

SAS/STAT[®] 14.1 User's Guide

The GLMPOWER Procedure

This document is an individual chapter from *SAS/STAT® 14.1 User's Guide*.

The correct bibliographic citation for this manual is as follows: SAS Institute Inc. 2015. *SAS/STAT® 14.1 User's Guide*. Cary, NC: SAS Institute Inc.

SAS/STAT® 14.1 User's Guide

Copyright © 2015, 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

July 2015

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.

Chapter 48

The GLMPOWER Procedure

Contents

Overview: GLMPOWER Procedure	3738
Getting Started: GLMPOWER Procedure	3739
Simple Two-Way ANOVA	3739
Incorporating Contrasts, Unbalanced Designs, and Multiple Means Scenarios	3743
Syntax: GLMPOWER Procedure	3745
PROC GLMPOWER Statement	3746
BY Statement	3747
CLASS Statement	3748
CONTRAST Statement	3748
MANOVA Statement	3749
MODEL Statement	3751
PLOT Statement	3751
POWER Statement	3756
REPEATED Statement	3763
WEIGHT Statement	3766
Details: GLMPOWER Procedure	3766
Specifying Value Lists in the POWER Statement	3766
Number-Lists	3767
Name-Lists	3767
Keyword-Lists	3767
Sample Size Adjustment Options	3767
Error and Information Output	3768
Displayed Output	3769
ODS Table Names	3770
Computational Methods and Formulas	3770
Contrasts in Fixed-Effect Univariate Models	3770
Adjustments for Covariates in Univariate Models	3772
Contrasts in Fixed-Effect Multivariate Models	3773
ODS Graphics	3781
Examples: GLMPOWER Procedure	3781
Example 48.1: One-Way ANOVA	3781
Example 48.2: Two-Way ANOVA with Covariate	3788
Example 48.3: Repeated Measures ANOVA	3794
References	3798

Overview: GLMPOWER Procedure

Power and sample size analysis optimizes the resource usage and design of a study, improving chances of conclusive results with maximum efficiency. The GLMPOWER procedure performs prospective power and sample size analysis for linear models, with a variety of goals:

- 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* analysis for a past study, which is not supported by this procedure.

The statistical analyses that are covered include Type III F tests and contrasts of fixed effects in univariate and multivariate linear models. For univariate models, you can specify covariates, which can be continuous or categorical. For multivariate models, you can choose among Wilks' likelihood ratio, Hotelling-Lawley trace, and Pillai's trace F tests for multivariate analysis of variance (MANOVA) and among uncorrected, Greenhouse-Geisser, Huynh-Feldt, and Box conservative F tests for the univariate approach to repeated measures. Tests and contrasts that involve random effects are not supported. For power and sample size analyses in a variety of other statistical situations, see Chapter 89, "[The POWER Procedure](#)."

Input for PROC GLMPOWER includes the following components, which are considered in study planning:

- design (including subject profiles and their allocation weights)
- statistical model and test
- between-subject contrasts of class effects
- within-subject contrasts (for multivariate models)
- significance level (alpha)
- surmised response means for subject profiles (often called "cell means")
- surmised variability (and correlation for multivariate models)
- power
- sample size

In order to identify power or sample size as the result parameter, you designate it by a missing value in the input. The procedure calculates this result value over one or more scenarios of input values for all other components.

You specify the design and the cell means by using an *exemplary data set*, a data set of artificial values that is constructed to represent the intended sampling design and the surmised response means in the underlying population. You specify the model and between-subject contrasts by using **MODEL** and **CONTRAST** statements similar to those in the GLM, ANOVA, and MIXED procedures. For multivariate models, you

specify the within-subject contrasts by using [MANOVA](#) and [REPEATED](#) statements similar to those in the GLM and MIXED procedures. You specify the remaining parameters by using the [POWER](#) statement, which is similar to analysis statements in the POWER procedure.

In addition to tabular results, PROC GLMPOWER 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 GLMPOWER 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 609 in Chapter 21, “[Statistical Graphics Using ODS](#).”

For specific information about the statistical graphics and options available with the GLMPOWER procedure, see the [PLOT](#) statement and the section “[ODS Graphics](#)” on page 3781.

The GLMPOWER procedure is one of several tools available in SAS/STAT software for power and sample size analysis. PROC POWER covers a variety of other analyses such as *t* tests, equivalence tests, confidence intervals, binomial proportions, multiple regression, one-way ANOVA, survival analysis, logistic regression, and the Wilcoxon rank-sum test. The Power and Sample Size application provides a user interface and implements many of the analyses supported in the procedures. For more information, see Chapter 89, “[The POWER Procedure](#),” and Chapter 90, “[The Power and Sample Size Application](#).”

The following sections of this chapter describe how to use PROC GLMPOWER and discuss the underlying statistical methodology. The section “[Getting Started: GLMPOWER Procedure](#)” on page 3739 introduces PROC GLMPOWER with examples of power computation for a two-way analysis of variance. The section “[Syntax: GLMPOWER Procedure](#)” on page 3745 describes the syntax of the procedure. The section “[Details: GLMPOWER Procedure](#)” on page 3766 summarizes the methods employed by PROC GLMPOWER and provides details on several special topics. The section “[Examples: GLMPOWER Procedure](#)” on page 3781 illustrates the use of the GLMPOWER 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 for linear models, see Castelleo and O’Brien (2001); O’Brien and Shieh (1992); Muller and Benignus (1992); O’Brien and Muller (1993). For additional discussion of general power and sample size concepts, see O’Brien and Castelleo (2007); Castelleo (2000); Muller and Benignus (1992); Lenth (2001).

Getting Started: GLMPOWER Procedure

Simple Two-Way ANOVA

This example demonstrates how to use PROC GLMPOWER to compute and plot power for each effect test in a two-way analysis of variance (ANOVA).

Suppose you are planning an experiment to study the effect of light exposure at three levels on the growth of two varieties of flowers. The planned data analysis is a two-way ANOVA with flower height (measured

at two weeks) as the response and a model consisting of the effects of light exposure, flower variety, and their interaction. You want to calculate the power of each effect test for a balanced design with a total of 60 specimens (10 for each combination of exposure and variety) with $\alpha = 0.05$ for each test.

As a first step, create an *exemplary data set* describing your conjectures about the underlying population means. You believe that the mean flower height for each combination of variety and exposure level (that is, for each design profile, or for each *cell* in the design) roughly follows [Table 48.1](#).

Table 48.1 Mean Flower Height (in cm) by Variety and Exposure

Variety	Exposure		
	1	2	3
1	14	16	21
2	10	15	16

The following statements create a data set named Exemplary containing these cell means.

```
data Exemplary;
  do Variety = 1 to 2;
    do Exposure = 1 to 3;
      input Height @@;
      output;
    end;
  end;
  datalines;
    14 16 21
    10 15 16
;
```

You also conjecture that the error standard deviation is about 5 cm.

Use the **DATA=** option in the **PROC GLMPOWER** statement to specify Exemplary as the exemplary data set. Identify the classification variables (Variety and Exposure) by using the **CLASS** statement. Specify the model by using the **MODEL** statement. Use the **POWER** statement to specify power as the result parameter and provide values for the other analysis parameters, error standard deviation and total sample size. The following SAS statements perform the power analysis:

```
proc glmpower data=Exemplary;
  class Variety Exposure;
  model Height = Variety | Exposure;
  power
    stddev = 5
    ntotal = 60
    power = .;
run;
```

The **MODEL** statement defines the full model including both main effects and the interaction. The **POWER=** option in the **POWER** statement identifies power as the result parameter with a missing value (**POWER=.**). The **STDDEV=** option specifies an error standard deviation of 5, and the **NTOTAL=** option specifies a total sample size of 60. The default value for the **ALPHA=** option sets the significance level to $\alpha = 0.05$.

Figure 48.1 shows the output.

Figure 48.1 Sample Size Analysis for Two-Way ANOVA

The GLMPower Procedure

Fixed Scenario Elements			
Dependent Variable		Height	
Error Standard Deviation		5	
Total Sample Size		60	
Alpha		0.05	
Error Degrees of Freedom		54	

Computed Power			
Index	Source	Test DF	Power
1	Variety	1	0.718
2	Exposure	2	0.957
3	Variety*Exposure	2	0.191

The power is about 0.72 for the test of the Variety effect. In other words, there is a probability of 0.72 that the test of the Variety effect will produce a significant result (given the assumptions for the means and error standard deviation). The power is 0.96 for the test of the Exposure effect and 0.19 for the interaction test.

