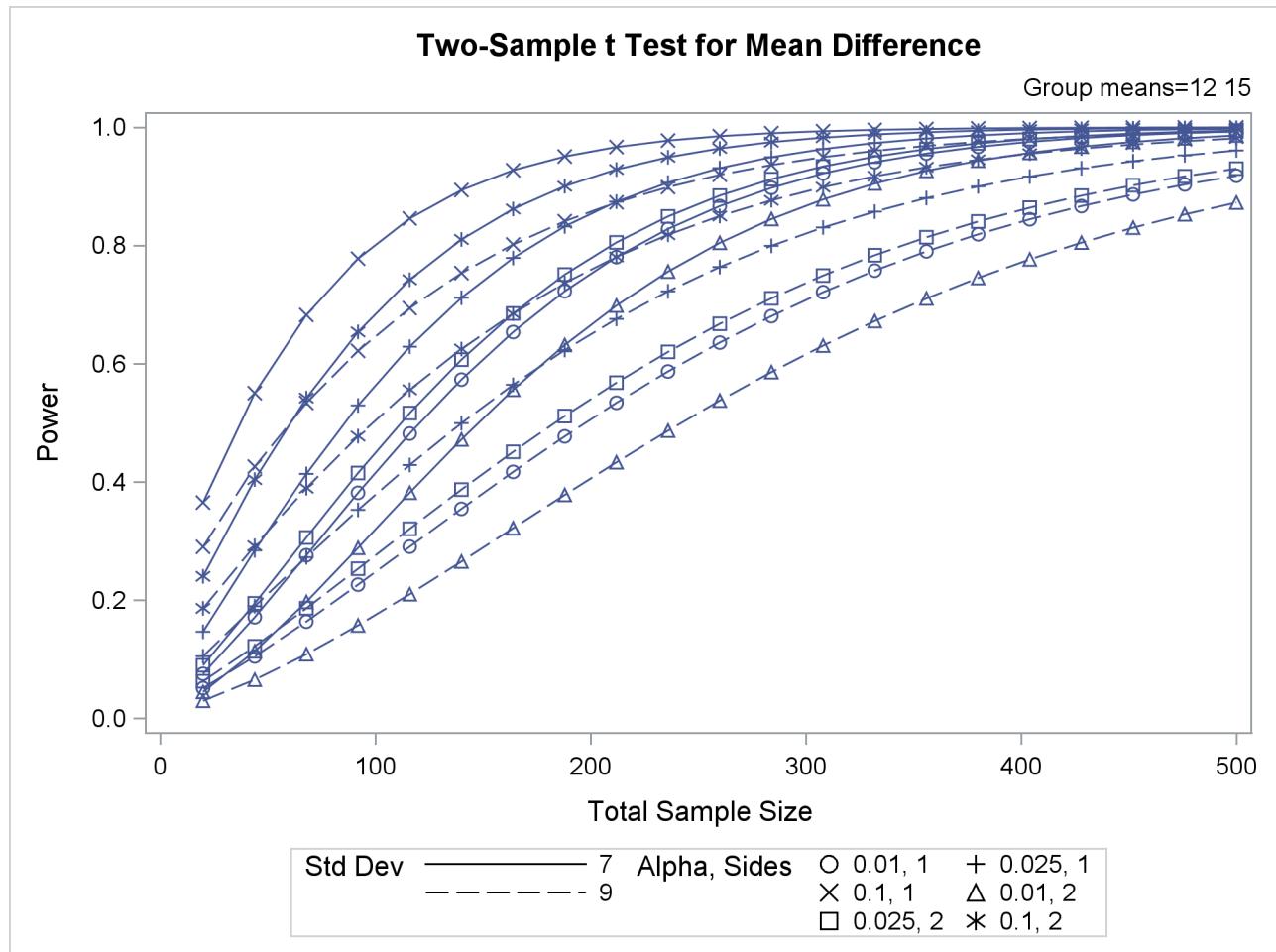
Output 90.8.15 Plot with Varying Color Instead of Panel

Finally, suppose you want to specify which features are used *and* which analysis parameters they are linked to. The following **PLOT** statement produces a two-panel plot (shown in [Output 90.8.16](#)) in which line style varies by standard deviation, plotting symbol varies by both alpha and sides, and panel varies by means:

```
plot x=n min=20 max=500
  vary (linestyle by stddev,
        symbol by alpha sides,
        panel by groupmeans);
```

Output 90.8.16 Plot with Features Explicitly Linked to Parameters

Output 90.8.16 continued

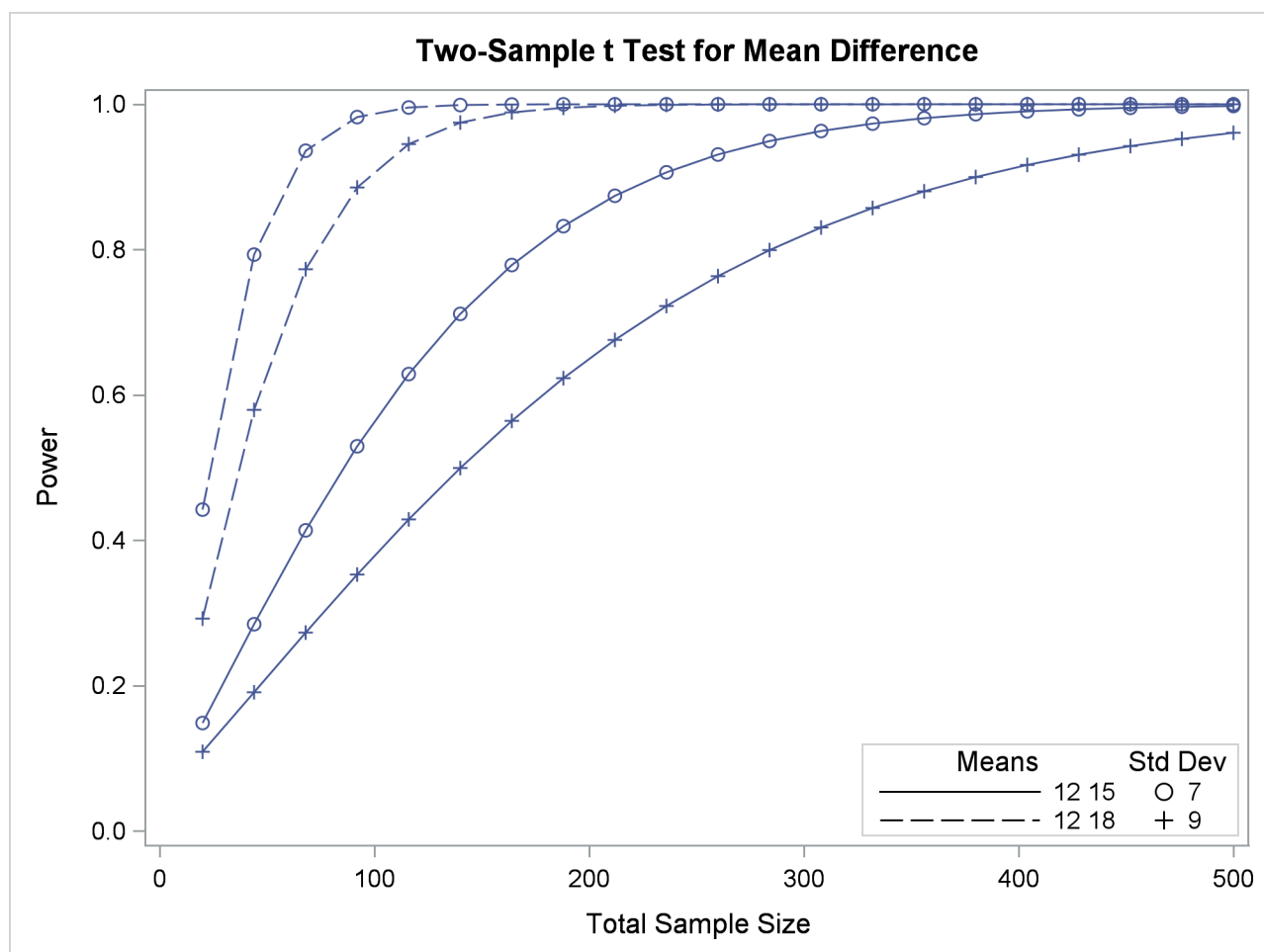


Choosing Key (Legend) Styles

The default style for the key (or “legend”) is one that displays the association between levels of features and levels of analysis parameters, located below the X axis. For example, [Output 90.8.5](#) demonstrates this style of key.

You can reproduce [Output 90.8.5](#) with the same key but a different location, inside the plotting region, by using the `POS=INSET` option within the `KEY=BYFEATURE` option in the `PLOT` statement. The following statements product the plot in [Output 90.8.17](#):

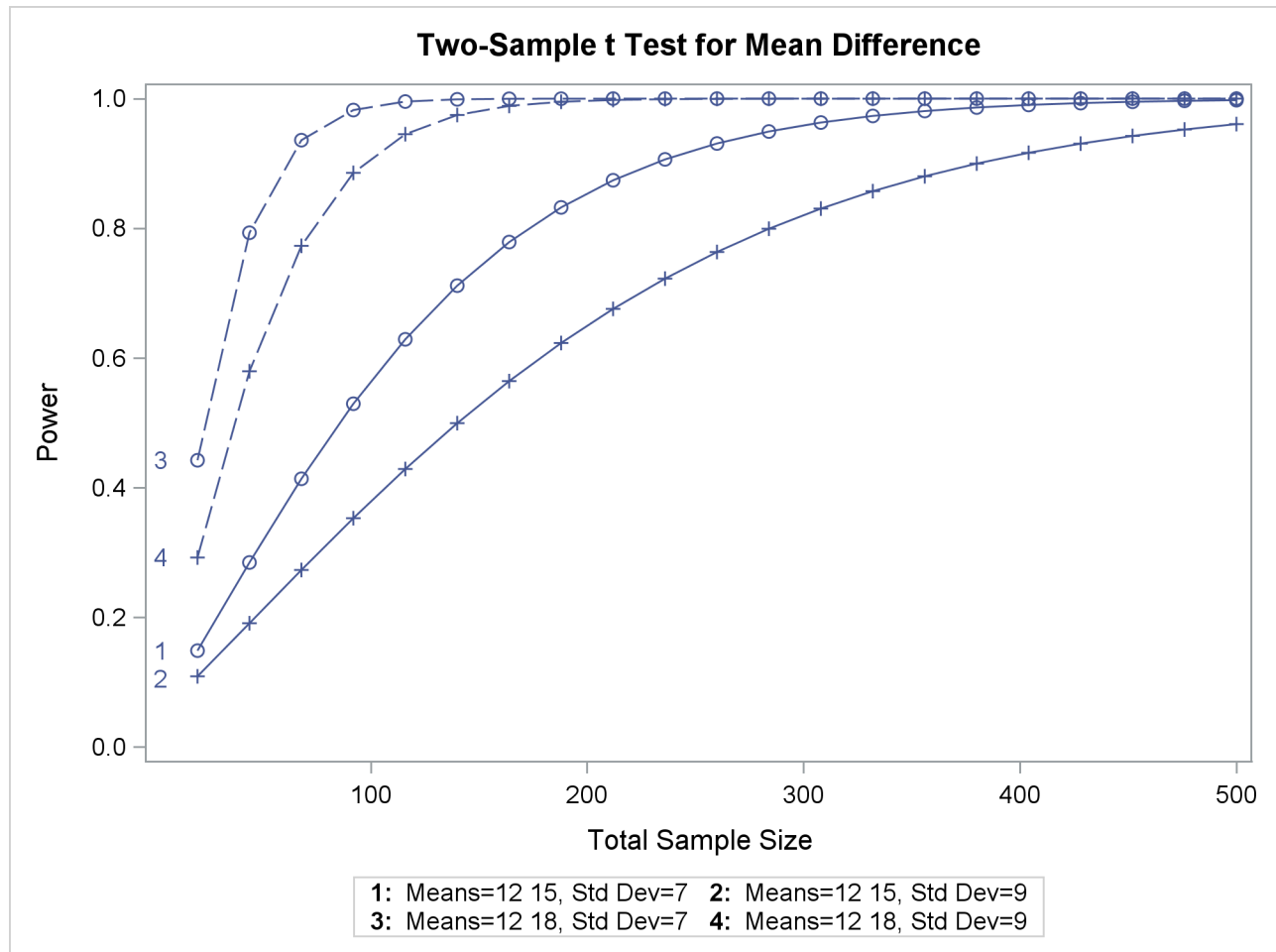
```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev        = 7 9
    power         = .
    ntotal        = 200;
  plot x=n min=20 max=500
    key = byfeature(pos=inset);
run;
```

Output 90.8.17 Plot with a By-Feature Key inside the Plotting Region

Alternatively, you can specify a key that identifies each individual curve separately by number by using the **KEY=BYCURVE** option in the **PLOT** statement:

```
plot x=n min=20 max=500
     key = bycurve;
```

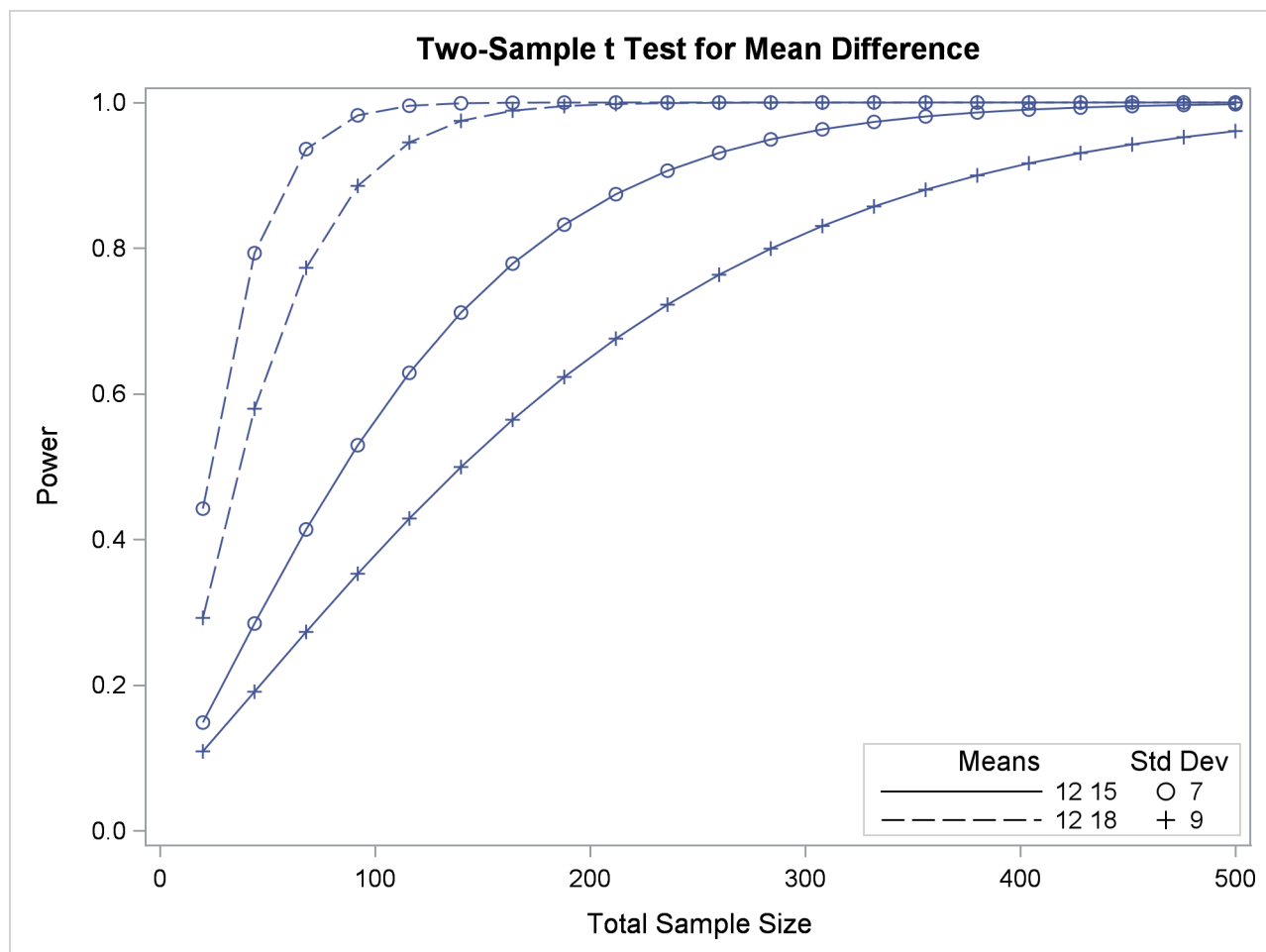
The resulting plot is shown in [Output 90.8.18](#).

Output 90.8.18 Plot with a Numbered By-Curve Key

Use the **NUMBERS=OFF** option within the **KEY=BYCURVE** option to specify a nonnumbered key that identifies curves with samples of line styles, symbols, and colors:

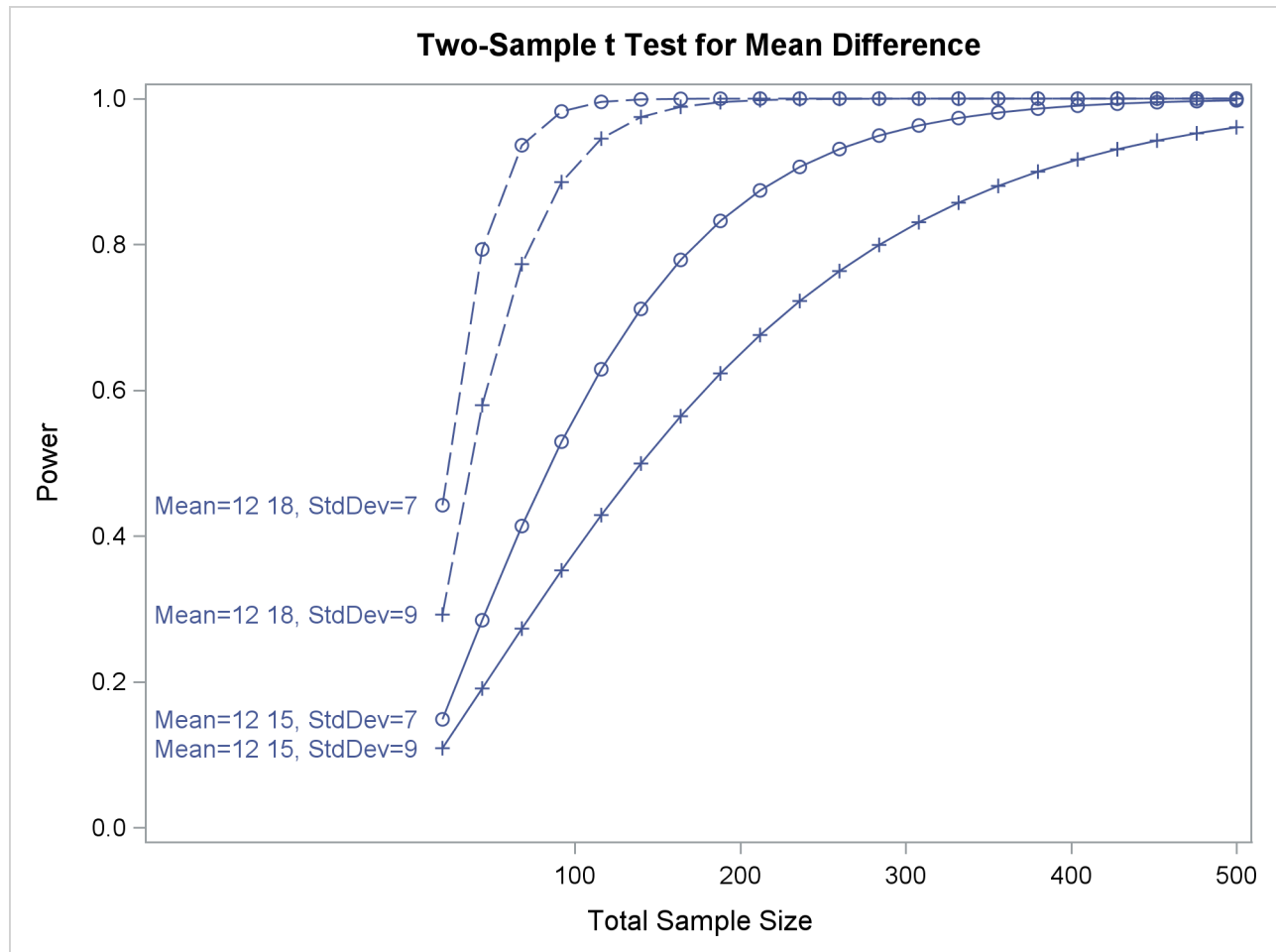
```
plot x=n min=20 max=500
    key = bycurve(numbers=off pos=inset);
```

The **POS=INSET** suboption places the key within the plotting region. The resulting plot is shown in Output 90.8.19.

Output 90.8.19 Plot with a Nonnumbered By-Curve Key

Finally, you can attach labels directly to curves with the **KEY=ONCURVES** option. The following **PLOT** statement produces [Output 90.8.20](#):

```
plot x=n min=20 max=500
    key = oncurves;
```

Output 90.8.20 Plot with Directly Labeled Curves

Modifying Symbol Locations

The default locations for plotting symbols are the points computed directly from the power and sample size algorithms. For example, [Output 90.8.5](#) shows plotting symbols corresponding to computed points. The curves connecting these points are interpolated (as indicated by the `INTERPOL=` option in the `PLOT` statement).

You can modify the locations of plotting symbols by using the `MARKERS=` option in the `PLOT` statement. The `MARKERS=ANALYSIS` option places plotting symbols at locations corresponding to the input specified in the analysis statement preceding the `PLOT` statement. You might prefer this as an alternative to using reference lines to highlight specific points. For example, you can reproduce [Output 90.8.5](#), but with the plotting symbols located at the sample sizes shown in [Output 90.8.1](#), by using the following statements:

```
proc power plotonly;
  twosamplemeans test=diff
    groupmeans    = 12 | 15 18
    stddev         = 7 9
    power          = .
    ntotal         = 232 382 60 98;
```

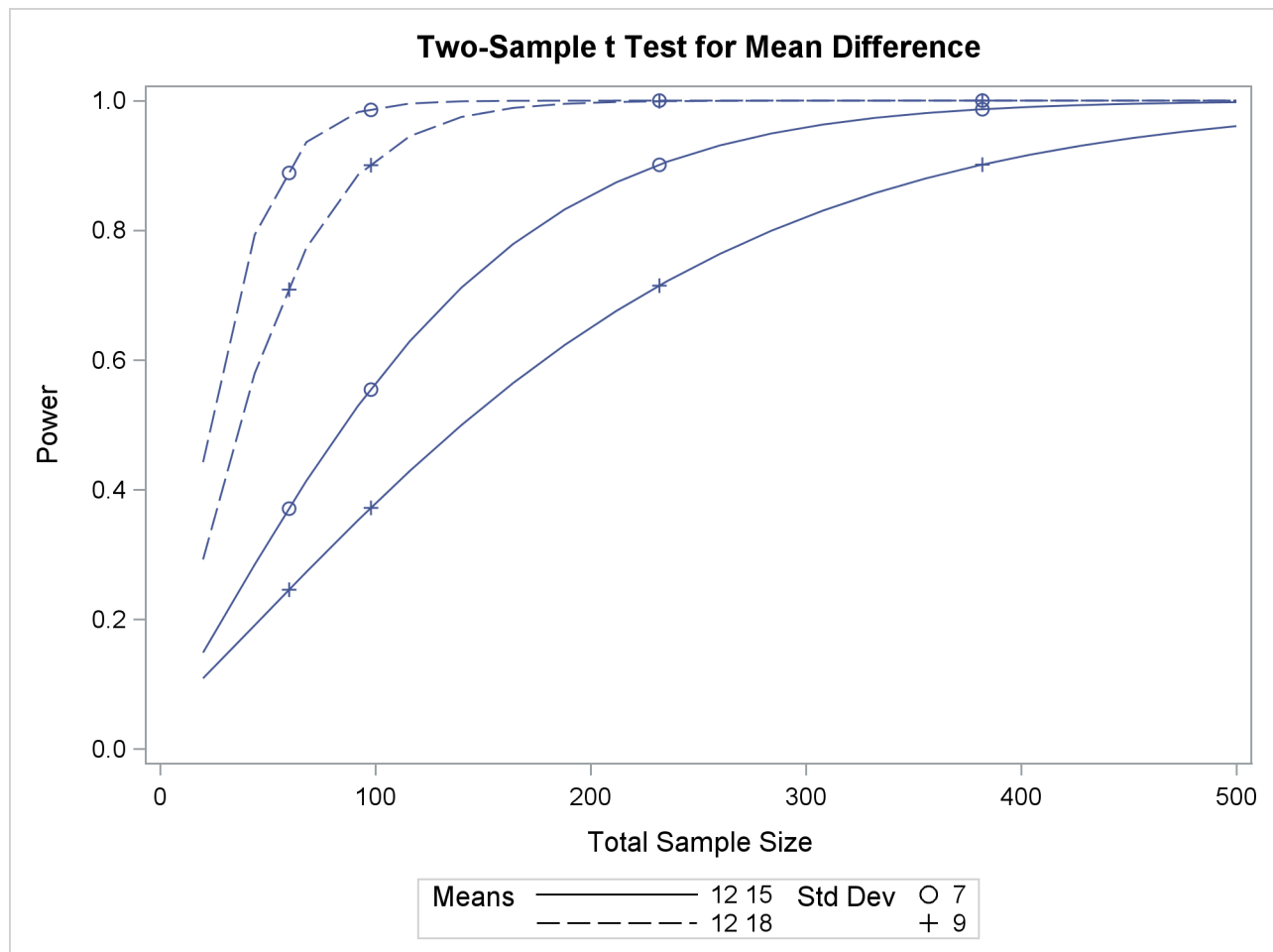
```

plot x=n min=20 max=500
    markers=analysis;
run;

```

The analysis statement here is the **TWOSAMPLEMEANS** statement. The **MARKERS=ANALYSIS** option in the **PLOT** statement causes the plotting symbols to occur at sample sizes specified by the **NTOTAL=** option in the **TWOSAMPLEMEANS** statement: 232, 382, 60, and 98. The resulting plot is shown in [Output 90.8.21](#).