Now, suppose you want to account for some of your uncertainty in conjecturing the true error standard deviation by evaluating the power at reasonable low and high values, 4 and 6.5. You also want to plot power for sample sizes between 30 and 90. The following statements perform the analysis:

```
ods graphics on;

proc glmpower data=Exemplary;
  class Variety Exposure;
  model Height = Variety | Exposure;
  power
    stddev = 4 6.5
    ntotal = 60
    power = .;
  plot x=n min=30 max=90;
run;

ods graphics off;
```

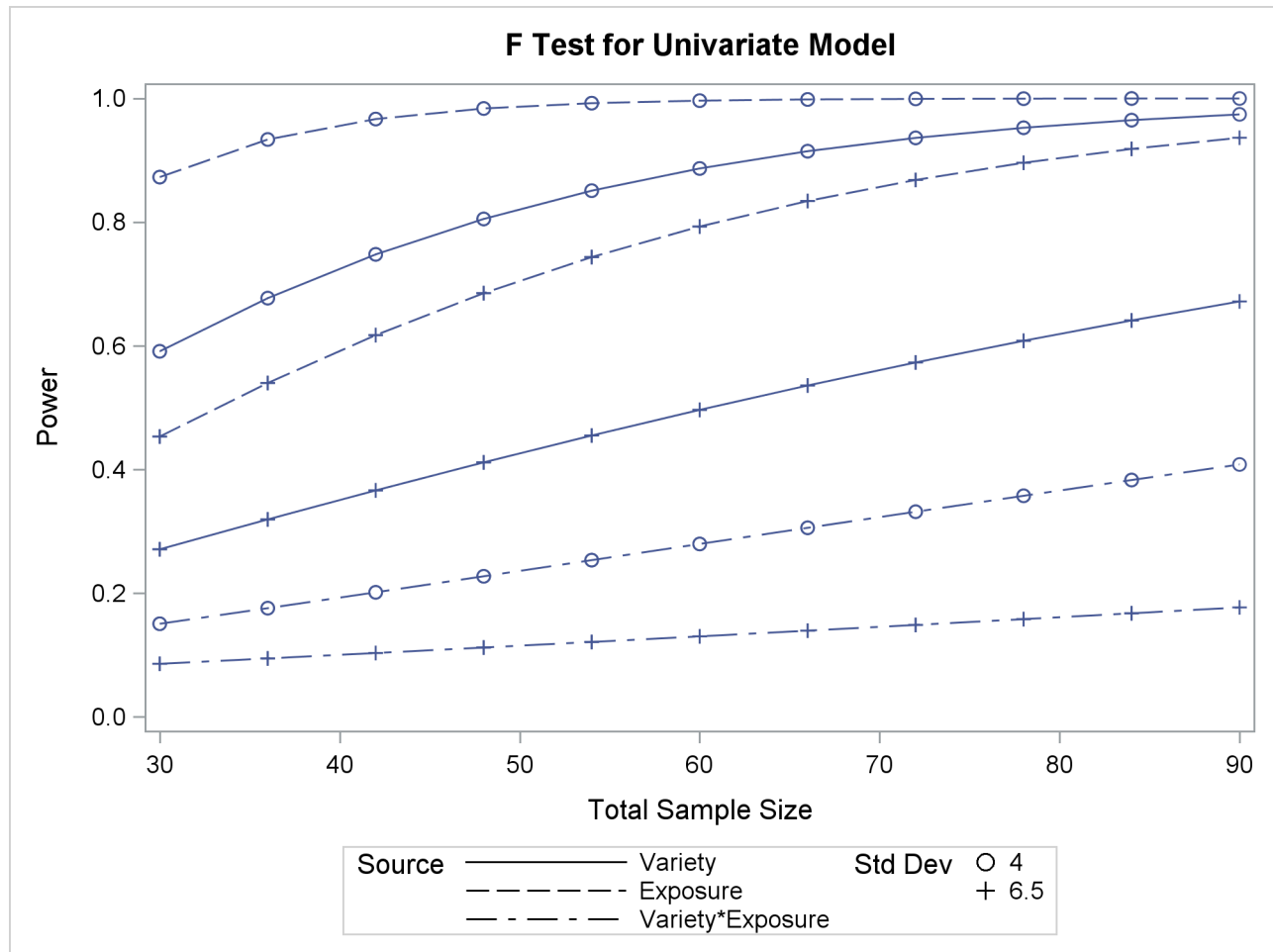
The **PLOT** statement with the **X=N** option requests a plot with sample size on the X axis. (The result parameter—in this case, power—is always plotted on the other axis.) The **MIN=** and **MAX=** options in the **PLOT** statement specify the sample size range. The ODS GRAPHICS ON statement enables ODS Graphics.

Figure 48.2 shows the output, and Figure 48.3 shows the plot.

Figure 48.2 Sample Size Analysis for Two-Way ANOVA with Input Ranges

The GLMPOWER Procedure				
Fixed Scenario Elements				
Dependent Variable		Height		
Total Sample Size		60		
Alpha		0.05		
Error Degrees of Freedom		54		
Computed Power				
Index	Source	Std Dev	Test DF	Power
1	Variety	4.0	1	0.887
2	Variety	6.5	1	0.496
3	Exposure	4.0	2	0.996
4	Exposure	6.5	2	0.793
5	Variety*Exposure	4.0	2	0.280
6	Variety*Exposure	6.5	2	0.130

Figure 48.2 reveals that the power ranges from about 0.130 to 0.996 for the different effect tests and scenarios for standard deviation, with a sample size of 60. In Figure 48.3, the line style identifies the effect test, and the plotting symbol identifies the standard deviation. The locations of the plotting symbols identify actual computed powers; the curves are linear interpolations of these points. Note that the computed points in the plot occur at sample size multiples of 6, because there are 6 cells in the design (and by default, sample sizes are rounded to produce integer cell sizes).

Figure 48.3 Plot of Power versus Sample Size for Two-Way ANOVA with Input Ranges

Incorporating Contrasts, Unbalanced Designs, and Multiple Means Scenarios

Suppose you want to compute power for the two-way ANOVA described in the section “Simple Two-Way ANOVA” on page 3739, but you want to additionally perform the following tasks:

- try an unbalanced sample size allocation with respect to Exposure, using twice as many samples for levels 2 and 3 as for level 1
- consider an additional, less optimistic scenario for the cell means, shown in [Table 48.2](#)
- test a contrast of Exposure comparing levels 1 and 3

Table 48.2 Additional Cell Means Scenario

Variety	Exposure		
	1	2	3
1	15	16	20
2	11	14	15

To specify the unbalanced design and the additional cell means scenario, you can add two new variables to the exemplary data set (Weight for the sample size weights, and HeightNew for the new cell means scenario). Change the name of the original cell means scenario to HeightOrig. The following statements define the exemplary data set:

```
data Exemplary;
  input Variety $ Exposure $ HeightOrig HeightNew Weight;
  datalines;
    1 1 14 15 1
    1 2 16 16 2
    1 3 21 20 2
    2 1 10 11 1
    2 2 15 14 2
    2 3 16 15 2
  ;
```

In PROC GLMPOWER, specify the name of the weight variable by using the **WEIGHT** statement, and specify the name of the cell means variables as dependent variables in the **MODEL** statement. Use the **CONTRAST** statement to specify the contrast as you would in PROC GLM. The following statements perform the sample size analysis.

```
proc glmpower data=Exemplary;
  class Variety Exposure;
  model HeightOrig HeightNew = Variety | Exposure;
  weight Weight;
  contrast 'Exposure=1 vs Exposure=3' Exposure 1 0 -1;
  power
    stddev = 5
    ntotal = 60
    power = .;
run;
```

Figure 48.4 shows the output.

Figure 48.4 Sample Size Analysis for More Complex Two-Way ANOVA

The GLMPOWER Procedure

Fixed Scenario Elements	
Weight Variable	Weight
Error Standard Deviation	5
Total Sample Size	60
Alpha	0.05
Error Degrees of Freedom	54

Figure 48.4 *continued*

Computed Power					Test	
Index	Dependent	Type	Source		DF	Power
1	HeightOrig	Effect	Variety		1	0.672
2	HeightOrig	Effect	Exposure		2	0.911
3	HeightOrig	Effect	Variety*Exposure		2	0.217
4	HeightOrig	Contrast	Exposure=1 vs Exposure=3		1	0.951
5	HeightNew	Effect	Variety		1	0.754
6	HeightNew	Effect	Exposure		2	0.633
7	HeightNew	Effect	Variety*Exposure		2	0.137
8	HeightNew	Contrast	Exposure=1 vs Exposure=3		1	0.705

The power of the contrast of Exposure levels 1 and 3 is about 0.95 for the original cell means scenario (HeightOrig) and only 0.71 for the new one (HeightNew). The power is higher for the test of Variety, but lower for the tests of Exposure and of Variety*Exposure for the new cell means scenario compared to the original one. Note also for the HeightOrig scenario that the power for the unbalanced design (Figure 48.4) compared to the balanced design (Figure 48.1) is slightly lower for the tests of Variety and Exposure, but slightly higher for the test of Variety*Exposure.

Syntax: GLMPOWER Procedure

The following statements are available in the GLMPOWER procedure:

```

PROC GLMPOWER < options > ;
  BY variables ;
  CLASS variables ;
  CONTRAST 'label' effect values < ... effect values > < / options > ;
  MANOVA 'label' < test-options > < / detail-options > ;
  MODEL dependent-variables = independent-effects ;
  PLOT < plot-options > < / graph-options > ;
  POWER < options > ;
  REPEATED factor-specification ;
  WEIGHT variable ;

```

The **PROC GLMPOWER** statement, the **MODEL** statement, and the **POWER** statement are required. If your model contains classification effects, the classification variables must be listed in a **CLASS** statement, and the **CLASS** statement must appear before the **MODEL** statement. In addition, **CONTRAST** and **POWER** statements must appear after the **MODEL** statement. **PLOT** statements must appear after the **POWER** statement that defines the analysis for the plot.

If you specify one or more **MANOVA** or **REPEATED** statements, then the model is assumed to be multivariate. Otherwise, a univariate model is assumed, in which case multiple dependent variables represent cell means scenarios for a single response.