Output 90.8.21 Plot with **MARKERS=ANALYSIS**



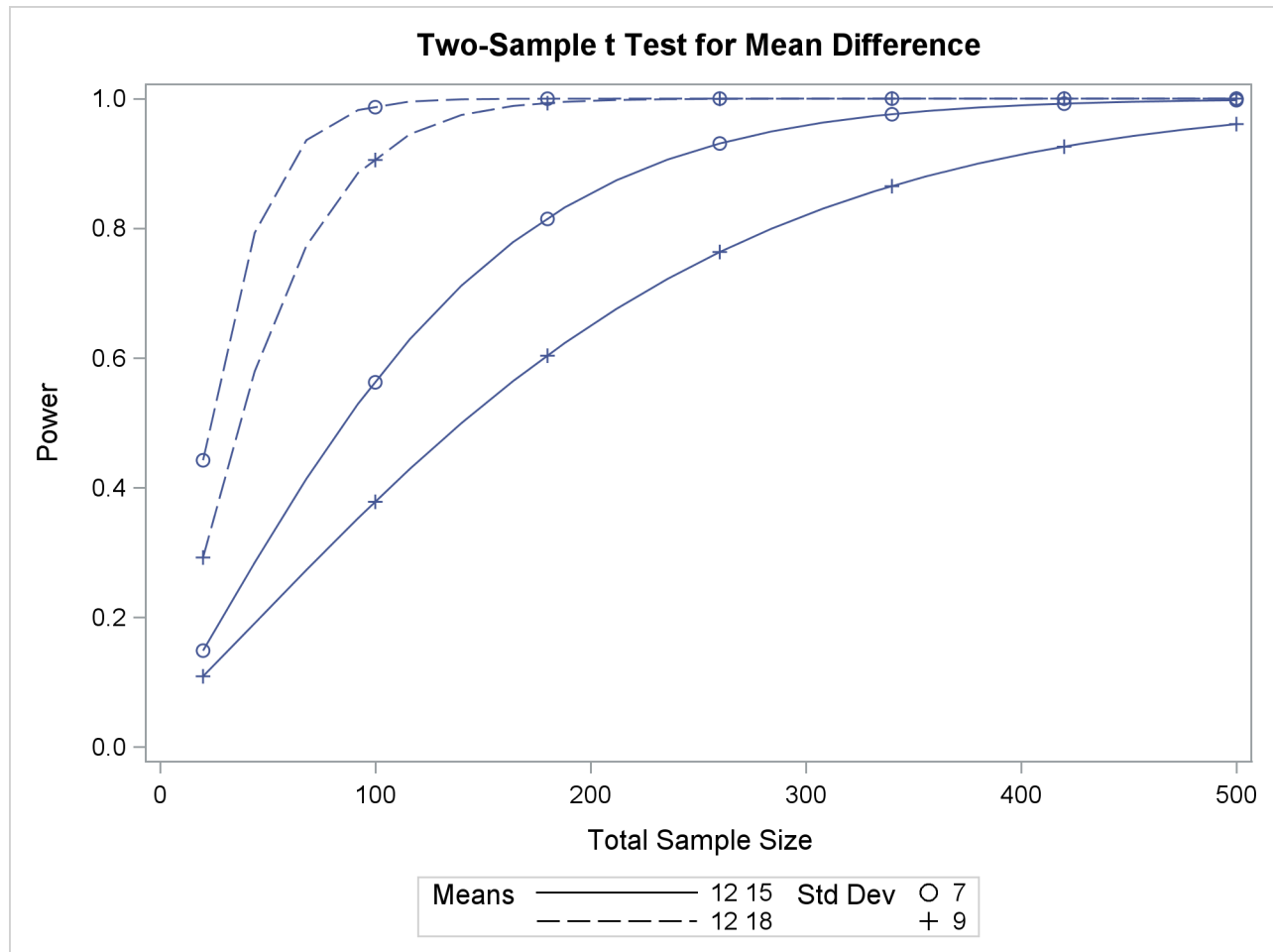
You can also use the **MARKERS=NICE** option to align symbols with the tick marks on one of the axes (the X axis when the **X=** option is used, or the Y axis when the **Y=** option is used):

```

plot x=n min=20 max=500
    markers=nice;

```

The plot created by this **PLOT** statement is shown in [Output 90.8.22](#).

Output 90.8.22 Plot with MARKERS=NICE

Note that the plotting symbols are aligned with the tick marks on the X axis because the `X=` option is specified.

Example 90.9: Binary Logistic Regression with Independent Predictors

Suppose you are planning an industrial experiment similar to the analysis in “Getting Started: LOGISTIC Procedure” on page 5382 in Chapter 73, “The LOGISTIC Procedure,” but for a different type of ingot. The primary test of interest is the likelihood ratio chi-square test of the effect of heating time on the readiness of the ingots for rolling. Ingots will be randomized independently into one of four different heating times (5, 10, 15, and 20 minutes) with allocation ratios 2:3:3:2 and three different soaking times (2, 4, and 6 minutes) with allocation ratios 2:2:1. The mass of each ingot will be measured as a covariate.

You want to know how many ingots you must sample to have a 90% chance of detecting an odds ratio as small as 1.2 for a five-minute heating time increase. The odds ratio is defined here as the odds of the ingot not being ready given a heating time of h minutes divided by the odds given a heating time of $h - 5$ minutes, for any time h . You will use a significance level of $\alpha = 0.1$ to balance Type I and Type II errors since you consider their importance to be roughly equal.

The distributions of heating time and soaking time are determined by the design, but you must conjecture the distribution of ingot mass. Suppose you expect its distribution to be approximately normal with mean 4 kg and standard deviation between 1 kg and 2 kg.

You are powering the study for an odds ratio of 1.2 for the heating time, but you must also conjecture odds ratios for soaking time and mass. You suspect that the odds ratio for a unit increase in soaking time is about 1.4, and the odds ratio for a unit increase in mass is between 1 and 1.3.

Finally, you must provide a guess for the average probability of an ingot not being ready for rolling, averaged across all possible design profiles. Existing data suggest that this probability lies between 0.15 and 0.25.

You decide to evaluate sample size at the two extremes of each parameter for which you conjectured a range. Use the following statements to perform the sample size determination:

```
proc power;
  logistic
    vardist("Heat") = ordinal((5 10 15 20) : (0.2 0.3 0.3 0.2))
    vardist("Soak") = ordinal((2 4 6) : (0.4 0.4 0.2))
    vardist("Mass1") = normal(4, 1)
    vardist("Mass2") = normal(4, 2)
    testpredictor = "Heat"
    covariates = "Soak" | "Mass1" "Mass2"
    responseprob = 0.15 0.25
    testoddsratio = 1.2
    units= ("Heat" = 5)
    covoddsratios = 1.4 | 1 1.3
    alpha = 0.1
    power = 0.9
    ntotal = .;
run;
```

The **VARDIST=** option is used to define the distributions of the predictor variables. The distributions of heating and soaking times are defined by the experimental design, with ordinal probabilities derived from the allocation ratios. The two conjectured standard deviations for the ingot mass are represented in the Mass1 and Mass2 distributions. The **TESTPREDICTOR=** option identifies the predictor being tested, and the **COVARIATES=** option specifies the scenarios for the remaining predictors in the model (soaking time and mass). The **RESPONSEPROB=** option specifies the overall response probability, and the **TESTODDSRATIO=** and **UNITS=** options indicate the odds ratio and increment for heating time. The **COVODDSRATIOS=** option specifies the scenarios for the odds ratios of soaking time and mass. The default **DEFAULTUNIT=1** option specifies a unit change for both of these odds ratios. The **ALPHA=** option sets the significance level, and the **POWER=** option defines the target power. Finally, the **NTOTAL=** option with a missing value (.) identifies the parameter to solve for.

Output 90.9.1 shows the results.

Output 90.9.1 Sample Sizes for Test of Heating Time in Logistic Regression

The POWER Procedure
Likelihood Ratio Chi-square Test for One Predictor

Fixed Scenario Elements										
Method		Shieh-O'Brien approximation								
Alpha		0.1								
Test Predictor		Heat								
Odds Ratio for Test Predictor		1.2								
Unit for Test Pred Odds Ratio		5								
Nominal Power		0.9								

Computed N Total										
Index	Response		Covariates	Cov		Cov		Total		N
	Prob			ORs		Units		N	Actual	
								Bins	Power	Total
1	0.15	Soak	Mass1	1.4	1.0	1	1	120	0.900	1878
2	0.15	Soak	Mass1	1.4	1.3	1	1	120	0.900	1872
3	0.15	Soak	Mass2	1.4	1.0	1	1	120	0.900	1878
4	0.15	Soak	Mass2	1.4	1.3	1	1	120	0.900	1857
5	0.25	Soak	Mass1	1.4	1.0	1	1	120	0.900	1342
6	0.25	Soak	Mass1	1.4	1.3	1	1	120	0.900	1348
7	0.25	Soak	Mass2	1.4	1.0	1	1	120	0.900	1342
8	0.25	Soak	Mass2	1.4	1.3	1	1	120	0.900	1369

The required sample size ranges from 1342 to 1878, depending on the unknown true values of the overall response probability, mass standard deviation, and soaking time odds ratio. The overall response probability clearly has the largest influence among these parameters, with a sample size increase of almost 40% going from 0.25 to 0.15.

Example 90.10: Wilcoxon-Mann-Whitney Test

Consider a hypothetical clinical trial to treat interstitial cystitis (IC), a painful, chronic inflammatory condition of the bladder with no known cause that most commonly affects women. Two treatments will be compared: lidocaine alone ("lidocaine") versus lidocaine plus a fictitious experimental drug called Mironel ("Mir+lido"). The design is balanced, randomized, double-blind, and female-only. The primary outcome is a measure of overall improvement at week 4 of the study, measured on a seven-point Likert scale as shown in [Table 90.40](#).

Table 90.40 Self-Report Improvement Scale

Compared to when I started this study, my condition is:	
Much worse	−3
Worse	−2
Slightly worse	−1
The same	0
Slightly better	+1
Better	+2
Much better	+3

The planned data analysis is a one-sided Wilcoxon-Mann-Whitney test with $\alpha = 0.05$ where the alternative hypothesis represents greater improvement for “Mir+lido.”

You are asked to graphically assess the power of the planned trial for sample sizes between 100 and 250, assuming that the conditional outcome probabilities given treatment are equal to the values in [Table 90.41](#).

Table 90.41 Conjectured Conditional Probabilities

Treatment	Response						
	−3	−2	−1	0	+1	+2	+3
Lidocaine	0.01	0.04	0.20	0.50	0.20	0.04	0.01
Mir+lido	0.01	0.03	0.15	0.35	0.30	0.10	0.06

Use the following statements to compute the power at sample sizes of 100 and 250 and generate a power curve:

```
ods graphics on;

proc power;
  twosamplewilcoxon
    vardist("lidocaine") = ordinal ((−3 −2 −1 0 1 2 3) :
                                   (.01 .04 .20 .50 .20 .04 .01))
    vardist("Mir+lido") = ordinal ((−3 −2 −1 0 1 2 3) :
                                   (.01 .03 .15 .35 .30 .10 .06))
    variables = "lidocaine" | "Mir+lido"
    sides = u
    ntotal = 100 250
    power = .;
  plot step=10;
run;

ods graphics off;
```

The **VARDIST=** option is used to define the distribution for each treatment, and the **VARIABLES=** option specifies the distributions to compare. The **SIDES=U** option corresponds to the alternative hypothesis that the second distribution ("Mir+lido") is more favorable. The **NTOTAL=** option specifies the total sample sizes of interest, and the **POWER=** option with a missing value (.) identifies the parameter to solve for. The default **GROUPWEIGHTS=** and **ALPHA=** options specify a balanced design and significance level $\alpha = 0.05$.

The **STEP=10** option in the **PLOT** statement requests a point for each sample size increment of 10. The default values for the **X=**, **MIN=**, and **MAX=** plot options specify a sample size range of 100 to 250 (the same as in the analysis) for the X axis.

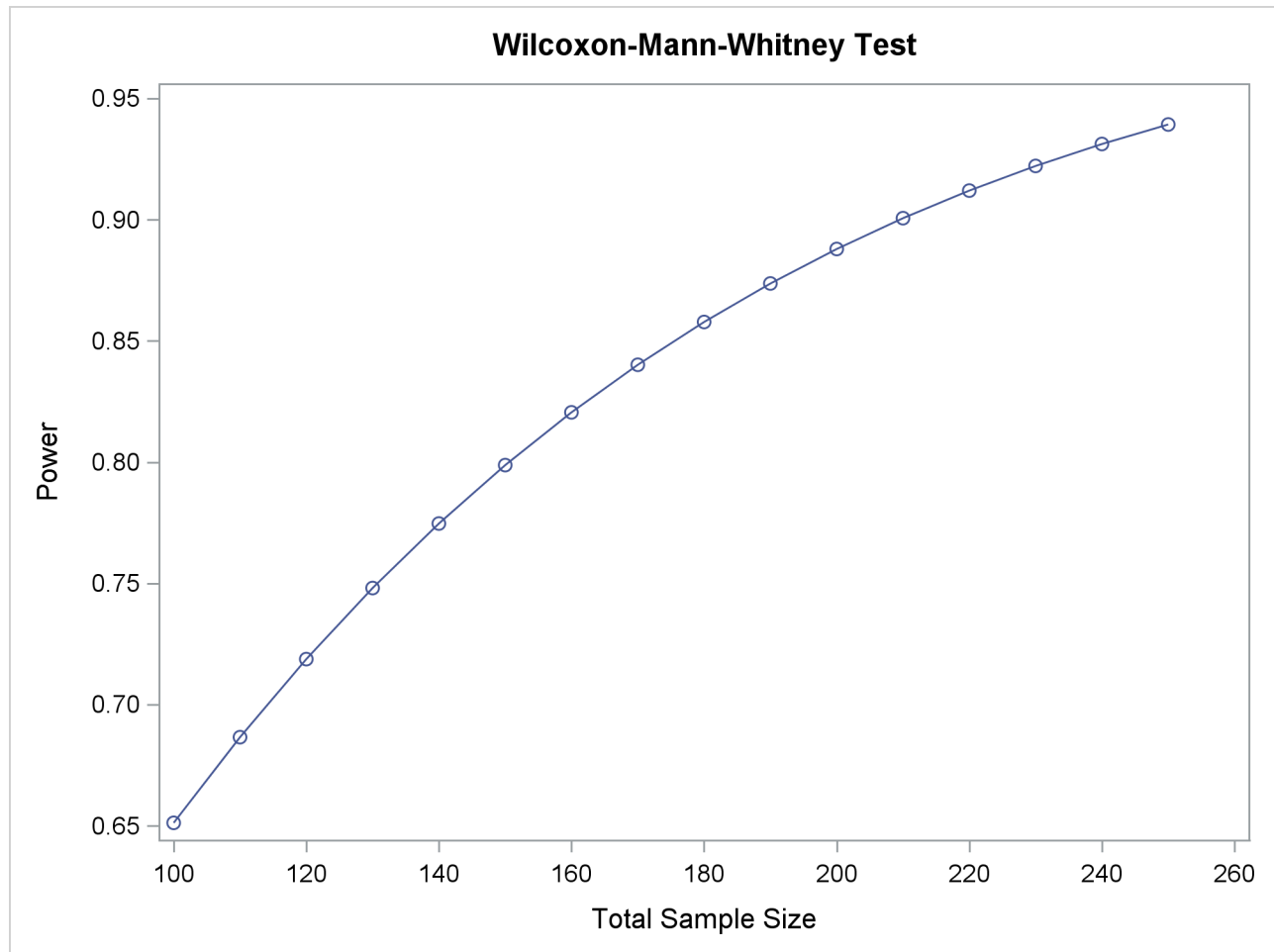
The tabular and graphical results are shown in [Output 90.10.1](#) and [Output 90.10.2](#), respectively.

Output 90.10.1 Power Values for Wilcoxon-Mann-Whitney Test

The POWER Procedure Wilcoxon-Mann-Whitney Test

Fixed Scenario Elements		
Method	O'Brien-Castellote approximation	
Number of Sides	U	
Group 1 Variable	lidocaine	
Group 2 Variable	Mir+lido	
Pooled Number of Bins	7	
Alpha	0.05	
Group 1 Weight	1	
Group 2 Weight	1	
NBins per Group	1000	

Computed Power		
N		
Index	Total	Power
1	100	0.651
2	250	0.939

Output 90.10.2 Plot of Power versus Sample Size for Wilcoxon Power Analysis

The achieved power ranges from 0.651 to 0.939, increasing with sample size.

Example 90.11: Logistic Regression Using the CUSTOM Statement

The framework in Lyles, Lin, and Williamson (2007) provides an effective strategy for using the **CUSTOM** statement to compute power or sample size for generalized linear models. The process involves the following steps:

- 1 Specify the SAS code that you would use to perform the data analysis.
- 2 Create an exemplary data set that resembles the nature of the data you expect to obtain, covering as many reasonable values of the response variable as possible.
- 3 Run the analysis on the exemplary data set and extract the appropriate test statistic.
- 4 If you are planning for a likelihood ratio test, run the analysis again, but this time with a reduced model (lacking the predictors being tested). Again extract the appropriate test statistic.

- 5** Run PROC POWER with the **CUSTOM** statement, specifying the noncentral chi-square distribution with primary noncentrality equal to the appropriate transformation of the test statistics.

Although the power computation is conditional (assuming fixed predictors), you can approximate unconditional power by using a **WEIGHT** statement in the data analysis and constructing the exemplary data to be representative of the conjectured distribution of the predictors.

Suppose you want to compute the power of the Wald chi-square and likelihood ratio tests for the interaction term in a logistic regression model that has two binary covariates, as in Lyles, Lin, and Williamson (2007, section 3.2). You plan to use a sample size $N = 100$ and significance level $\alpha = 0.05$. You have conjectured values for the β_i regression coefficients as shown in Table 90.42.

Table 90.42 Conjectured Regression Coefficient Values

Source	Coefficient	Value
Intercept	β_0	0
X1	β_1	-1.5
X2	β_2	0.5
X1*X2	β_3	2

You expect the covariates to be distributed according to a multinomial distribution with probabilities as shown in Table 90.43.

Table 90.43 Conjectured Covariate Distribution

X1	X2	Probability
0	0	0.2
0	1	0.25
1	0	0.25
1	1	0.3

The strategy for the power computation is summarized as follows:

- 1** Create an exemplary data set to represent both the β regression coefficients and the covariate distribution.
- 2** Use the LOGISTIC procedure to compute the primary noncentrality values and degrees of freedom you need for the power calculation.
- 3** Fetch the primary noncentrality values and degrees of freedom from PROC LOGISTIC output and pass them into the **CUSTOM** statement in PROC POWER along with the other power analysis parameters.

The following steps describe this strategy in detail:

- 1** First create a data set that has one row for each unique design profile and compute the response probability that is associated with each profile according to your conjectured regression coefficients values. The following DATA step computes these probabilities as values of a variable PY:

```

* Compute P(Yj|Xi) values from conjectured betas;
data Exemplary;
  _B0 = 0; _B1 = -1.5; _B2 = 0.5; _B3 = 2;
  do X1 = 0 to 1;
    do X2 = 0 to 1;
      _pi = logistic(_B0 + _B1*X1 + _B2*X2 + _B3*X1*X2);
      do Y = 0 to 1;
        PY = _pi**Y * (1-_pi)**(1-Y);
        output;
      end;
    end;
  end;
  keep X1 X2 Y PY;
run;

proc print data=Exemplary;
run;

```

The output data set is shown in [Output 90.11.1](#).

Output 90.11.1 Response Probabilities Produced from Regression Coefficients

Obs	X1	X2	Y	PY
1	0	0	0	0.50000
2	0	0	1	0.50000
3	0	1	0	0.37754
4	0	1	1	0.62246
5	1	0	0	0.81757
6	1	0	1	0.18243
7	1	1	0	0.26894
8	1	1	1	0.73106

Next expand this data set into an exemplary data set that represents your assumed covariate distribution. The following DATA step replicates each design profile the appropriate number of times according to the multinomial probabilities in [Table 90.43](#).

```

* Expand according to relative allocations of design profiles;
data Exemplary;
  set Exemplary;
  if      (X1 = 0 and X2 = 0) then _alloc=4;
  else if (X1 = 0 and X2 = 1) then _alloc=5;
  else if (X1 = 1 and X2 = 0) then _alloc=5;
  else                                     _alloc=6;
  do _i = 1 to _alloc;
    output;
  end;
  keep X1 X2 Y PY;
run;

```

```
proc print data=Exemplary;
run;
```

Output 90.11.2 shows the resulting exemplary data set.

Output 90.11.2 Exemplary Data Set

Obs	X1	X2	Y	PY
1	0	0	0	0.50000
2	0	0	0	0.50000
3	0	0	0	0.50000
4	0	0	0	0.50000
5	0	0	1	0.50000
6	0	0	1	0.50000
7	0	0	1	0.50000
8	0	0	1	0.50000
9	0	1	0	0.37754
10	0	1	0	0.37754
11	0	1	0	0.37754
12	0	1	0	0.37754
13	0	1	0	0.37754
14	0	1	1	0.62246
15	0	1	1	0.62246
16	0	1	1	0.62246
17	0	1	1	0.62246
18	0	1	1	0.62246
19	1	0	0	0.81757
20	1	0	0	0.81757
21	1	0	0	0.81757
22	1	0	0	0.81757
23	1	0	0	0.81757
24	1	0	1	0.18243
25	1	0	1	0.18243
26	1	0	1	0.18243
27	1	0	1	0.18243
28	1	0	1	0.18243
29	1	1	0	0.26894
30	1	1	0	0.26894
31	1	1	0	0.26894
32	1	1	0	0.26894
33	1	1	0	0.26894
34	1	1	0	0.26894
35	1	1	1	0.73106
36	1	1	1	0.73106
37	1	1	1	0.73106
38	1	1	1	0.73106
39	1	1	1	0.73106
40	1	1	1	0.73106

- 2 Next use the exemplary data as the input data set in PROC LOGISTIC, assigning PY as the WEIGHT variable. First fit the full model and save both the parameter estimates and fit statistics output tables as data sets. (You will later use a parameter estimate for the Wald chi-square power analysis and a fit statistic for the likelihood ratio test.)

```
* Run full model;
proc logistic data=Exemplary;
  ods output parameterestimates=fitFull fitstatistics=statsFull;
  weight PY;
  model Y = X1 X2 X1*X2;
run;
```

Output 90.11.3 and Output 90.11.4 show these tables from the PROC LOGISTIC output.

Output 90.11.3 Parameter Estimates for Full Model

The LOGISTIC Procedure

Analysis of Maximum Likelihood Estimates					
Parameter	DF	Estimate	Standard Error	Wald Chi-Square	Pr > ChiSq
Intercept	1	2.6E-17	1.0000	0.0000	1.0000
X1	1	1.4999	1.5300	0.9611	0.3269
X2	1	-0.5000	1.3605	0.1351	0.7132
X1*X2	1	-1.9999	2.0099	0.9901	0.3197

Output 90.11.4 Fit Statistics for Full Model

Model Fit Statistics		
Criterion	Intercept Only	Intercept and Covariates
AIC	29.692	31.911
SC	31.381	38.666
-2 Log L	27.692	23.911

The numbers you need for the power analysis for the Wald chi-square test of X1*X2 are the degrees of freedom (1) and the Wald chi-square test statistic (0.9901). The number you need for the likelihood ratio test is the -2LogL value (23.911). The following DATA step saves these values to data sets:

```
* Fetch Wald chi-square statistic for test of X1*X2;
data fitFull;
  set fitFull;
  where Variable = "X1*X2";
  keep DF WaldChiSq;
run;
```

```

* Fetch -2LogL for full model;
data statsFull;
  set statsFull;
  where Criterion = "-2 Log L";
  Neg2LogLFull = InterceptAndCovariates;
  keep Neg2LogLFull;
run;

```

Next run the reduced model (dropping the interaction term, the one you want to test) and save the fit statistics table.

```

* Run reduced model;
proc logistic data=Exemplary;
  ods output fitstatistics=statsRed;
  weight PY;
  model Y = X1 X2;
run;

```

Output 90.11.5 shows the fit statistics table from the PROC LOGISTIC output.

Output 90.11.5 Fit Statistics for Reduced Model

The LOGISTIC Procedure

Model Fit Statistics		
Criterion	Intercept Only	Intercept and Covariates
AIC	29.692	30.939
SC	31.381	36.005
-2 Log L	27.692	24.939

The number you need for the power analysis for the likelihood ratio test is the -2LogL value (24.939). The following DATA step saves these values to data sets:

```

* Fetch -2LogL for reduced model;
data statsRed;
  set statsRed;
  where Criterion = "-2 Log L";
  Neg2LogLRed = InterceptAndCovariates;
  keep Neg2LogLRed;
run;

```

Now you are ready to compute the primary noncentrality values. You need to divide both the Wald chi-square and -2LogL difference by the “effective sample size” of the exemplary data set—that is, the number of rows for each response value (20). The following statements compute the primary noncentrality and store all three numbers that you need for PROC POWER in macro variables:

```

* Compute primary noncentralities for Wald and LR tests;
data PrimNC;
  merge fitFull statsFull statsRed;
  WaldPrimNC = WaldChiSq / 20;
  LRPrimNC = -(Neg2LogLFull-Neg2LogLRed) / 20;
  keep DF WaldPrimNC LRPrimNC;
  call symput ("DF", DF);
  call symput ("WaldPrimNC", WaldPrimNC);
  call symput ("LRPrimNC", LRPrimNC);
run;

proc print data=PrimNC;
run;

```

The primary noncentrality values thus computed are shown in [Output 90.11.6](#).

Output 90.11.6 Primary Noncentrality Values Derived from PROC LOGISTIC

Obs	DF	WaldPrimNC	LRPrimNC
1	1	0.049506	0.051396

3 Perform the power analysis by using the [CUSTOM](#) statement in PROC POWER as follows:

```

proc power;
  custom
    dist    = chisquare
    primnc  = &WaldPrimNC &LRPrimNC
    testdf  = &DF
    ntotal  = 100
    power   = .;
run;

```

The computed powers for the two tests, shown in [Output 90.11.7](#), match the results in Lyles, Lin, and Williamson (2007, section 3.2).