You can use multiple **CONTRAST**, **MANOVA**, **REPEATED**, **POWER**, and **PLOT** statements. Each **CONTRAST** statement defines a separate between-subject contrast. Each **MANOVA** or **REPEATED** statement

defines a separate within-subject contrast for a multivariate model. Each **POWER** statement produces a separate analysis and uses the information that is contained in the **CLASS**, **MODEL**, **WEIGHT**, **CONTRAST**, **MANOVA**, and **REPEATED** statements. Each **PLOT** statement refers to the previous **POWER** statement and generates a separate graph (or set of graphs).

Table 48.3 summarizes the basic functions of each statement in PROC GLMPOWER. The syntax of each statement in Table 48.3 is described in the following pages.

Table 48.3 Statements in the GLMPOWER Procedure

Statement	Description
PROC GLMPOWER	Invokes procedure and specifies exemplary data set
BY	Specifies variables to define subgroups for the analysis
CLASS	Declares classification variables
CONTRAST	Defines between-subject linear tests of model parameters
MANOVA	Defines within-subject linear tests of model parameters for multivariate models, in terms of contrast matrix coefficients
MODEL	Defines model and specifies dependent variables; for univariate models, multiple dependent variables represent cell means scenarios for a single response
PLOT	Displays graphs for preceding POWER statement
POWER	Identifies parameter to solve for and provides one or more scenarios for values of other analysis parameters
REPEATED	Defines within-subject linear tests of model parameters for multivariate models, in terms of common repeated measures transformations of the dependent variables
WEIGHT	Specifies variable for allocating sample sizes to different subject profiles

PROC GLMPOWER Statement

PROC GLMPOWER < options> ;

The **PROC GLMPOWER** statement invokes the GLMPOWER procedure. You can specify the following options.

DATA=SAS-data-set

names a SAS data set to be used as the exemplary data set, which is an artificial data set constructed to represent the intended sampling design and the conjectured response means for the underlying population.

ORDER=DATA | FORMATTED | FREQ | INTERNAL

specifies the sort order for the levels of the classification variables (which are specified in the [CLASS](#) statement).

This option applies to the levels for all classification variables, except when you use the (default) `ORDER=FORMATTED` option with numeric classification variables that have no explicit format. In that case, the levels of such variables are ordered by their internal value.

The `ORDER=` option can take the following values:

Value of <code>ORDER=</code>	Levels Sorted By
DATA	Order of appearance in the input data set
FORMATTED	External formatted value, except for numeric variables with no explicit format, which are sorted by their unformatted (internal) value
FREQ	Descending frequency count; levels with the most observations come first in the order
INTERNAL	Unformatted value

By default, `ORDER=FORMATTED`. For `ORDER=FORMATTED` and `ORDER=INTERNAL`, the sort order is machine-dependent.

For more information about sort order, see the chapter on the `SORT` procedure in the *Base SAS Procedures Guide* and the discussion of BY-group processing in *SAS Language Reference: Concepts*.

PLOTONLY

specifies that only graphical results from the [PLOT](#) statement be produced.

BY Statement

BY variables ;

You can specify a BY statement with `PROC GLMPower` to obtain separate analyses of observations in groups that are defined by the BY variables. When a BY statement appears, the procedure expects the input data set to be sorted in order of the BY variables. If you specify more than one BY statement, only the last one specified is used.

If your input data set is not sorted in ascending order, use one of the following alternatives:

- Sort the data by using the `SORT` procedure with a similar BY statement.

- Specify the NOTSORTED or DESCENDING option in the BY statement for the GLMPOWER procedure. The NOTSORTED option does not mean that the data are unsorted but rather that the data are arranged in groups (according to values of the BY variables) and that these groups are not necessarily in alphabetical or increasing numeric order.
- Create an index on the BY variables by using the DATASETS procedure (in Base SAS software).

Because sorting the data changes the order in which PROC GLMPOWER reads observations, the sort order for the levels of the classification variables might be affected if you have also specified **ORDER=DATA** in the **PROC GLMPOWER** statement. This, in turn, affects specifications in **CONTRAST** statements.

For more information about BY-group processing, see the discussion in *SAS Language Reference: Concepts*. For more information about the DATASETS procedure, see the discussion in the *Base SAS Procedures Guide*.

CLASS Statement

CLASS *variables* ;

The **CLASS** statement names the classification variables to be used in the analysis. If you use the **CLASS** statement, it must appear before the **MODEL** statement.

Classification variables can be either character or numeric. By default, class levels are determined from the entire set of formatted values of the **CLASS** variables.

CONTRAST Statement

CONTRAST *'label' effect values <... effect values> </ options>* ;

The **CONTRAST** statement enables you to define custom Type III hypothesis tests by specifying an **L** vector or matrix for testing either the hypothesis $\mathbf{L}\boldsymbol{\beta} = 0$ (for univariate models) or the hypothesis $\mathbf{LBM} = 0$ (for multivariate models). The **L** matrix consists of one or more between-subject contrasts.

To use this feature, you must be familiar with the details of the model parameterization that PROC GLM uses. For more information, see the section “[Parameterization of PROC GLM Models](#)” on page 3596 in Chapter 46, “[The GLM Procedure](#).” All the elements of the **L** matrix can be given, or if only certain portions of the **L** matrix are given, PROC GLMPOWER constructs the remaining elements from the context (in a manner similar to that in rule 4 in the section “[Construction of Least Squares Means](#)” on page 3629 in Chapter 46, “[The GLM Procedure](#)”).

There is no limit to the number of **CONTRAST** statements that you can specify, but they must appear after the **MODEL** statement. Each power analysis includes tests for all **CONTRAST** statements.

You can specify the following arguments:

- | | |
|---------------|---|
| <i>label</i> | identifies the contrast in the output. A label is required for every contrast that is specified. Labels must be enclosed in single or double quotation marks. |
| <i>effect</i> | identifies an effect that appears in the MODEL statement, or the INTERCEPT effect. You do not need to include all effects that appear in the MODEL statement. |

values are constants that are elements of the **L** matrix associated with the effect.

You can specify the following option in the **CONTRAST** statement after a slash (/):

SINGULAR=*number*

tunes the estimability checking. If $\text{ABS}(\mathbf{L} - \mathbf{LH}) > C \times \text{number}$ for any row in the contrast, then **L** is declared nonestimable. **H** is the $(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}$ matrix, and *C* is $\text{ABS}(\mathbf{L})$ except for rows where **L** is zero, and then it is 1. The default value for the **SINGULAR=** option is 10^{-4} . Values for the **SINGULAR=** option must be between 0 and 1.

As stated previously, the **CONTRAST** statement enables you to define custom hypothesis tests. If the hypothesis is testable in a univariate model, then the hypothesis sum of squares, $SS(H_0: \mathbf{L}\boldsymbol{\beta} = 0)$, is computed as

$$(\mathbf{Lb})'(\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}')^{-1}(\mathbf{Lb})$$

where $\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y}$.

For testable hypotheses in a multivariate model, the usual multivariate tests are defined by using

$$\mathbf{H} = \mathbf{M}'(\mathbf{LB})'(\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}')^{-1}(\mathbf{LB})\mathbf{M}$$

where $\mathbf{B} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}$ and **Y** is the matrix of multivariate responses or dependent variables.

The degrees of freedom associated with the hypothesis are equal to the row rank of **L**. The sum of squares computed in this situation is equivalent to the sum of squares computed using an **L** matrix with any row deleted that is a linear combination of previous rows.

Multiple-degrees-of-freedom hypotheses can be specified by separating the rows of the **L** matrix with commas.

MANOVA Statement

MANOVA *'label'* < test-options > < / detail-options > ;

If the **MODEL** statement includes more than one dependent variable, you can indicate a multivariate model and define transformations of dependent variables by using the **MANOVA** statement.

The **MANOVA** statement enables you to define custom Type III hypothesis tests by specifying an **M** vector or matrix for testing the hypothesis $\mathbf{L}\boldsymbol{\beta}\mathbf{M} = 0$. The **L** matrix consists of one or more between-subject contrasts that involve the model effects, and the **M** matrix consists of one or more within-subject contrasts.

To use this feature, you must be familiar with the details of multivariate model and contrast parameterizations that are used in PROC GLM. For more information, see the sections “[Multivariate Analysis of Variance](#)” on page 3632 and “[Repeated Measures Analysis of Variance](#)” on page 3633 in Chapter 46, “[The GLM Procedure](#).” For information about the power and sample size computational methods and formulas, see the section “[Contrasts in Fixed-Effect Multivariate Models](#)” on page 3773.

You can use either the **MANOVA** statement or the **REPEATED** statement with any of the tests for multivariate models that are supported in the **MTEST=** option in the **POWER** statement. For handling repeated measures on the same experimental unit, you would usually use the **REPEATED** statement instead of the **MANOVA** statement. But you can use the **MANOVA** statement in repeated measures situations, in addition to situations where you have clusters or multiple outcome variables. The differences between the **MANOVA** and **REPEATED** statements are as follows:

- You can use the **MANOVA** statement to construct any **M** matrix, but you must specify the coefficients explicitly (except for the default identity matrix).
- You can use the **REPEATED** statement to specify commonly used contrasts by using keywords rather than coefficients, but you are limited to only those forms of the **M** matrix.

There is no limit to the number of **MANOVA** statements that you can specify. Each power analysis includes tests for all **MANOVA** statements.

The *label* identifies the dependent variable transformation in the output. The label serves the same purpose as the *factor-name* in the **REPEATED** statement, enabling you to use the **MANOVA** statement for tests of within-subject effects and within-subject-by-between-subject interactions. A label is required for every transformation that is specified. Labels must be enclosed in single or double quotation marks.