Output 90.11.7 Computed Power Values

**The POWER Procedure
Custom Test**

Fixed Scenario Elements	
Distribution	Chi-square
Method	Default
Test Degrees of Freedom	1
Total Sample Size	100
Alpha	0.05
Primary Noncentrality Multiplier	1
Critical Value Multiplier	1

Output 90.11.7 *continued*

Computed Power		
Primary		
Index	NC	Power
1	0.0495	0.605
2	0.0514	0.621

The following statements validate the PROC POWER results by using the QUANTILE and SDF functions in the DATA step:

```
data Power;
  set PrimNC;
  Crit = quantile("chisq", 1-0.05, DF);
  WaldPower = sdf("chisq", Crit, DF, 100 * WaldPrimNC);
  LRPower = sdf("chisq", Crit, DF, 100 * LRPrimNC);
  keep WaldPower LRPower;
run;

proc print data=Power;
run;
```

The computed powers for the two tests, shown in [Output 90.11.8](#), match the results in Lyles, Lin, and Williamson (2007, section 3.2).

Output 90.11.8 Validation of Power Values

Obs	WaldPower	LRPower
1	0.60452	0.62063

References

- Agresti, A. (1980). "Generalized Odds Ratios for Ordinal Data." *Biometrics* 36:59–67.
- Agresti, A., and Coull, B. A. (1998). "Approximate Is Better Than 'Exact' for Interval Estimation of Binomial Proportions." *American Statistician* 52:119–126.
- Anderson, T. W. (1984). *An Introduction to Multivariate Statistical Analysis*. 2nd ed. New York: John Wiley & Sons.
- Beal, S. L. (1989). "Sample Size Determination for Confidence Intervals on the Population Means and on the Difference between Two Population Means." *Biometrics* 45:969–977.
- Blackwelder, W. C. (1982). "'Proving the Null Hypothesis' in Clinical Trials." *Controlled Clinical Trials* 3:345–353.
- Brown, L. D., Cai, T. T., and DasGupta, A. (2001). "Interval Estimation for a Binomial Proportion." *Statistical Science* 16:101–133.
- Cantor, A. B. (1997). *Extending SAS Survival Analysis Techniques for Medical Research*. Cary, NC: SAS Institute Inc.

- Castelloe, J. M. (2000). "Sample Size Computations and Power Analysis with the SAS System." In *Proceedings of the Twenty-Fifth Annual SAS Users Group International Conference*. Cary, NC: SAS Institute Inc. <http://www2.sas.com/proceedings/sugi25/25/st/25p265.pdf>.
- Castelloe, J. M., and O'Brien, R. G. (2001). "Power and Sample Size Determination for Linear Models." In *Proceedings of the Twenty-Sixth Annual SAS Users Group International Conference*. Cary, NC: SAS Institute Inc. <http://www2.sas.com/proceedings/sugi26/p240-26.pdf>.
- Chernick, M. R., and Liu, C. Y. (2002). "The Saw-Toothed Behavior of Power versus Sample Size and Software Solutions: Single Binomial Proportion Using Exact Methods." *American Statistician* 56:149–155.
- Chow, S.-C., Shao, J., and Wang, H. (2003). *Sample Size Calculations in Clinical Research*. Boca Raton, FL: CRC Press.
- Connor, R. J. (1987). "Sample Size for Testing Differences in Proportions for the Paired-Sample Design." *Biometrics* 43:207–211.
- Diegert, C., and Diegert, K. V. (1981). "Note on Inversion of Casagrande-Pike-Smith Approximate Sample-Size Formula for Fisher-Irwin Test on 2×2 Tables." *Biometrics* 37:595.
- Diletti, D., Hauschke, D., and Steinijans, V. W. (1991). "Sample Size Determination for Bioequivalence Assessment by Means of Confidence Intervals." *International Journal of Clinical Pharmacology, Therapy, and Toxicology* 29:1–8.
- DiSantostefano, R. L., and Muller, K. E. (1995). "A Comparison of Power Approximations for Satterthwaite's Test." *Communications in Statistics—Simulation and Computation* 24:583–593.
- Farrington, C. P., and Manning, G. (1990). "Test Statistics and Sample Size Formulae for Comparative Binomial Trials with Null Hypothesis of Non-zero Risk Difference or Non-unity Relative Risk." *Statistics in Medicine* 9:1447–1454.
- Fisher, R. A. (1921). "On the 'Probable Error' of a Coefficient of Correlation Deduced from a Small Sample." *Metron* 1:3–32.
- Fleiss, J. L., Tytun, A., and Ury, H. K. (1980). "A Simple Approximation for Calculating Sample Sizes for Comparing Independent Proportions." *Biometrics* 36:343–346.
- Gatsonis, C., and Sampson, A. R. (1989). "Multiple Correlation: Exact Power and Sample Size Calculations." *Psychological Bulletin* 106:516–524.
- Hocking, R. R. (1985). *The Analysis of Linear Models*. Monterey, CA: Brooks/Cole.
- Hsieh, F. Y. (1989). "Sample Size Tables for Logistic Regression." *Statistics in Medicine* 8:795–802.
- Hsieh, F. Y., and Lavori, P. W. (2000). "Sample-Size Calculations for the Cox Proportional Hazards Regression Model with Nonbinary Covariates." *Controlled Clinical Trials* 21:552–560.
- Johnson, N. L., and Kotz, S. (1970). *Distributions in Statistics: Continuous Univariate Distributions*. New York: John Wiley & Sons.
- Johnson, N. L., Kotz, S., and Balakrishnan, N. (1994). *Continuous Univariate Distributions*. 2nd ed. Vol. 1. New York: John Wiley & Sons.

- Johnson, N. L., Kotz, S., and Balakrishnan, N. (1995). *Continuous Univariate Distributions*. 2nd ed. Vol. 2. New York: John Wiley & Sons.
- Johnson, N. L., Kotz, S., and Kemp, A. W. (1992). *Univariate Discrete Distributions*. 2nd ed. New York: John Wiley & Sons.
- Jones, R. M., and Miller, K. S. (1966). "On the Multivariate Lognormal Distribution." *Journal of Industrial Mathematics* 16:63–76.
- Kolassa, J. E. (1995). "A Comparison of Size and Power Calculations for the Wilcoxon Statistic for Ordered Categorical Data." *Statistics in Medicine* 14:1577–1581.
- Kotz, S., Balakrishnan, N., and Johnson, N. L. (2000). *Continuous Multivariate Distributions*. 2nd ed. New York: Wiley-Interscience.
- Lachin, J. M. (1992). "Power and Sample Size Evaluation for the McNemar Test with Application to Matched Case-Control Studies." *Statistics in Medicine* 11:1239–1251.
- Lakatos, E. (1988). "Sample Sizes Based on the Log-Rank Statistic in Complex Clinical Trials." *Biometrics* 44:229–241.
- Lenth, R. V. (2001). "Some Practical Guidelines for Effective Sample Size Determination." *American Statistician* 55:187–193.
- Lyles, R. H., Lin, H.-M., and Williamson, J. M. (2007). "A Practical Approach to Computing Power for Generalized Linear Models with Nominal, Count, or Ordinal Responses." *Statistics in Medicine* 26:1632–1648.
- Maxwell, S. E. (2000). "Sample Size and Multiple Regression Analysis." *Psychological Methods* 5:434–458.
- Miettinen, O. S. (1968). "The Matched Pairs Design in the Case of All-or-None Responses." *Biometrics* 24:339–352.
- Moser, B. K., Stevens, G. R., and Watts, C. L. (1989). "The Two-Sample T Test versus Satterthwaite's Approximate F Test." *Communications in Statistics—Theory and Methods* 18:3963–3975.
- Muller, K. E., and Benignus, V. A. (1992). "Increasing Scientific Power with Statistical Power." *Neurotoxicology and Teratology* 14:211–219.
- O'Brien, R. G., and Casteloe, J. M. (2006). "Exploiting the Link between the Wilcoxon-Mann-Whitney Test and a Simple Odds Statistic." In *Proceedings of the Thirty-First Annual SAS Users Group International Conference*. Cary, NC: SAS Institute Inc. <http://www2.sas.com/proceedings/sugi31/209-31.pdf>.
- O'Brien, R. G., and Casteloe, J. M. (2007). "Sample-Size Analysis for Traditional Hypothesis Testing: Concepts and Issues." In *Pharmaceutical Statistics Using SAS: A Practical Guide*, edited by A. Dmitrienko, C. Chuang-Stein, and R. D'Agostino, 237–271. Cary, NC: SAS Institute Inc.
- O'Brien, R. G., and Muller, K. E. (1993). "Unified Power Analysis for t -Tests through Multivariate Hypotheses." In *Applied Analysis of Variance in Behavioral Science*, edited by L. K. Edwards, 297–344. New York: Marcel Dekker.
- Owen, D. B. (1965). "A Special Case of a Bivariate Non-central t -Distribution." *Biometrika* 52:437–446.

- Pagano, M., and Gauvreau, K. (1993). *Principles of Biostatistics*. Belmont, CA: Wadsworth.
- Phillips, K. F. (1990). "Power of the Two One-Sided Tests Procedure in Bioequivalence." *Journal of Pharmacokinetics and Biopharmaceutics* 18:137–144.
- Satterthwaite, F. E. (1946). "An Approximate Distribution of Estimates of Variance Components." *Biometrics Bulletin* 2:110–114.
- Schork, M. A., and Williams, G. W. (1980). "Number of Observations Required for the Comparison of Two Correlated Proportions." *Communications in Statistics—Simulation and Computation* 9:349–357.
- Schuirmann, D. J. (1987). "A Comparison of the Two One-Sided Tests Procedure and the Power Approach for Assessing the Equivalence of Average Bioavailability." *Journal of Pharmacokinetics and Biopharmaceutics* 15:657–680.
- Self, S. G., Mauritsen, R. H., and Ohara, J. (1992). "Power Calculations for Likelihood Ratio Tests in Generalized Linear Models." *Biometrics* 48:31–39.
- Senn, S. (1993). *Cross-Over Trials in Clinical Research*. New York: John Wiley & Sons.
- Shieh, G. (2000). "A Comparison of Two Approaches for Power and Sample Size Calculations in Logistic Regression Models." *Communications in Statistics—Simulation and Computation* 29:763–791.
- Shieh, G., and O'Brien, R. G. (1998). "A Simpler Method to Compute Power for Likelihood Ratio Tests in Generalized Linear Models." In *Proceedings of the Annual Joint Statistical Meetings of the American Statistical Association*. Alexandria, VA: ASA.
- Stuart, A., and Ord, J. K. (1994). *Distribution Theory*. Vol. 1 of Kendall's Advanced Theory of Statistics. 6th ed. Baltimore: Edward Arnold.
- Walters, D. E. (1979). "In Defence of the Arc Sine Approximation." *The Statistician* 28:219–232.
- Wellek, S. (2003). *Testing Statistical Hypotheses of Equivalence*. Boca Raton, FL: Chapman & Hall/CRC.

Subject Index

- F* distribution
 - power and sample size (POWER), 7231, 7339
- T* distribution
 - power and sample size (POWER), 7231, 7340
- AB/BA crossover designs
 - power and sample size (POWER), 7410
- actual power
 - POWER procedure, 7225, 7332, 7334
- alpha level
 - POWER procedure, 7326
- alternative hypothesis, 7326
- analysis of variance
 - power and sample size (POWER), 7231, 7267, 7271, 7272, 7369, 7396
- analysis statements
 - POWER procedure, 7226
- bar (l) operator
 - POWER procedure, 7330
- binomial proportion confidence interval
 - power and sample size (POWER), 7362–7365
- binomial proportion confidence interval precision
 - power and sample size (POWER), 7260
- binomial proportion test
 - power and sample size (POWER), 7231, 7253, 7258, 7347, 7348, 7350, 7401
- bioequivalence, *see* equivalence tests
- ceiling sample size
 - POWER procedure, 7225, 7334
- chi-square distribution
 - power and sample size (POWER), 7231, 7339
- chi-square tests
 - power and sample size (POWER), 7231, 7253, 7259, 7292, 7298, 7348, 7350, 7382, 7401
- confidence intervals
 - means, power and sample size (POWER), 7261, 7267, 7279, 7287, 7300, 7309, 7368, 7378, 7389, 7423
- contrasts
 - power and sample size (POWER), 7231, 7267, 7268, 7271, 7369, 7396
- correlated proportions, *see* McNemar's test
- correlation coefficients
 - power and sample size (POWER), 7231, 7249, 7339, 7345, 7346
- Cox proportional hazards regression
 - power and sample size (POWER), 7227, 7231, 7337
- crossover designs
 - power and sample size (POWER), 7410
- effect size
 - power and sample size (POWER), 7290
- equivalence tests
 - power and sample size (POWER), 7259, 7261, 7266, 7279, 7287, 7300, 7309, 7367, 7376, 7387, 7388, 7410
- Farrington-Manning test
 - power and sample size (POWER), 7292, 7299, 7380, 7381
- Fisher's exact test
 - power and sample size (POWER), 7292, 7299, 7383
- Fisher's *z* test for correlation
 - power and sample size (POWER), 7249, 7252, 7345, 7417
- fractional sample size
 - POWER procedure, 7225, 7334
- Gehan test
 - power and sample size (POWER), 7231, 7310, 7321, 7389
- generalized linear regression
 - power and sample size (POWER), 7231, 7456
- GLMPOWER procedure
 - compared to other procedures, 7218
- graphics, *see* plots
 - POWER procedure, 7395
- graphs, *see* plots
- grouped-name-lists
 - POWER procedure, 7329
- grouped-number-lists
 - POWER procedure, 7329
- half-width, confidence intervals, 7327
- keyword-lists
 - POWER procedure, 7329
- likelihood-ratio chi-square test
 - power and sample size (POWER), 7231, 7237, 7292, 7300, 7382, 7456
- log-rank test for homogeneity
 - power and sample size (POWER), 7231, 7310, 7319, 7389, 7421

- logistic regression
 - power and sample size (POWER), 7231, 7237, 7340, 7451, 7456
- lognormal data
 - power and sample size (POWER), 7263, 7266, 7281, 7286, 7287, 7302, 7308, 7309, 7366, 7367, 7374, 7376, 7386, 7388, 7413
- Mann-Whitney-Wilcoxon test, *see* Wilcoxon-Mann-Whitney (rank-sum) test
- McNemar's test
 - power and sample size (POWER), 7231, 7272, 7277, 7279, 7372
- means
 - power and sample size (POWER), 7231, 7261, 7279, 7300, 7307, 7383
- name-lists
 - POWER procedure, 7329
- nominal power
 - POWER procedure, 7225, 7332, 7334
- Noncentral F distribution
 - power and sample size (POWER), 7231, 7339
- Noncentral T distribution
 - power and sample size (POWER), 7231, 7340
- noncentral chi-square distribution
 - power and sample size (POWER), 7231, 7339
- noncentral distributions
 - power and sample size (POWER), 7231, 7456
- noncentrality
 - power and sample size (POWER), 7231, 7456
- noninferiority tests
 - power and sample size (POWER), 7260, 7413
- nonparametric tests
 - power and sample size (POWER), 7231, 7321, 7326
- Normal distribution
 - power and sample size (POWER), 7231, 7340
- null hypothesis, 7326
- number-lists
 - POWER procedure, 7329
- odds ratio
 - power and sample size (POWER), 7231, 7292, 7298, 7378, 7382
- ODS graph names
 - POWER procedure, 7395
- ODS Graphics
 - POWER procedure, 7395
- one-sample t -test
 - power and sample size (POWER), 7219, 7231, 7261, 7266, 7365–7367
- one-way ANOVA
 - power and sample size (POWER), 7231, 7267, 7271, 7272, 7369, 7396

- paired proportions, *see* McNemar's test
- paired t test
 - power and sample size (POWER), 7231, 7279, 7286, 7374
- paired-difference t test, *see* paired t test
- partial correlations
 - power and sample size (POWER), 7231, 7249, 7253, 7345, 7346, 7417
- Pearson chi-square test
 - power and sample size (POWER), 7292, 7298, 7382
- Pearson correlation statistics
 - power and sample size (POWER), 7231, 7249, 7345, 7346, 7417
- plots
 - power and sample size (POWER), 7218, 7226, 7227, 7288, 7426
- Poisson regression
 - power and sample size (POWER), 7231
- power
 - overview of power concepts (POWER), 7326
 - See POWER procedure, 7216
- power curves, *see* plots
- POWER procedure, 7231
 - F distribution, 7231, 7339
 - T distribution, 7231, 7340
 - AB/BA crossover designs, 7410
 - actual alpha, 7334
 - actual power, 7225, 7332, 7334
 - actual prob(width), 7334
 - analysis of variance, 7231, 7267, 7271, 7272, 7369, 7396
 - analysis statements, 7226
 - bar (|) operator, 7330
 - binomial proportion confidence interval, 7362–7365
 - binomial proportion confidence interval precision, 7260
 - binomial proportion tests, 7231, 7253, 7258, 7347, 7348, 7350, 7401
 - ceiling sample size, 7225, 7334
 - chi-square distribution, 7231, 7339
 - chi-square tests, 7231, 7253
 - compared to other procedures, 7218
 - computational methods, 7335
 - computational resources, 7335
 - confidence intervals for means, 7261, 7267, 7279, 7287, 7300, 7309, 7368, 7378, 7389, 7423
 - contrasts, analysis of variance, 7231, 7267, 7268, 7271, 7369, 7396
 - correlated proportions, 7231, 7272, 7277, 7279, 7372
 - correlation, 7231, 7249, 7339, 7345, 7346, 7417

Cox proportional hazards regression, 7227, 7231, 7337
 crossover designs, 7410
 displayed output, 7334
 effect size, 7290
 equivalence tests, 7259, 7261, 7266, 7279, 7287, 7300, 7309, 7367, 7376, 7387, 7388, 7410
 Farrington-Manning test, 7292, 7299, 7380, 7381
 Fisher's exact test, 7292, 7299, 7383
 Fisher's z test for correlation, 7249, 7252, 7345, 7417
 fractional sample size, 7225, 7334
 Gehan test, 7231, 7310, 7321, 7389
 generalized linear regression, 7231, 7456
 graphics, 7395
 grouped-name-lists, 7329
 grouped-number-lists, 7329
 introductory example, 7219
 keyword-lists, 7329
 likelihood-ratio chi-square test, 7231, 7237, 7292, 7300, 7382, 7456
 log-rank test for comparing survival curves, 7231, 7310, 7319, 7389, 7421
 logistic regression, 7231, 7237, 7340, 7451, 7456
 lognormal data, 7263, 7266, 7281, 7286, 7287, 7302, 7308, 7309, 7366, 7367, 7374, 7376, 7386, 7388, 7413
 McNemar's test, 7231, 7272, 7277, 7279, 7372
 name-lists, 7329
 nominal power, 7225, 7332, 7334
 Noncentral F distribution, 7231, 7339
 Noncentral T distribution, 7231, 7340
 noncentral chi-square distribution, 7231, 7339
 noncentral distributions, 7231, 7456
 noncentrality, 7231, 7456
 noninferiority tests, 7260, 7413
 normal distribution, 7231, 7340
 notation for formulas, 7336
 number-lists, 7329
 odds ratio, 7231, 7292, 7298, 7378, 7382
 ODS graph names, 7395
 ODS Graphics, 7395
 ODS table names, 7335
 one-sample t test, 7219, 7231, 7261, 7266, 7365, 7366
 one-way ANOVA, 7231, 7267, 7271, 7272, 7369, 7396
 overview of power concepts, 7326
 paired proportions, 7231, 7272, 7277, 7279, 7372
 paired t test, 7231, 7279, 7286, 7374
 partial correlation, 7231, 7249, 7253, 7345, 7346, 7417
 Pearson chi-square test, 7292, 7298, 7382

Pearson correlation, 7249, 7252, 7345, 7346, 7417
 plots, 7218, 7226, 7227, 7288, 7426
 Poisson regression, 7231
 primary noncentrality, 7231, 7456
 regression, 7231, 7245, 7249, 7343, 7417, 7456
 relative risk, 7231, 7292, 7298, 7378, 7381, 7382
 risk difference, 7380
 sample size adjustment, 7332
 sample size deflation, 7231
 sample size inflation, 7231
 statistical graphics, 7395
 summary of analyses, 7327
 summary of statements, 7226
 superiority tests, 7260
 survival analysis, 7227, 7231, 7310, 7319, 7389
 t test for correlation, 7231, 7249, 7253, 7346
 t tests, 7231, 7261, 7266, 7279, 7286, 7300, 7307, 7365, 7366, 7374, 7383, 7386, 7426
 Tarone-Ware test, 7231, 7310, 7321, 7389
 two-sample t test, 7221, 7231, 7300, 7307, 7308, 7383, 7385, 7386, 7426
 value lists, 7329
 Wald chi-square test, 7231, 7456
 Wilcoxon-Mann-Whitney (rank-sum) test, 7231, 7321, 7326, 7393
 Wilcoxon-Mann-Whitney test, 7453
 z test, 7231, 7253, 7259, 7348, 7350
 zero-inflated models, 7231
 precision, confidence intervals, 7327
 primary noncentrality
 power and sample size (POWER), 7231, 7456
 prospective power, 7217
 rank-sum test, *see* Wilcoxon-Mann-Whitney (rank-sum) test
 regression
 power and sample size (POWER), 7231, 7245, 7249, 7343, 7417, 7456
 relative risk
 power and sample size (POWER), 7231, 7292, 7298, 7378, 7381, 7382
 retrospective power, 7217
 risk difference
 power and sample size (POWER), 7380
 sample size
 overview of power concepts (POWER), 7326
 See POWER procedure, 7216
 sample size adjustment
 POWER procedure, 7332
 sample size deflation
 power and sample size (POWER), 7231
 sample size inflation

- power and sample size (POWER), 7231
- Satterthwaite t test
 - power and sample size (POWER), 7300, 7308, 7385
- sawtooth power function, 7401
- statistical graphics
 - POWER procedure, 7395
- superiority tests
 - power and sample size (POWER), 7260
- survival analysis
 - power and sample size (POWER), 7227, 7231, 7310, 7319, 7389
- t test
 - power and sample size (POWER), 7231, 7261, 7266, 7279, 7286, 7300, 7307, 7365, 7366, 7374, 7383
- t test for correlation
 - power and sample size (POWER), 7231, 7249, 7253, 7346
- Tarone-Ware test for homogeneity
 - power and sample size (POWER), 7231, 7310, 7321, 7389
- two-sample t -test
 - power and sample size (POWER), 7221, 7231, 7300, 7307, 7308, 7383, 7385, 7386, 7426
- Type I error, 7326
- Type II error, 7326
- value lists
 - POWER procedure, 7329
- Wald chi-square test
 - power and sample size (POWER), 7231, 7456
- Welch t test
 - power and sample size (POWER), 7300, 7308, 7385
- width, confidence intervals, 7327
- Wilcoxon rank-sum test, *see* Wilcoxon-Mann-Whitney (rank-sum) test
- Wilcoxon-Mann-Whitney (rank-sum) test
 - power and sample size (POWER), 7231, 7321, 7393
- Wilcoxon-Mann-Whitney test
 - power and sample size (POWER), 7326, 7453
- z test
 - power and sample size (POWER), 7231, 7253, 7259, 7348, 7350
- zero-inflated models
 - power and sample size (POWER), 7231

Syntax Index

ACCRUALRATEPERGROUP= option
TWOSAMPLESURVIVAL statement (POWER),
[7311](#)

ACCRUALRATETOTAL= option
TWOSAMPLESURVIVAL statement (POWER),
[7312](#)

ACCRUALTIME= option
TWOSAMPLESURVIVAL statement (POWER),
[7312](#)

ALPHA= option
COXREG statement (POWER), [7228](#)
CUSTOM statement (POWER), [7233](#)
LOGISTIC statement (POWER), [7238](#)
MULTREG statement (POWER), [7246](#)
ONECORR statement (POWER), [7250](#)
ONESAMPLEFREQ statement (POWER), [7255](#)
ONESAMPLEMEANS statement (POWER),
[7262](#)
ONEWAYANOVA statement (POWER), [7268](#)
PAIREFREQ statement (POWER), [7273](#)
PAIREDMEANS statement (POWER), [7281](#)
TWOSAMPLEFREQ statement (POWER), [7294](#)
TWOSAMPLEMEANS statement (POWER),
[7302](#)
TWOSAMPLESURVIVAL statement (POWER),
[7312](#)
TWOSAMPLEWILCOXON statement
(POWER), [7323](#)

CI= option
ONESAMPLEFREQ statement (POWER), [7255](#)
ONESAMPLEMEANS statement (POWER),
[7262](#)
PAIREDMEANS statement (POWER), [7281](#)
TWOSAMPLEMEANS statement (POWER),
[7302](#)

CONTRAST= option
ONEWAYANOVA statement (POWER), [7268](#)

CORR= option
CUSTOM statement (POWER), [7233](#)
LOGISTIC statement (POWER), [7238](#)
ONECORR statement (POWER), [7250](#)
PAIREFREQ statement (POWER), [7273](#)
PAIREDMEANS statement (POWER), [7281](#)

COVARIATES= option
LOGISTIC statement (POWER), [7238](#)

COVODDSRATIOS= option
LOGISTIC statement (POWER), [7239](#)

COVREGCOEFFS= option
LOGISTIC statement (POWER), [7239](#)

COXREG statement
POWER procedure, [7227](#)

CRITMULT= option
CUSTOM statement (POWER), [7233](#)

CURVE= option
TWOSAMPLESURVIVAL statement (POWER),
[7312](#)

CUSTOM statement
POWER procedure, [7231](#)

CV= option
ONESAMPLEMEANS statement (POWER),
[7263](#)
PAIREDMEANS statement (POWER), [7281](#)
TWOSAMPLEMEANS statement (POWER),
[7302](#)

DEFAULTNBINS= option
LOGISTIC statement (POWER), [7239](#)

DEFAULTUNIT= option
LOGISTIC statement (POWER), [7239](#)

DESCRIPTION= option
PLOT statement (POWER), [7292](#)

DISCPROPDIFF= option
PAIREFREQ statement (POWER), [7274](#)

DISCPROPORTIONS= option
PAIREFREQ statement (POWER), [7274](#)

DISCPRORATIO= option
PAIREFREQ statement (POWER), [7274](#)

DIST= option
CUSTOM statement (POWER), [7233](#)
ONECORR statement (POWER), [7251](#)
ONESAMPLEMEANS statement (POWER),
[7263](#)
PAIREFREQ statement (POWER), [7274](#)
PAIREDMEANS statement (POWER), [7281](#)
TWOSAMPLEMEANS statement (POWER),
[7302](#)