Test Options

You can specify the following *test-option* in the **MANOVA** statement:

M=*equation, ..., equation* | (*row-of-matrix, ..., row-of-matrix*)

specifies a transformation matrix for the dependent variables that are listed in the **MODEL** statement.

The equations in the **M=** specification are of the form

$$c_1 \times \text{dependent-variable} \pm c_2 \times \text{dependent-variable} \\ \dots \pm c_n \times \text{dependent-variable}$$

where the c_i values are coefficients of the various *dependent-variables*. If the value of a given c_i is 1, it can be omitted; in other words, $1 \times Y$ is the same as Y . Equations should involve two or more dependent variables.

Alternatively, you can input the transformation matrix directly by entering the elements of the matrix, using commas to separate the rows and parentheses to surround the matrix. When you use this alternate form of input, the number of elements in each row must equal the number of dependent variables. Although these combinations actually represent the columns of the **M** matrix, they are displayed by rows.

When you include an **M=** specification, the tests are based on the variables that are defined by the equations in the specification, not the original dependent variables. If you omit the **M=** option, the tests are based on the original dependent variables in the **MODEL** statement. Omitting the **M=** option is equivalent to specifying an identity matrix, as in the following example, which assumes three dependent variables:

```
MANOVA 'Identity' M=(1 0 0,
                     0 1 0,
                     0 0 1);
```

For more examples of the **M=** option, see the section “**Examples**” on page 3571 in Chapter 46, “**The GLM Procedure**.” The syntax and functionality of the **M=** option in PROC GLM are the same as in PROC GLMPOWER.

Detail Options

You can specify the following *detail-option* in the **MANOVA** statement after a slash (/):

ORTH

requests that the transformation matrix in the **M=** specification of the **MANOVA** statement be orthogonalized by rows before the analysis.

MODEL Statement

MODEL *dependent-variables = independent-effects ;*

The **MODEL** statement names the dependent variables and independent effects. If one or more **MANOVA** or **REPEATED** statements are specified, then multiple dependent variables define a multivariate model. In the absence of the **MANOVA** and **REPEATED** statements, a univariate model is assumed, and multiple dependent variables represent different scenarios for the cell means.

The *independent-effects* can involve classification variables, continuous variables, or both. You can include main effects and interactions by using the effects notation of PROC GLM; for more information, see the section “[Specification of Effects](#)” on page 3593 in Chapter 46, “[The GLM Procedure](#).” For any model effect that involves classification variables (main effects and interactions), the number of levels cannot exceed 32,767. If no independent effects are specified, only an intercept term is fit. You can specify only one **MODEL** statement, and it must appear before the **POWER** statement if the **EFFECTS=** option is specified in the **POWER** statement.

For a univariate model, you can account for covariates without specifying them explicitly in the model by using the **NCOVARIATES=** option and either the **CORRXY=** or **PROPVARREDUCTION=** option in the **POWER** statement. For a multivariate model, you must explicitly specify any covariates in the **MODEL** statement.

The values of dependent variables in the exemplary data set (the data set named by the **DATA=** option in the **PROC GLMPOWER** statement) are surmised response means across subject profiles. For a univariate model, multiple dependent variables correspond to multiple scenarios for these cell means.

The **MODEL** statement is required. You can specify only one **MODEL** statement.

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 **POWER** statement. The *plot-options* define the plot characteristics, and the *graph-options* are like those in SAS/GRAPH software. If ODS Graphics is enabled, then the **PLOT** statement uses ODS Graphics to create graphs. For example:

```
ods graphics on;

proc glmpower data=Exemplary;
  class Variety Exposure;
  model Height = Variety | Exposure;
  power
    stddev = 4 6.5
    ntotal = 60
    power = .;
  plot x=n min=30 max=90;
run;

ods graphics off;
```

Otherwise, traditional graphics are produced. For example:

```
ods graphics off;

proc glmpower data=Exemplary;
  class Variety Exposure;
  model Height = Variety | Exposure;
  power
    stddev = 4 6.5
    ntotal = 60
    power = .;
  plot x=n min=30 max=90;
run;
```

For more information about enabling and disabling ODS Graphics, see the section “[Enabling and Disabling ODS Graphics](#)” on page 609 in Chapter 21, “[Statistical Graphics Using ODS](#).”

Table 48.4 summarizes the options available in the PLOT statement.

Table 48.4 PLOT Statement Options

Option	Description
Plot Options	
INTERPOL=	Specifies the type of curve to draw
KEY=	Specifies the style of key for the plot
MARKERS=	Specifies the locations for plotting symbols
MAX=	Specifies the maximum of the range of values
MIN=	Specifies the minimum of the range of values
NPOINTS=	Specifies the number of values
STEP=	Specifies the increment between values
VARY	Specifies how plot features should be linked to varying analysis parameters
X=	Specifies a plot with the requested type of parameter on the X axis
XOPTS=	Specifies plot characteristics pertaining to the X axis
Y=	Specifies a plot with the requested type of parameter on the Y axis
YOPTS=	Specifies plot characteristics pertaining to the Y axis
Graph Options	
DESCRIPTION=	Specifies a descriptive string
NAME=	Specifies a name for the catalog entry for the plot

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 with 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 **POWER** 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 **POWER** 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 **POWER** 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 *features* are COLOR, LINESYLE, 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 **POWER** statement preceding the **PLOT** statement, *or* one of the following keywords:

- **SOURCE**, which represents the model effects and contrasts in a univariate model and the between-subject effects and contrasts in a multivariate model
- **DEPENDENT**, which represents the cell means scenarios in a univariate model or the dependent variable transformations in a multivariate model

If the name represents an analysis option that is specified in the **POWER** statement, then it must be the *primary* name for the analysis option—that is, the one that is 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 are automatically selected for you.

X=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=N**, sample size is assigned to the X axis. When **X=POWER**, power is assigned to the X axis. You cannot use the **X=** and **Y=** options simultaneously. The default is **X=POWER**, unless the result parameter is power, in which case the default is **X=N**.

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 the POWER Statement](#)” on page 3766.

Y=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=N**, sample size is assigned to the Y axis. When **Y=POWER**, power is assigned to the Y axis. 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 the POWER Statement](#)” on page 3766.

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 GLMPOWER 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 GLMPOWER. 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 GLMPOWER.

POWER Statement

POWER < options > ;

The **POWER** statement performs power and sample size analyses for the Type III *F* tests that are specified in the **MTEST=** option, for each effect in the model that is defined by the **MODEL** statement, and for the contrasts that are defined by all **CONTRAST**, **MANOVA**, and **REPEATED** statements. The **POWER** statement must appear after the **MODEL** statement if the **EFFECTS=** option is used in the **POWER** statement.

For information about the power and sample size computational methods and formulas, see the section “Computational Methods and Formulas” on page 3770.

Summary of Options

Table 48.5 summarizes the options available in the **POWER** statement.

Table 48.5 POWER Statement Options

Option	Description
Specify test statistic	
MTEST=	Specifies the test statistic for a multivariate model
UEPSDEF=	Specifies the form of the Huynh-Feldt epsilon for MTEST=HF
Specify analysis information	
ALPHA=	Specifies the level of significance of each test
EFFECTS	Specifies the model effects
Specify covariates for a univariate model	
CORRXY=	Specifies multiple correlation (ρ) between covariates and response
NCOVARIATES=	Specifies additional degrees of freedom due to covariates
PROVARREDUCTION=	Specifies proportional variance reduction (r) due to covariates
Specify variability	
CORRMAT=	Specifies the correlation matrix of the dependent variables in a multivariate model
CORRS=	Specifies the correlations among the dependent variables in a multivariate model
COVMAT=	Specifies the covariance matrix of the dependent variables in a multivariate model
MATRIX=	Defines a matrix or vector
SQRTVAR=	Specifies the vector of error standard deviations for each dependent variable in a multivariate model
STDDEV=	Specifies the common error standard deviation
Specify sample size	
NFRACTIONAL	Enables fractional input and output for sample sizes
NTOTAL=	Specifies the sample size
Specify power	
POWER=	Specifies power
Choose computational method	
METHOD=	Specifies the computational method for multivariate tests

Table 48.5 *continued*

Option	Description
Control ordering in output	
DEPENDENT	Specifies the location of the Dependent or Transformation column in the output
OUTPUTORDER=	Controls ordering in output

Table 48.6 summarizes the valid result parameters.

Table 48.6 Summary of Result Parameters in the POWER Statement

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

Dictionary of Options

ALPHA=*number-list*

specifies the level of significance of each test. The default is 0.05, which corresponds to the usual $0.05 \times 100\% = 5\%$ level of significance. Note that this is a test-wise significance level with the same value for all tests, not incorporating any corrections for multiple testing. For information about specifying the *number-list*, see the section “[Specifying Value Lists in the POWER Statement](#)” on page 3766.

CORRMAT=*name-list*

specifies the correlation matrix of the dependent variables in a multivariate model, by using labels that are specified in the **MATRIX=** option. The corresponding matrices that are defined in the **MATRIX=** option must have either a lower triangular form that *includes* the diagonal of 1’s or a linear exponent autoregressive (LEAR) correlation structure. The matrix must be positive definite. You can use the **CORRMAT=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements. For information about specifying the *name-list*, see the section “[Specifying Value Lists in the POWER Statement](#)” on page 3766.