EQUIVBOUNDS= option
ONESAMPLEFREQ statement (POWER), [7255](#)

EVENTPROB= option
COXREG statement (POWER), [7229](#)

EVENTSTOTAL= option
COXREG statement (POWER), [7229](#)
TWOSAMPLESURVIVAL statement (POWER),
[7313](#)

FOLLOWUPTIME= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7313

GROUP ACCRUAL RATES= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7314

GROUP LOSS= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7314

GROUP LOSS EX PHAZARDS= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7314

GROUP MEANS= option
 ONE WAY ANOVA statement (POWER), 7269
 TWO SAMPLE MEANS statement (POWER),
 7303

GROUP MED LOSS TIMES= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7314

GROUP MED SURV TIMES= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7314

GROUP NS= option
 ONE WAY ANOVA statement (POWER), 7269
 TWO SAMPLE FREQ statement (POWER), 7294
 TWO SAMPLE MEANS statement (POWER),
 7303
 TWO SAMPLE SURVIVAL statement (POWER),
 7315
 TWO SAMPLE WILCOXON statement
 (Power), 7323

GROUP PROPORTIONS= option
 TWO SAMPLE FREQ statement (POWER), 7294

GROUP STD DEVS= option
 TWO SAMPLE MEANS statement (POWER),
 7303

GROUP SURV EX PHAZARDS= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7315

GROUP SURVIVAL= option
 TWO SAMPLE SURVIVAL statement (POWER),
 7315

GROUP WEIGHTS= option
 ONE WAY ANOVA statement (POWER), 7269
 TWO SAMPLE FREQ statement (POWER), 7294
 TWO SAMPLE MEANS statement (POWER),
 7303
 TWO SAMPLE SURVIVAL statement (POWER),
 7315
 TWO SAMPLE WILCOXON statement
 (Power), 7323

HALFWIDTH= option

ONE SAMPLE FREQ statement (POWER), 7255
 ONE SAMPLE MEANS statement (POWER),
 7263
 PAIRED MEANS statement (POWER), 7281
 TWO SAMPLE MEANS statement (POWER),
 7303

HAZARD RATIO= option
 COX REG statement (POWER), 7229
 TWO SAMPLE SURVIVAL statement (POWER),
 7315

INTERCEPT= option
 LOGISTIC statement (POWER), 7240

INTERPOL= option
 PLOT statement (POWER), 7288

KEY= option
 PLOT statement (POWER), 7288

LOGISTIC statement
 POWER procedure, 7237

LOWER= option
 ONE SAMPLE FREQ statement (POWER), 7256
 ONE SAMPLE MEANS statement (POWER),
 7263
 PAIRED MEANS statement (POWER), 7282
 TWO SAMPLE MEANS statement (POWER),
 7303

MARGIN= option
 ONE SAMPLE FREQ statement (POWER), 7256

MARKERS= option
 PLOT statement (POWER), 7289

MAX= option
 PLOT statement (POWER), 7289

MEAN= option
 ONE SAMPLE MEANS statement (POWER),
 7263

MEAN DIFF= option
 PAIRED MEANS statement (POWER), 7282
 TWO SAMPLE MEANS statement (POWER),
 7304

MEAN RATIO= option
 PAIRED MEANS statement (POWER), 7282
 TWO SAMPLE MEANS statement (POWER),
 7304

METHOD= option
 ONE SAMPLE FREQ statement (POWER), 7256
 PAIRED FREQ statement (POWER), 7274

MIN= option
 PLOT statement (POWER), 7289

MODEL= option
 MULT REG statement (POWER), 7246
 ONE CORR statement (POWER), 7251

MODEL DF= option

CUSTOM statement (POWER), [7233](#)
 MULTREG statement
 POWER procedure, [7245](#)
 NAME= option
 PLOT statement (POWER), [7292](#)
 NBINS= option
 LOGISTIC statement (POWER), [7240](#)
 TWO SAMPLE WILCOXON statement
 (POWER), [7323](#)
 NFRACTIONAL option
 COXREG statement (POWER), [7229](#)
 CUSTOM statement (POWER), [7234](#)
 LOGISTIC statement (POWER), [7240](#)
 MULTREG statement (POWER), [7246](#)
 ONECORR statement (POWER), [7251](#)
 ONESAMPLEMEANS statement (POWER),
 [7263](#)
 ONEWAYANOVA statement (POWER), [7269](#)
 PAIREDFREQ statement (POWER), [7274](#)
 PAIREDMEANS statement (POWER), [7282](#)
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 TWO SAMPLE MEANS statement (POWER),
 [7304](#)
 TWO SAMPLE SURVIVAL statement (POWER),
 [7315](#)
 NFRACTIONAL= option
 ONESAMPLEFREQ statement (POWER), [7256](#)
 TWO SAMPLE WILCOXON statement
 (POWER), [7323](#)
 NFULLPREDICTORS= option
 MULTREG statement (POWER), [7246](#)
 NOINT option
 MULTREG statement (POWER), [7246](#)
 NPAIRS= option
 PAIREDFREQ statement (POWER), [7274](#)
 PAIREDMEANS statement (POWER), [7282](#)
 NPARTIALVARS= option
 ONECORR statement (POWER), [7251](#)
 NPERGROUP= option
 ONEWAYANOVA statement (POWER), [7269](#)
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 TWO SAMPLE MEANS statement (POWER),
 [7304](#)
 TWO SAMPLE SURVIVAL statement (POWER),
 [7316](#)
 TWO SAMPLE WILCOXON statement
 (POWER), [7323](#)
 NPOINTS= option
 PLOT statement (POWER), [7290](#)
 NREDUCEDPREDICTORS= option
 MULTREG statement (POWER), [7247](#)
 NSUBINTERVAL= option

TWO SAMPLE SURVIVAL statement (POWER),
 [7316](#)
 NTESTPREDICTORS= option
 MULTREG statement (POWER), [7247](#)
 NTOTAL= option
 COXREG statement (POWER), [7229](#)
 CUSTOM statement (POWER), [7234](#)
 LOGISTIC statement (POWER), [7240](#)
 MULTREG statement (POWER), [7247](#)
 ONECORR statement (POWER), [7251](#)
 ONESAMPLEFREQ statement (POWER), [7256](#)
 ONESAMPLEMEANS statement (POWER),
 [7263](#)
 ONEWAYANOVA statement (POWER), [7269](#)
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 TWO SAMPLE MEANS statement (POWER),
 [7304](#)
 TWO SAMPLE SURVIVAL statement (POWER),
 [7316](#)
 TWO SAMPLE WILCOXON statement
 (POWER), [7323](#)
 NULLCONTRAST= option
 ONEWAYANOVA statement (POWER), [7269](#)
 NULLCORR= option
 ONECORR statement (POWER), [7251](#)
 NULLDIFF= option
 PAIREDMEANS statement (POWER), [7282](#)
 TWO SAMPLE MEANS statement (POWER),
 [7304](#)
 NULLDISCPRORATIO= option
 PAIREDFREQ statement (POWER), [7274](#)
 NULLMEAN= option
 ONESAMPLEMEANS statement (POWER),
 [7263](#)
 NULLODDSRATIO= option
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 NULLPROPORTION= option
 ONESAMPLEFREQ statement (POWER), [7256](#)
 NULLPROPORTIONDIFF= option
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 NULLRATIO= option
 PAIREDMEANS statement (POWER), [7282](#)
 TWO SAMPLE MEANS statement (POWER),
 [7304](#)
 NULLRELATIVERISK= option
 TWO SAMPLE FREQ statement (POWER), [7295](#)
 ODDSRATIO= option
 PAIREDFREQ statement (POWER), [7275](#)
 TWO SAMPLE FREQ statement (POWER), [7296](#)
 ONECORR statement
 POWER procedure, [7249](#)
 ONESAMPLEFREQ statement
 POWER procedure, [7253](#)

- ONESAMPLEMEANS statement
 - POWER procedure, 7261
- ONEWAYANOVA statement
 - POWER procedure, 7267
- OUTPUTORDER= option
 - COXREG statement (POWER), 7229
 - CUSTOM statement (POWER), 7234
 - LOGISTIC statement (POWER), 7240
 - MULTREG statement (POWER), 7247
 - ONECORR statement (POWER), 7251
 - ONESAMPLEFREQ statement (POWER), 7257
 - ONESAMPLEMEANS statement (POWER), 7263
 - ONEWAYANOVA statement (POWER), 7270
 - PAIREDFREQ statement (POWER), 7275
 - PAIREDMEANS statement (POWER), 7282
 - TWOSAMPLEFREQ statement (POWER), 7296
 - TWOSAMPLEMEANS statement (POWER), 7305
 - TWOSAMPLESURVIVAL statement (POWER), 7316
 - TWOSAMPLEWILCOXON statement (POWER), 7324
- PAIREDCVS= option
 - PAIREDMEANS statement (POWER), 7283
- PAIREDFREQ statement
 - POWER procedure, 7272
- PAIREDMEANS statement
 - POWER procedure, 7279
- PAIREDMEANS= option
 - PAIREDMEANS statement (POWER), 7283
- PAIREDPROPORTIONS= option
 - PAIREDFREQ statement (POWER), 7275
- PAIREDSTDDEVS= option
 - PAIREDMEANS statement (POWER), 7283
- PARTIALCORR= option
 - MULTREG statement (POWER), 7248
- PLOT statement
 - POWER procedure, 7288
- PLOTONLY= option
 - PROC POWER statement, 7227
- PNCMULT= option
 - CUSTOM statement (POWER), 7234
- POWER procedure
 - syntax, 7226
- POWER procedure, COXREG statement, 7227
 - ALPHA= option, 7228
 - EVENTPROB= option, 7229
 - EVENTSTOTAL= option, 7229
 - HAZARDRATIO= option, 7229
 - NFRACTIONAL option, 7229
 - NTOTAL= option, 7229
 - OUTPUTORDER= option, 7229
 - POWER= option, 7230
 - RSQUARE= option, 7230
 - SIDES= option, 7230
 - STDDEV= option, 7230
 - TEST= option, 7230
- POWER procedure, CUSTOM statement, 7231
 - ALPHA= option, 7233
 - CORR= option, 7233
 - CRITMULT= option, 7233
 - DIST= option, 7233
 - MODELDF= option, 7233
 - NFRACTIONAL option, 7234
 - NTOTAL= option, 7234
 - OUTPUTORDER= option, 7234
 - PNCMULT= option, 7234
 - POWER= option, 7234
 - PRIMNC= option, 7235
 - SIDES= option, 7235
 - TESTDF= option, 7235
- POWER procedure, LOGISTIC statement, 7237
 - ALPHA= option, 7238
 - CORR= option, 7238
 - COVARIATES= option, 7238
 - COVODDSRATIOS= option, 7239
 - COVREGCOEFFS= option, 7239
 - DEFAULTNBINS= option, 7239
 - DEFAULTUNIT= option, 7239
 - INTERCEPT= option, 7240
 - NBINS= option, 7240
 - NFRACTIONAL option, 7240
 - NTOTAL= option, 7240
 - OUTPUTORDER= option, 7240
 - POWER= option, 7241
 - RESPONSEPROB= option, 7241
 - TEST= option, 7241
 - TESTODDSRATIO= option, 7241
 - TESTPREDICTOR= option, 7241
 - TESTREGCOEFF= option, 7241
 - UNITS= option, 7242
 - VARDIST= option, 7242
- POWER procedure, MULTREG statement, 7245
 - ALPHA= option, 7246
 - MODEL= option, 7246
 - NFRACTIONAL option, 7246
 - NFULLPREDICTORS= option, 7246
 - NOINT option, 7246
 - NREDUCEDPREDICTORS= option, 7247
 - NTESTPREDICTORS= option, 7247
 - NTOTAL= option, 7247
 - OUTPUTORDER= option, 7247
 - PARTIALCORR= option, 7248
 - POWER= option, 7248
 - RSQUAREDIFF= option, 7248
 - RSQUAREFULL= option, 7248

RSQUAREREDUCED= option, 7248
 TEST= option, 7248
 POWER procedure, ONEWAYANOVA statement, 7249
 ALPHA= option, 7250
 CORR= option, 7250
 DIST= option, 7251
 MODEL= option, 7251
 NFRACTIONAL option, 7251
 NPARTIALVARS= option, 7251
 NTOTAL= option, 7251
 NULLCORR= option, 7251
 OUTPUTORDER= option, 7251
 POWER= option, 7252
 SIDES= option, 7252
 TEST= option, 7252
 POWER procedure, ONESAMPLEFREQ statement, 7253
 ALPHA= option, 7255
 CI= option, 7255
 EQUIVBOUNDS= option, 7255
 HALFWIDTH= option, 7255
 LOWER= option, 7256
 MARGIN= option, 7256
 METHOD= option, 7256
 NFRACTIONAL= option, 7256
 NTOTAL= option, 7256
 NULLPROPORTION= option, 7256
 OUTPUTORDER= option, 7257
 POWER= option, 7257
 PROBWIDTH= option, 7257
 PROPORTION= option, 7257
 SIDES= option, 7257
 TEST= option, 7258
 UPPER= option, 7258
 VAREST= option, 7258
 POWER procedure, ONESAMPLEMEANS statement, 7261
 ALPHA= option, 7262
 CI= option, 7262
 CV= option, 7263
 DIST= option, 7263
 HALFWIDTH= option, 7263
 LOWER= option, 7263
 MEAN= option, 7263
 NFRACTIONAL option, 7263
 NTOTAL= option, 7263
 NULLMEAN= option, 7263
 OUTPUTORDER= option, 7263
 POWER= option, 7264
 PROBTYPE= option, 7264
 PROBWIDTH= option, 7264
 SIDES= option, 7265
 STDDEV= option, 7265
 TEST= option, 7265