CORRS=*name-list*

specifies the correlations among the dependent variables in a multivariate model, by using labels that are specified in the **MATRIX=** option. The corresponding matrices that are defined in the **MATRIX=** option must have either a lower triangular form that *excludes* the diagonal of 1’s or a LEAR correlation structure. The matrix must be positive definite. You can use the **CORRS=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements. For information about specifying the *name-list*, see the section “[Specifying Value Lists in the POWER Statement](#)” on page 3766.

COVMAT=*name-list*

specifies the covariance matrix of the dependent variables in a multivariate model, by using labels that are specified in the **MATRIX=** option. The corresponding matrices that are defined in the **MATRIX=** option must have a lower triangular form that includes the diagonal of error variances. The matrix must be positive definite. You can use the **COVMAT=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements. For information about specifying the *name-list*, see the section “Specifying Value Lists in the POWER Statement” on page 3766.

CORRXY=*number-list*

specifies the multiple correlation (ρ) between all covariates and the response for a univariate model. The error standard deviation that is given by the **STDDEV=** option is consequently reduced by multiplying it by a factor of $(1 - \rho^2)^{\frac{1}{2}}$, provided that the number of covariates (as determined by the **NCOVARIATES=** option) is greater than 0. You cannot use the **CORRXY=** and **PROPVARREDUCTION=** options simultaneously. You cannot use the **CORRXY=** option when you have a multivariate model—that is, in the presence of a **MANOVA** or **REPEATED** statement. For information about specifying the *number-list*, see the section “Specifying Value Lists in the POWER Statement” on page 3766.

DEPENDENT

specifies the location of the Dependent column (for a univariate model) or the Transformation column (for a multivariate model) in the output when you specify the **OUTPUTORDER=REVERSE** option or **OUTPUTORDER=SYNTAX** option, according to its relative position in the **POWER** statement.

EFFECTS <= <(effect ... effect) >>

specifies the model effects to include in the power analysis. By default, or if the **EFFECTS** keyword is specified without the equal sign (=), all model effects are included. Specify **EFFECTS=()** to exclude all model effect tests from the power analysis. You can include main effects and interactions by using the effects notation of PROC GLM; for more information, see the section “Specification of Effects” on page 3593 in Chapter 46, “The GLM Procedure.” The **MODEL** statement must appear before the **POWER** statement if the **EFFECTS** option is used.

MATRIX('label')=*matrix-specification*

defines a matrix or vector that you can use along with the **CORRMAT=**, **CORRS=**, **COVMAT=**, and **SQRTVAR=** options when you have a multivariate model.

The *matrix-specification* can have one of the following three forms:

1. Raw values

(*values*)

specifies the values of a matrix or vector in one of the following forms:

- a matrix in lower triangular form, for use with the **CORRMAT=** or **COVMAT=** option
- a matrix in strictly lower triangular form, for use with the **CORRS=** option
- a vector, for use with the **SQRTVAR=** option

A matrix in lower triangular form contains the diagonal the and values below it. For example, you can represent a 3×3 correlation matrix for use with the **CORRMAT=** option as follows:


```
MATRIX ('MyCorrMat') = ( 1
                        0.5 1
                        0.2 0.5 1)
```

You can represent the same correlation matrix in strictly lower triangular form for use with the **CORRS=** option as follows:

```
MATRIX ('MyCorrs') = (0.5
                      0.2 0.5)
```

An example of a vector for use with the **SQRTVAR=** option is as follows:

```
MATRIX ('MySqrtVar') = (3.2 4.5 3.7)
```

2. Linear exponent autoregressive (LEAR) correlation structure

LEAR (*base-corr*, *corr-decay* <, *nlevels* <, *level-values* >>)

specifies a LEAR correlation structure for use with the **CORRMAT=** or **CORRS=** option. The LEAR structure is useful for characterizing exponentially decaying within-subject correlation that decays at a rate slower or faster than AR(1). Special cases include compound symmetry, first-order autoregressive (AR(1)), and first-order moving average correlation structures.

The LEAR correlation structure is related to the spatial covariance structures in PROC MIXED and is discussed in Simpson et al. (2010).

The *base-corr* (ρ) is the correlation between variables whose *level-values* are one unit apart, and it must satisfy $0 \leq \rho < 1$. The *corr-decay* (δ) is the correlation decay rate, which must be nonnegative. The default value for the number of levels, *nlevels* (n), is the number of dependent variables. The n *level-values*, denoted $\{l_1, \dots, l_n\}$, must be distinct. The default *level-values* are $\{1, \dots, n\}$. Let d_{jk} denote the distance between levels j and k , ($d_{jk} = |l_j - l_k|$). Let $d_{\min} = \min(d_{jk} : j \neq k)$ and $d_{\max} = \max(d_{jk} : j \neq k)$. The (i, j) element of the correlation matrix according to the LEAR model is defined as

$$\rho_{jk} = \begin{cases} 1 & \text{if } j = k \\ \rho^{d_{\min} + \delta[(d_{jk} - d_{\min}) / (d_{\max} - d_{\min})]} & \text{if } j \neq k \text{ and } d_{\min} \neq d_{\max} \\ \rho & \text{if } j \neq k \text{ and } d_{\min} = d_{\max} \end{cases}$$

Compound symmetry is the special case $\delta = 0$. AR(1) is the special case $\delta = d_{\max} - d_{\min}$. As $\delta \rightarrow \infty$, the model approaches the first-order moving average model.

3. Kronecker product

'*matrix-name*' @ '*matrix-name*' < @ ... @ '*matrix-name*' >

specifies a direct (Kronecker) product of two or more matrices for use with the **CORRMAT=**, **CORRS=**, or **COVMAT=** option. This form is useful when you have more than one type of distinction among dependent variables. For example, suppose you have a three-level repeated measurement factor with correlation that is 0.4 for neighboring measurements and that decays slightly more slowly than AR(1)

across more distant measurements. You also have four clusters that you believe satisfy compound symmetry with a correlation of 0.3. Your level values are the same as the default for the LEAR model. You can specify this correlation structure as follows:

```
MATRIX ('RepMeasures') = LEAR (0.4, 1.5, 3)
MATRIX ('Clusters') = LEAR (0.3, 0, 4)
MATRIX ('FullCorr') = 'RepMeasures' @ 'Clusters'
CORRMAT = 'FullCorr'
```

You can use the **MATRIX** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements.

METHOD=MULLERPETERSON | MP

METHOD=OBRIENSHIEH | OS

specifies the power computation method for the multivariate tests (**MTEST=HLT**, **MTEST=PT**, and **MTEST=WILKS**). **METHOD=OBRIENSHIEH** (the default) is based on O'Brien and Shieh (1992), and **METHOD=MULLERPETERSON** is based on Muller and Peterson (1984). For information about the associated power and sample size computational methods and formulas, see the section “**Multivariate Tests**” on page 3775.

If the dependent variable transformation consists of a single contrast ($r_M = 1$), then the two methods are identical and compute exact power. If $r_M > 1$ but the model effect or between-subject contrast has only one degree of freedom ($r_L = 1$), then **METHOD=OBRIENSHIEH** computes exact results and **METHOD=MULLERPETERSON** computes approximate results. If $r_M > 1$ and $r_L > 1$, then both methods compute approximate results.

You can use the **METHOD=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements.

MTEST=test-list

specifies the form of the F test for a multivariate model. Seven keywords are available, as discussed in the following paragraphs: **BOX**, **GG**, **HF**, **HLT**, **PT**, **UNCORR**, and **WILKS**. For information about specifying the *keyword-list*, see the section “**Specifying Value Lists in the POWER Statement**” on page 3766.

Three of these tests are multivariate, corresponding to the default **MSTAT=FAPPROX** option in the **MANOVA** and **REPEATED** statements in PROC GLM:

- **MTEST=HLT** (the default) is the Hotelling-Lawley trace
- **MTEST=PT** is Pillai's trace
- **MTEST=WILKS** is Wilks' lambda

For more information about these multivariate tests, see the section “**Multivariate Tests**” on page 94 in Chapter 4, “**Introduction to Regression Procedures**.” For information about the associated power and sample size computational methods and formulas, see the section “**Multivariate Tests**” on page 3775.

The other four tests are univariate, corresponding to the univariate approach to repeated measures in the **REPEATED** statement in PROC GLM:

- **MTEST=UNCORR** is the uncorrected univariate F test, assuming sphericity ($\varepsilon = 1$).

- **MTEST=GG** is the F test with the Greenhouse-Geisser adjustment, estimating ε by its maximum likelihood estimate $\hat{\varepsilon}$.
- **MTEST=HF** is the F test with the Huynh-Feldt adjustment as specified by the **UEPSDEF=** option, using an approximately unbiased estimate $\hat{\varepsilon}$.
- **MTEST=BOX** is the F test with Box's conservative adjustment, estimating sphericity by its smallest possible value $1/r_M$, where r_M is the number of within-subject contrasts in the dependent variable transformation.

For more information about these univariate tests, see the section “[Hypothesis Testing in Repeated Measures Analysis](#)” on page 3636 in Chapter 46, “[The GLM Procedure](#).” For information about the associated power and sample size computational methods and formulas, see the section “[Univariate Tests](#)” on page 3778.

These tests are all of the form $\mathbf{L}\boldsymbol{\beta}\mathbf{M} = 0$, where $\mathbf{L}$ is a between-subject contrast, $\boldsymbol{\beta}$ is the matrix of model parameters, and $\mathbf{M}$ is a within-subject contrast.

You can use the **MTEST=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements.

NCOVARIATES=*number-list*

NCOVARIATE=*number-list*

NCOV=*number-list*

NCOV=*number-list*

specifies the number of additional degrees of freedom to accommodate covariate effects—both class and continuous—not listed in the **MODEL** statement, for a univariate model. The error degrees of freedom are consequently reduced by the value of the **NCOVARIATES=** option, and the error standard deviation (whose unadjusted value is provided with the **STDDEV=** option) is reduced according to the value of the **CORRXY=** or **PROPVARREDUCTION=** option. You cannot use the **NCOVARIATES=** option when you have a multivariate model—that is, in the presence of a **MANOVA** or **REPEATED** statement. For information about specifying the *number-list*, see the section “[Specifying Value Lists in the POWER Statement](#)” on page 3766.

NFRACTIONAL

NFRAC

enables fractional input and output for sample sizes. See the section “[Sample Size Adjustment Options](#)” on page 3767 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=.**). The term “sample size” here refers to the number of independent sampling units. Values for the sample size for a univariate model must be no smaller than the model degrees of freedom (counting the covariates if present). The minimum required sample size for a multivariate model depends on the analysis and computational method; for more information, see the section “[Contrasts in Fixed-Effect Multivariate Models](#)” on page 3773. For information about specifying the *number-list*, see the section “[Specifying Value Lists in the POWER Statement](#)” on page 3766.

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:

- **DEPENDENT**
- **EFFECTS**
- weight variable (from the **WEIGHT** statement)
- **ALPHA=**
- **NCOVARIATES=**
- **CORRXY=**
- **PROPVARREDUCTION=**
- **STDDEV=**
- **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 **POWER** 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 **POWER** statement.

POWER=number-list

specifies the desired power of each test or requests a solution for the power by specifying a missing value (**POWER=.**). The power is expressed as a probability (for example, 0.9) rather than a percentage. Note that this is a test-wise power with the same value for all tests, without any correction for multiple testing. For information about specifying the *number-list*, see the section “Specifying Value Lists in the **POWER** Statement” on page 3766.

PROPVARREDUCTION=number-list**PVRED=number-list**

specifies the proportional reduction (r) in total R square incurred by the covariates—in other words, the amount of additional variation explained by the covariates—for a univariate model. The error standard deviation that is given by the **STDDEV=** option is consequently reduced by multiplying it by a factor of $(1 - r)^{\frac{1}{2}}$, provided that the number of covariates (as determined by the **NCOVARIATES=** option) is greater than 0. You cannot use the **PROPVARREDUCTION=** and **CORRXY=** options simultaneously. You cannot use the **PROPVARREDUCTION=** option when you have a multivariate model—that is, in the presence of a **MANOVA** or **REPEATED** statement. For information about specifying the *number-list*, see the section “Specifying Value Lists in the **POWER** Statement” on page 3766.

SQRTVAR=name-list

specifies the vector of standard deviations—that is, the square roots of the variances—of the dependent variables in a multivariate model, by using labels that are specified using the **MATRIX=** option. The standard deviation values must be positive. You can use the **SQRTVAR=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements. For information about specifying the *name-list*, see the section “Specifying Value Lists in the **POWER** Statement” on page 3766.

STDDEV=*number-list*

specifies the error standard deviation, or root MSE. For a multivariate model, each value in the *number-list* is taken to be a common value for all dependent variables. If covariates are specified by using the **NCOVARIATES=** option, then the **STDDEV=** option denotes the error standard deviation before accounting for these covariates. For information about specifying the *number-list*, see the section “Specifying Value Lists in the **POWER** Statement” on page 3766.

UEPSDEF=*unbiased-epsilon-definition*

specifies the type of adjustment for **MTEST=HF**. The default is **UEPSDEF=HFL**, which corresponds to the corrected form of the Huynh-Feldt adjustment (Huynh and Feldt 1976; Lecoutre 1991). Other alternatives are **UEPSDEF=HF**; the uncorrected Huynh-Feldt adjustment; and **UEPSDEF=CM**, the adjustment of Chi et al. (2012). For more information about these adjustments, see the section “Hypothesis Testing in Repeated Measures Analysis” on page 3636 in Chapter 46, “The GLM Procedure.” You can use the **UEPSDEF=** option only when you have a multivariate model—that is, in the presence of one or more **MANOVA** or **REPEATED** statements. For information about the associated power and sample size computational methods and formulas, see the section “Univariate Tests” on page 3778.

Restrictions on Option Combinations

To specify the variability in a multivariate model, choose one of the following parameterizations:

- covariance matrix (using the **MATRIX=** and **COVMAT=** options)
- standard deviations and correlations (using the **MATRIX=**, **SQRTVAR=**, and **CORRS=** options)
- common standard deviation and correlations (using the **STDDEV=**, **MATRIX=**, and **CORRS=** options)
- standard deviations and correlation matrix (using the **MATRIX=**, **SQRTVAR=**, and **CORRMAT=** options)
- common standard deviation and correlation matrix (using the **STDDEV=**, **MATRIX=**, and **CORRMAT=** options)

For the relationship between covariates and response in a univariate model, specify either the multiple correlation (by using the **CORRXY=** option) or the proportional reduction in total R square (by using the **PROPVARREDUCTION=** option).

REPEATED Statement

REPEATED *factor-specification* ;

If the **MODEL** statement includes more than one dependent variable, you can indicate a multivariate model and define transformations of dependent variables by using the **REPEATED** statement.

The **REPEATED** statement enables you to define custom Type III hypothesis tests by choosing from among several transformations of the dependent variables: contrast, Helmert, identity, mean, polynomial, and profile. You can specify a transformation for each repeated factor (often called a *within-subject factor*), and each combination of repeated factors produces an **M** vector or matrix for testing the hypothesis $\mathbf{L}\boldsymbol{\beta}\mathbf{M} = 0$. The **L** matrix consists of one or more between-subject contrasts that involve the model effects, and the **M** matrix

consists of one or more within-subject contrasts that involve the repeated factors. There is no limit to the number of repeated factors that you can specify.

Usually, the variables on the left side of the equation in the **MODEL** statement represent one repeated response variable. This does not mean that you are limited to listing only one factor in the **REPEATED** statement. For example, one repeated response variable (wellness rating) might be measured six times (implying variables Y1 to Y6 on the left side of the equal sign in the **MODEL** statement), with the associated within-subject factors rater and time (implying two factors listed in the **REPEATED** statement). However, designs that have two or more repeated response variables can be handled by using the **IDENTITY** transformation.

To use this feature, you must be familiar with the details of multivariate model and contrast parameterizations that PROC GLM uses. For more information, see the sections “[Repeated Measures Analysis of Variance](#)” on page 3633 and “[Multivariate Analysis of Variance](#)” on page 3632 in Chapter 46, “[The GLM Procedure](#).” For information about the power and sample size computational methods and formulas, see the section “[Contrasts in Fixed-Effect Multivariate Models](#)” on page 3773.

If you specify one or more **REPEATED** statements, then a “Mean(Dep)” transformation is added to the power analysis. This transformation is the mean of the dependent variables, the same transformation that is used implicitly in the “Tests of Hypotheses for Between Subjects Effects” table in PROC GLM. In addition, the Intercept model effect is included in the power analysis. If the **REPEATED** statement is not specified, then tests that involve the Intercept are excluded from the power analysis.

You can use either the **REPEATED** statement or the **MANOVA** statement along with any of the tests for multivariate models that are supported in the **MTEST=** option in the **POWER** statement. The **REPEATED** statement is usually used for handling repeated measurements on the same experimental unit, but you can also use the **REPEATED** statement for other situations, such as clusters or multiple outcome variables. The differences between the **REPEATED** and **MANOVA** statements are as follows:

- You can use the **REPEATED** statement to specify commonly used contrasts by using keywords rather than coefficients, but you are limited to only those forms of the **M** matrix.
- You can use the **MANOVA** statement to construct any **M** matrix, but you must specify the coefficients explicitly (except for the default identity matrix).

There is no limit to the number of **REPEATED** statements that you can specify. Each power analysis includes tests for all **REPEATED** statements and also (if you specify at least one **REPEATED** statement) the extra “Mean(Dep)” transformation that was previously mentioned.

The simplest form of the **REPEATED** statement requires only a *factor-name*. When you have two or more repeated factors, you must specify the *factor-name* and number of levels (*levels*) for each factor. Optionally, you can specify the actual values for the levels (*level-values*) and a *transformation* that defines single-degree-of-freedom contrasts. When you specify more than one within-subject factor, the *factor-names* (and associated level and transformation information) must be separated by a comma in the **REPEATED** statement.

The *factor-specification* for the **REPEATED** statement can include any number of individual factor specifications, separated by commas, of the following form:

factor-name levels < (*level-values*) > < *transformation* >

where

- factor-name* names a factor to be associated with the dependent variables. The name should not be the same as any variable name that already exists in the data set being analyzed and should conform to the usual conventions of SAS variable names.
- When you specify more than one factor, list the dependent variables in the **MODEL** statement so that the within-subject factors that you define in the **REPEATED** statement are nested; that is, the first factor you define in the **REPEATED** statement should be the one whose values change least frequently.
- levels* gives the number of levels associated with the factor being defined. When there is only one within-subject factor, the number of levels is equal to the number of dependent variables. In this case, *levels* is optional. When more than one within-subject factor is defined, however, *levels* is required, and the product of the number of levels of all the factors must equal the number of dependent variables in the **MODEL** statement.
- (level-values)* gives values that correspond to levels of a repeated measures factor. These values are used as spacings for constructing orthogonal polynomial contrasts if you specify a **POLYNOMIAL** transformation. The number of values that you specify must correspond to the number of levels for that factor in the **REPEATED** statement. Enclose the *level-values* in parentheses.