UPPER= option, 7265
 POWER procedure, ONEWAYANOVA statement, 7267
 ALPHA= option, 7268
 CONTRAST= option, 7268
 GROUPMEANS= option, 7269
 GROUPNS= option, 7269
 GROUPWEIGHTS= option, 7269
 NFRACTIONAL option, 7269
 NPERGROUP= option, 7269
 NTOTAL= option, 7269
 NULLCONTRAST= option, 7269
 OUTPUTORDER= option, 7270
 POWER= option, 7270
 SIDES= option, 7270
 STDDEV= option, 7270
 TEST= option, 7271
 POWER procedure, PAIREDFREQ statement, 7272
 ALPHA= option, 7273
 CORR= option, 7273
 DISCPROPDIF= option, 7274
 DISCPROPORTIONS= option, 7274
 DISCPROPRATIO= option, 7274
 DIST= option, 7274
 METHOD= option, 7274
 NFRACTIONAL option, 7274
 NPAIRS= option, 7274
 NULLDISCPROPRATIO= option, 7274
 ODDSRATIO= option, 7275
 OUTPUTORDER= option, 7275
 PAIREDPROPORTIONS= option, 7275
 POWER= option, 7275
 PROPORTIONDIFF= option, 7275
 REFPROPORTION= option, 7276
 RELATIVERISK= option, 7276
 SIDES= option, 7276
 TEST= option, 7276
 TOTALPROPDISC= option, 7276
 POWER procedure, PAIREDMEANS statement, 7279
 ALPHA= option, 7281
 CI= option, 7281
 CORR= option, 7281
 CV= option, 7281
 DIST= option, 7281
 HALFWIDTH= option, 7281
 LOWER= option, 7282
 MEANDIFF= option, 7282
 MEANRATIO= option, 7282
 NFRACTIONAL option, 7282
 NPAIRS= option, 7282
 NULLDIFF= option, 7282
 NULLRATIO= option, 7282
 OUTPUTORDER= option, 7282
 PAIREDCVS= option, 7283

PAIREDMEANS= option, 7283
 PAIREDSTDDEVS= option, 7283
 POWER= option, 7284
 PROBTYP= option, 7284
 PROBWIDTH= option, 7284
 SIDES= option, 7284
 STDDEV= option, 7285
 TEST= option, 7285
 UPPER= option, 7285
 POWER procedure, PLOT statement, 7288
 DESCRIPTION= option, 7292
 INTERPOL= option, 7288
 KEY= option, 7288
 MARKERS= option, 7289
 MAX= option, 7289
 MIN= option, 7289
 NAME= option, 7292
 NPOINTS= option, 7290
 STEP= option, 7290
 VARY option, 7290
 X= option, 7290
 XOPTS= option, 7291
 Y= option, 7292
 YOPTS= option, 7292
 POWER procedure, PROC POWER statement, 7227
 PLOTONLY= option, 7227
 POWER procedure, TWOSAMPLEFREQ statement, 7292
 ALPHA= option, 7294
 GROUPNS= option, 7294
 GROUPPROPORTIONS= option, 7294
 GROUPWEIGHTS= option, 7294
 NFRACTIONAL option, 7295
 NPERGROUP= option, 7295
 NTOTAL= option, 7295
 NULLODDSRATIO= option, 7295
 NULLPROPORTIONDIFF= option, 7295
 NULLRELATIVERISK= option, 7295
 ODDSRATIO= option, 7296
 OUTPUTORDER= option, 7296
 POWER= option, 7296
 PROPORTIONDIFF= option, 7296
 REFPROPORTION= option, 7296
 RELATIVERISK= option, 7297
 SIDES= option, 7297
 TEST= option, 7297
 POWER procedure, TWOSAMPLEMEANS statement, 7300
 ALPHA= option, 7302
 CI= option, 7302
 CV= option, 7302
 DIST= option, 7302
 GROUPMEANS= option, 7303
 GROUPNS= option, 7303
 GROUPSTDDEVS= option, 7303
 GROUPWEIGHTS= option, 7303
 HALFWIDTH= option, 7303
 LOWER= option, 7303
 MEANDIFF= option, 7304
 MEANRATIO= option, 7304
 NFRACTIONAL option, 7304
 NPERGROUP= option, 7304
 NTOTAL= option, 7304
 NULLDIFF= option, 7304
 NULLRATIO= option, 7304
 OUTPUTORDER= option, 7305
 POWER= option, 7305
 PROBTYP= option, 7305
 PROBWIDTH= option, 7306
 SIDES= option, 7306
 STDDEV= option, 7306
 TEST= option, 7306
 UPPER= option, 7307
 POWER procedure, TWOSAMPLESURVIVAL statement, 7310
 ACCRUALRATEPERGROUP= option, 7311
 ACCRUALRATETOTAL= option, 7312
 ACCRUALTIME= option, 7312
 ALPHA= option, 7312
 CURVE= option, 7312
 EVENTSTOTAL= option, 7313
 FOLLOWUPTIME= option, 7313
 GROUPACCRUALRATES= option, 7314
 GROUPLOSS= option, 7314
 GROUPLOSSEXPHAZARDS= option, 7314
 GROUPMEDLOSSTIMES= option, 7314
 GROUPMEDSURVTIMES= option, 7314
 GROUPNS= option, 7315
 GROUPSURVEXPHAZARDS= option, 7315
 GROUPSURVIVAL= option, 7315
 GROUPWEIGHTS= option, 7315
 HAZARDRATIO= option, 7315
 NFRACTIONAL option, 7315
 NPERGROUP= option, 7316
 NSUBINTERVAL= option, 7316
 NTOTAL= option, 7316
 OUTPUTORDER= option, 7316
 POWER= option, 7317
 REFSURVEXPHAZARD= option, 7317
 REFSURVIVAL= option, 7317
 SIDES= option, 7317
 TEST= option, 7317
 TOTALTIME= option, 7317
 POWER procedure, TWOSAMPLEWILCOXON statement, 7321
 ALPHA= option, 7323
 GROUPNS= option, 7323
 GROUPWEIGHTS= option, 7323

NBINS= option, [7323](#)
 NFRACTIONAL= option, [7323](#)
 NPERGROUP= option, [7323](#)
 NTOTAL= option, [7323](#)
 OUTPUTORDER= option, [7324](#)
 POWER= option, [7324](#)
 SIDES= option, [7324](#)
 TEST= option, [7324](#)
 VARDIST= option, [7324](#)
 VARIABLES= option, [7325](#)
 POWER= option
 COXREG statement (POWER), [7230](#)
 CUSTOM statement (POWER), [7234](#)
 LOGISTIC statement (POWER), [7241](#)
 MULTREG statement (POWER), [7248](#)
 ONECORR statement (POWER), [7252](#)
 ONESAMPLEFREQ statement (POWER), [7257](#)
 ONESAMPLEMEANS statement (POWER),
 [7264](#)
 ONEWAYANOVA statement (POWER), [7270](#)
 PAIREDFREQ statement (POWER), [7275](#)
 PAIREDMEANS statement (POWER), [7284](#)
 TWOSAMPLEFREQ statement (POWER), [7296](#)
 TWOSAMPLEMEANS statement (POWER),
 [7305](#)
 TWOSAMPLESURVIVAL statement (POWER),
 [7317](#)
 TWOSAMPLEWILCOXON statement
 (POWER), [7324](#)
 PRIMNC= option
 CUSTOM statement (POWER), [7235](#)
 PROBTYP= option
 ONESAMPLEMEANS statement (POWER),
 [7264](#)
 PAIREDMEANS statement (POWER), [7284](#)
 TWOSAMPLEMEANS statement (POWER),
 [7305](#)
 PROBWIDTH= option
 ONESAMPLEFREQ statement (POWER), [7257](#)
 ONESAMPLEMEANS statement (POWER),
 [7264](#)
 PAIREDMEANS statement (POWER), [7284](#)
 TWOSAMPLEMEANS statement (POWER),
 [7306](#)
 PROC POWER statement, *see* POWER procedure
 PROPORTION= option
 ONESAMPLEFREQ statement (POWER), [7257](#)
 PROPORTIONDIFF= option
 PAIREDFREQ statement (POWER), [7275](#)
 TWOSAMPLEFREQ statement (POWER), [7296](#)
 REFPROPORTION= option
 PAIREDFREQ statement (POWER), [7276](#)
 TWOSAMPLEFREQ statement (POWER), [7296](#)
 REFSURVEXPHAZARD= option
 TWOSAMPLESURVIVAL statement (POWER),
 [7317](#)
 REFSURVIVAL= option
 TWOSAMPLESURVIVAL statement (POWER),
 [7317](#)
 RELATIVERISK= option
 PAIREDFREQ statement (POWER), [7276](#)
 TWOSAMPLEFREQ statement (POWER), [7297](#)
 RESPONSEPROB= option
 LOGISTIC statement (POWER), [7241](#)
 RSQUARE= option
 COXREG statement (POWER), [7230](#)
 RSQUAREDIFF= option
 MULTREG statement (POWER), [7248](#)
 RSQUAREFULL= option
 MULTREG statement (POWER), [7248](#)
 RSQUAREREDUCED= option
 MULTREG statement (POWER), [7248](#)
 SIDES= option
 COXREG statement (POWER), [7230](#)
 CUSTOM statement (POWER), [7235](#)
 ONECORR statement (POWER), [7252](#)
 ONESAMPLEFREQ statement (POWER), [7257](#)
 ONESAMPLEMEANS statement (POWER),
 [7265](#)
 ONEWAYANOVA statement (POWER), [7270](#)
 PAIREDFREQ statement (POWER), [7276](#)
 PAIREDMEANS statement (POWER), [7284](#)
 TWOSAMPLEFREQ statement (POWER), [7297](#)
 TWOSAMPLEMEANS statement (POWER),
 [7306](#)
 TWOSAMPLESURVIVAL statement (POWER),
 [7317](#)
 TWOSAMPLEWILCOXON statement
 (POWER), [7324](#)
 STDDEV= option
 COXREG statement (POWER), [7230](#)
 ONESAMPLEMEANS statement (POWER),
 [7265](#)
 ONEWAYANOVA statement (POWER), [7270](#)
 PAIREDMEANS statement (POWER), [7285](#)
 TWOSAMPLEMEANS statement (POWER),
 [7306](#)
 STEP= option
 PLOT statement (POWER), [7290](#)
 TEST= option
 COXREG statement (POWER), [7230](#)
 LOGISTIC statement (POWER), [7241](#)
 MULTREG statement (POWER), [7248](#)
 ONECORR statement (POWER), [7252](#)
 ONESAMPLEFREQ statement (POWER), [7258](#)

ONESAMPLEMEANS statement (POWER),
 7265
 ONEWAYANOVA statement (POWER), 7271
 PAIREDFREQ statement (POWER), 7276
 PAIREDMEANS statement (POWER), 7285
 TWOSAMPLEFREQ statement (POWER), 7297
 TWOSAMPLEMEANS statement (POWER),
 7306
 TWOSAMPLESURVIVAL statement (POWER),
 7317
 TWOSAMPLEWILCOXON statement
 (POWER), 7324
 TESTDF= option
 CUSTOM statement (POWER), 7235
 TESTODDSRATIO= option
 LOGISTIC statement (POWER), 7241
 TESTPREDICTOR= option
 LOGISTIC statement (POWER), 7241
 TESTREGCOEFF= option
 LOGISTIC statement (POWER), 7241
 TOTALPROPDISC= option
 PAIREDFREQ statement (POWER), 7276
 TOTALTIME= option
 TWOSAMPLESURVIVAL statement (POWER),
 7317
 TWOSAMPLEFREQ statement
 POWER procedure, 7292
 TWOSAMPLEMEANS statement
 POWER procedure, 7300
 TWOSAMPLESURVIVAL statement
 POWER procedure, 7310
 TWOSAMPLEWILCOXON statement
 POWER procedure, 7321

 UNITS= option
 LOGISTIC statement (POWER), 7242
 UPPER= option
 ONESAMPLEFREQ statement (POWER), 7258
 ONESAMPLEMEANS statement (POWER),
 7265
 PAIREDMEANS statement (POWER), 7285
 TWOSAMPLEMEANS statement (POWER),
 7307

 VARDIST= option
 LOGISTIC statement (POWER), 7242
 TWOSAMPLEWILCOXON statement
 (POWER), 7324
 VAREST= option
 ONESAMPLEFREQ statement (POWER), 7258
 VARIABLES= option
 TWOSAMPLEWILCOXON statement
 (POWER), 7325
 VARY option

 PLOT statement (POWER), 7290

 X= option
 PLOT statement (POWER), 7290
 XOPTS= option
 PLOT statement (POWER), 7291

 Y= option
 PLOT statement (POWER), 7292
 YOPTS= option
 PLOT statement (POWER), 7292