The following *transformation* keywords define single-degree-of-freedom contrasts for factors that you specify in the **REPEATED** statement. Because the number of contrasts that are generated is always one less than the number of levels of the factor, you have some control over which contrast is omitted from the analysis by which transformation you select. The only exception is the **IDENTITY** transformation, which is not composed of contrasts and has the same degrees of freedom as the factor has levels. By default, **PROC GLMPOWER** uses the **CONTRAST** transformation.

CONTRAST< (*ordinal-reference-level*) >

generates contrasts between levels of the factor and a reference level. By default, **PROC GLMPOWER** uses the last level as the reference level; you can optionally specify a reference level in parentheses after the keyword **CONTRAST**. The reference level corresponds to the ordinal value of the level rather than the level value that is specified. For example, to generate contrasts between the first level of a factor and the other levels, specify **CONTRAST(1)**.

HELMERT

generates contrasts between each level of the factor and the mean of subsequent levels.

IDENTITY

generates an identity transformation that corresponds to the associated factor. This transformation is *not* composed of contrasts; it has n degrees of freedom for an n -level factor, instead of $n - 1$ degrees of freedom.

MEAN< (*ordinal-reference-level*) >

generates contrasts between levels of the factor and the mean of all other levels of the factor. Specifying a reference level eliminates the contrast between that level and the mean. When no reference level is specified, the contrast that involves the last level is omitted. For an example, see the **CONTRAST** transformation.

POLYNOMIAL

generates orthogonal polynomial contrasts. Level values, if provided, are used as spacings in the construction of the polynomials; otherwise, equal spacing is assumed.

PROFILE

generates contrasts between adjacent levels of the factor.

Examples

When you specify more than one factor, list the dependent variables in the **MODEL** statement so that the within-subject factors that you define in the **REPEATED** statement are nested; that is, the first factor you define in the **REPEATED** statement should be the one whose values change least frequently. For example, assume that two raters submit a wellness rating at each of three times, for a total of six dependent variables for each subject. Consider the following statements:

```
proc glm;
  class treatment;
  model Y1-Y6 = treatment;
  repeated rater 2, time 3;
run;
```

The variables are listed in the **MODEL** statement as Y1 through Y6, so the **REPEATED** statement in the preceding statements implies the following structure:

Dependent Variables						
	Y1	Y2	Y3	Y4	Y5	Y6
Value of rater	1	1	1	2	2	2
Value of time	1	2	3	1	2	3

WEIGHT Statement

WEIGHT *variable* ;

The **WEIGHT** statement names a variable that provides a profile weight (“cell weight”) for each observation in the exemplary data set specified by the **DATA=** option in the **PROC GLMPOWER** statement.

If the **WEIGHT** statement is not used, then a balanced design is assumed with default cell weights of 1.

Details: GLMPOWER Procedure

Specifying Value Lists in the POWER Statement

To specify one or more scenarios for an analysis parameter (or set of parameters) in the **POWER** statement, you provide a list of values for the option that corresponds to the parameter(s). To identify the parameter you want to solve for, you place a missing value in the appropriate list.

There are three basic types of such lists: *number-lists*, *name-lists*, and *keyword-lists*. Scenarios for scalar-valued parameters, such as power, are represented by a *number-list*. Scenarios for named parameters, such as correlation matrices, are represented by a *name-list*. Some parameters, such as the test statistic for a multivariate model, have values that are represented by one or more keywords in a *keyword-list*.

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. For example, you can specify four scenarios (30, 50, 70, and 100) for a total sample size in either 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 = .
```

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 correlation matrix in a multivariate model as follows:

```
CORRMAT = "Corr A" "Corr B"
```

Keyword-Lists

A *keyword-list* is a list of one or more keywords, separated by spaces. For example, you can specify both the multivariate Hotelling-Lawley trace and uncorrected univariate *F* test for a multivariate model as follows:

```
MTEST = HLT UNCORR
```

Sample Size Adjustment Options

By default, PROC GLMPower rounds sample sizes conservatively (down in the input, up in the output) so that all total sizes *and* sample sizes for individual design profiles 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, all design profile sizes are adjusted to be multiples of their corresponding weights. If a design profile is present more than once in the exemplary data set, then the weights for that design profile are summed. For example, if a particular design profile is present twice in the exemplary data set with weight values 2 and 6, then all sample sizes for this design profile become multiples of $2 + 6 = 8$.

With the **NFRACTIONAL** option, sample size input is not rounded, and sample size output is 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 explains reasons for missing results and flags numerical results that are bounds rather than exact answers.

The Info column provides further information about Error 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 one-way ANOVA:

```
data MyExemp;
  input A $ Y1 Y2;
  datalines;
    1   10 11
    2   12 11
    3   15 11
  ;

proc glmpower data=MyExemp;
  class A;
  model Y1 Y2 = A;
  power
    stddev = 2
    ntotal = 3 10
    power = .;
  ods output output=Power;
run;

proc print noobs data=Power;
  var NominalNTotal NTotal Dependent Power Error Info;
run;
```

The output is shown in [Figure 48.5](#).

Figure 48.5 Error and Information Columns

Nominal	NTotal	NTotal	Dependent	Power	Error	Info
3	3	Y1	.	Invalid input	Error DF=0	
10	9	Y1	0.557		Input N adjusted	
3	3	Y2	.	Invalid input	Error DF=0 / No effect	
10	9	Y2	0.050		Input N adjusted / No effect	

The sample size of 3 specified with the **NTOTAL=** option causes an “Invalid input” message in the Error column and an “Error DF=0” message in the Info column, because a sample size of 3 is so small that there are no degrees of freedom left for the error term. The sample size of 10 causes an “Input N adjusted” message in the Info column, because it is rounded down to 9 to produce integer group sizes of 3 per cell. The cell means scenario represented by the dependent variable Y2 causes a “No effect” message to appear in the Info column, because the means in this scenario are all equal.

Displayed Output

If you use the **PLOTONLY** option in the **PROC GLMPower** statement, the procedure displays only graphical output. Otherwise, the displayed output of the GLMPower procedure includes the following:

- the “Fixed Scenario Elements” table, which shows all applicable single-valued analysis parameters, in the following order: the dependent variable that represents the cell means scenario (for a univariate model) or the dependent variable transformation (for a multivariate model), the source of the test (that is, the model effect or between-subject contrast), the weight variable, parameters that are input explicitly, parameters that are supplied with defaults, and ancillary results
- an output table that shows the following when applicable (in order): the index of the scenario, the dependent variable that represents the cell means scenario (for a univariate model) or the dependent variable transformation (for a multivariate model), the type of the test, the source of the test (that is, the model effect or between-subject contrast), all multivalued input, ancillary results, the primary computed result, and error descriptions
- plots (if requested)

The exception to these ordering conventions is that the **DEPENDENT** and **EFFECTS=** options can be used along with the **OUTPUTORDER=SYNTAX** or **OUTPUTORDER=REVERSE** option in the **POWER** statement to specify the relative location of the output for dependent variable and type and source of test.

Ancillary results include the following:

- Actual Power, the achieved power, if it differs from the input (Nominal) power 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 **POWER** statement
- test or numerator degrees of freedom in the test’s critical value
- error or denominator degrees of freedom in the test’s critical value
- Effect, the combination of the within-subject Transformation contrast and between-subject Source test or contrast in a multivariate model

If sample size is the result parameter and the **NFRACTIONAL** option is used in the **POWER** statement, then both “Fractional” and “Ceiling” sample size results are displayed. Fractional sample sizes correspond to the “Nominal” values of power. Ceiling sample sizes are simply the fractional sample sizes rounded up to the nearest integer; they correspond to “Actual” values of power.

The noncentrality parameter is computed and stored in a hidden column called Noncentrality in the “Output” table. If a univariate test for a multivariate model is specified (that is, one of **MTEST=BOX**, **MTEST=GG**, **MTEST=HF**, or **MTEST=UNCORR**), then the numerator and denominator degrees of freedom that are used in the noncentral F approximation of the test statistic distribution are computed and stored in hidden columns called NumNCDF and DenNCDF, respectively, in the “Output” table. These are the only tests for which the degrees of freedom in the noncentral F approximation of the test statistic are different from those in the critical value.

ODS Table Names

PROC GLMPOWER 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 48.7. For more information about ODS, see Chapter 20, “Using the Output Delivery System.”

Table 48.7 ODS Tables Produced by PROC GLMPOWER

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 GLMPOWER output.)	PLOT

Computational Methods and Formulas

This section describes the approaches that PROC GLMPOWER uses to compute power and sample size.

Contrasts in Fixed-Effect Univariate Models

The univariate linear model has the form

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$$

where $\mathbf{y}$ is the $N \times 1$ vector of responses, $\mathbf{X}$ is the $N \times k$ design matrix, $\boldsymbol{\beta}$ is the $k \times 1$ vector of model parameters that correspond to the columns of $\mathbf{X}$, and $\boldsymbol{\epsilon}$ is an $N \times 1$ vector of errors with

$$\epsilon_1, \dots, \epsilon_N \sim N(0, \sigma^2) \quad (\text{iid})$$

In PROC GLMPOWER, the model parameters $\boldsymbol{\beta}$ are not specified directly, but rather indirectly as $\mathbf{y}^*$, which represents either conjectured response means or typical response values for each design profile. The $\mathbf{y}^*$ values are manifested as the dependent variable in the **MODEL** statement. The vector $\boldsymbol{\beta}$ is obtained from $\mathbf{y}^*$ according to the least squares equation,

$$\boldsymbol{\beta} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y}^*$$

Note that, in general, there is not a one-to-one mapping between $\mathbf{y}^*$ and $\boldsymbol{\beta}$. Many different scenarios for $\mathbf{y}^*$ might lead to the same $\boldsymbol{\beta}$. If you specify $\mathbf{y}^*$ with the intention of representing cell means, keep in mind that PROC GLMPOWER allows scenarios that are *not* valid cell means according to the model that is specified in the **MODEL** statement. For example, if $\mathbf{y}^*$ exhibits an interaction effect but the corresponding interaction term is left out of the model, then the cell means ($\mathbf{X}\boldsymbol{\beta}$) that are derived from $\boldsymbol{\beta}$ differ from $\mathbf{y}^*$. In particular, the cell means that are derived in this way are the projection of $\mathbf{y}^*$ onto the model space.

It is convenient in power analysis to parameterize the design matrix $\mathbf{X}$ in three parts, $\{\ddot{\mathbf{X}}, \mathbf{w}, N\}$, defined as follows:

1. The $q \times k$ essence design matrix $\ddot{\mathbf{X}}$ is the collection of unique rows of $\mathbf{X}$. Its rows are sometimes referred to as “design profiles.” Here, $q \leq N$ is defined simply as the number of unique rows of $\mathbf{X}$.
2. The $q \times 1$ weight vector $\mathbf{w}$ reveals the relative proportions of design profiles, and $\mathbf{W} = \text{diag}(\mathbf{w})$. Row i of $\ddot{\mathbf{X}}$ is to be included in the design w_i times for every w_j times that row j is included. The weights are assumed to be standardized (that is, they sum up to 1).
3. The total sample size is N . This is the number of rows in $\mathbf{X}$. If you gather $Nw_i = n_i$ copies of the i th row of $\ddot{\mathbf{X}}$, for $i = 1, \dots, q$, then you end up with $\mathbf{X}$.

The preceding quantities are derived from PROC GLMPOWER syntax as follows:

- Values for $\ddot{\mathbf{X}}$, $\mathbf{y}^*$, and $\mathbf{w}$ are specified in the exemplary data set (from using the **DATA=** option in the **PROC GLMPOWER** statement), and the corresponding variables are identified in the **CLASS**, **MODEL**, and **WEIGHT** statements.
- N is specified in the **NTOTAL=** option in the **POWER** statement.

It is useful to express the crossproduct matrix $\mathbf{X}'\mathbf{X}$ in terms of these three parts,

$$\mathbf{X}'\mathbf{X} = N\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}}$$

because this expression factors out the portion (N) that depends on sample size and the portion ($\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}}$) that depends only on the design structure.

A general linear hypothesis for the univariate model has the form

$$H_0: \mathbf{L}\boldsymbol{\beta} = \boldsymbol{\theta}_0$$

$$H_A: \mathbf{L}\boldsymbol{\beta} \neq \boldsymbol{\theta}_0$$

where $\mathbf{L}$ is an $l \times k$ contrast matrix with rank r_L and $\boldsymbol{\theta}_0$ is the null value (usually just a vector of zeros).

Note that model effect tests are just contrasts that use special forms of $\mathbf{L}$. Thus, this scheme covers both effect tests (which are specified in the **MODEL** statement and the **EFFECTS=** option in the **POWER** statement) and custom contrasts (which are specified in the **CONTRAST** statement).

The model degrees of freedom DF_M are equal to the rank of $\mathbf{X}$, denoted r_X . The error degrees of freedom DF_E are equal to $N - r_X$. The sample size N must be at least DF_M plus the number of covariates.

The test statistic is

$$F = \frac{\left(\frac{SS_H}{r_L} \right)}{\hat{\sigma}^2}$$

where

$$\begin{aligned} \text{SS}_H &= \frac{1}{N} (\mathbf{L}\hat{\boldsymbol{\beta}} - \boldsymbol{\theta}_0)' (\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}')^{-1} (\mathbf{L}\hat{\boldsymbol{\beta}} - \boldsymbol{\theta}_0) \\ \hat{\boldsymbol{\beta}} &= (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y} \\ \hat{\sigma}^2 &= \frac{1}{\text{DF}_E} (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})' (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) \end{aligned}$$

Under H_0 , $F \sim F(r_L, \text{DF}_E)$. Under H_A , F is distributed as $F(r_L, \text{DF}_E, \lambda)$ with noncentrality

$$\lambda = N (\mathbf{L}\boldsymbol{\beta} - \boldsymbol{\theta}_0)' \left(\mathbf{L}(\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}})^{-1}\mathbf{L}' \right)^{-1} (\mathbf{L}\boldsymbol{\beta} - \boldsymbol{\theta}_0) \sigma^{-2}$$

The value of σ is specified in the **STDDEV=** option in the **POWER** statement.

Muller and Peterson (1984) give the exact power of the test as

$$\text{power} = P(F(r_L, \text{DF}_E, \lambda) \geq F_{1-\alpha}(r_L, \text{DF}_E))$$

The value of α is specified in the **ALPHA=** option in the **POWER** statement.

Sample size is computed by inverting the power equation.

See Muller and Benignus (1992) and O'Brien and Shieh (1992) for additional discussion.

Adjustments for Covariates in Univariate Models

If you specify covariates in a univariate model (whether continuous or categorical), then two adjustments are made in order to compute approximate power in the presence of the covariates. Let n_v denote the number of covariates (counting dummy variables for categorical covariates individually) as specified in the **NCOVARIATES=** option in the **POWER** statement. In other words, n_v is the total degrees of freedom used by the covariates. The adjustments are as follows:

1. The error degrees of freedom decrease by n_v .
2. The error standard deviation σ shrinks by a factor of $(1 - \rho^2)^{\frac{1}{2}}$ (if the **CORRXY=** option is used to specify the correlation ρ between covariates and response) or $(1 - r)^{\frac{1}{2}}$ (if the **PROPVARREDUCTION=** option is used to specify the proportional reduction in total R^2 incurred by the covariates). Let σ^* represent the updated value of σ .

As a result of these changes, the power is computed as

$$\text{power} = P(F(r_L, \text{DF}_E - n_v, \lambda^*) \geq F_{1-\alpha}(r_L, N - r_x - n_v))$$

where λ^* is calculated using σ^* rather than σ :

$$\lambda^* = N (\mathbf{L}\boldsymbol{\beta} - \boldsymbol{\theta}_0)' \left(\mathbf{L}(\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}})^{-1}\mathbf{L}' \right)^{-1} (\mathbf{L}\boldsymbol{\beta} - \boldsymbol{\theta}_0) (\sigma^*)^{-2}$$

Contrasts in Fixed-Effect Multivariate Models

The multivariate model has the form

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$$

where $\mathbf{Y}$ is the $N \times p$ vector of responses, for $p > 1$; $\mathbf{X}$ is the $N \times k$ design matrix; $\boldsymbol{\beta}$ is the $k \times p$ matrix of model parameters that correspond to the columns of $\mathbf{X}$ and $\mathbf{Y}$; and $\boldsymbol{\epsilon}$ is an $N \times p$ vector of errors, where

$$\epsilon_1, \dots, \epsilon_N \sim N(0, \boldsymbol{\Sigma}) \quad (\text{iid})$$

In PROC GLMPower, the model parameters $\boldsymbol{\beta}$ are not specified directly, but rather indirectly as $\mathbf{Y}^*$, which represents either conjectured response means or typical response values for each design profile. The $\mathbf{Y}^*$ values are manifested as the collection of dependent variables in the **MODEL** statement. The matrix $\boldsymbol{\beta}$ is obtained from $\mathbf{Y}^*$ according to the least squares equation,

$$\boldsymbol{\beta} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}^*$$

Note that, in general, there is not a one-to-one mapping between $\mathbf{Y}^*$ and $\boldsymbol{\beta}$. Many different scenarios for $\mathbf{Y}^*$ might lead to the same $\boldsymbol{\beta}$. If you specify $\mathbf{Y}^*$ with the intention of representing cell means, keep in mind that PROC GLMPower allows scenarios that are *not* valid cell means according to the model that is specified in the **MODEL** statement. For example, if $\mathbf{Y}^*$ exhibits an interaction effect but the corresponding interaction term is left out of the model, then the cell means ($\mathbf{X}\boldsymbol{\beta}$) that are derived from $\boldsymbol{\beta}$ differ from $\mathbf{Y}^*$. In particular, the cell means that are derived in this way are the projection of $\mathbf{Y}^*$ onto the model space.

It is convenient in power analysis to parameterize the design matrix $\mathbf{X}$ in three parts, $\{\ddot{\mathbf{X}}, \mathbf{W}, N\}$, defined as follows:

1. The $q \times k$ essence design matrix $\ddot{\mathbf{X}}$ is the collection of unique rows of $\mathbf{X}$. Its rows are sometimes referred to as “design profiles.” Here, $q \leq N$ is defined simply as the number of unique rows of $\mathbf{X}$.
2. The $q \times 1$ weight vector $\mathbf{w}$ reveals the relative proportions of design profiles, and $\mathbf{W} = \text{diag}(\mathbf{w})$. Row i of $\ddot{\mathbf{X}}$ is to be included in the design w_i times for every w_j times that row j is included. The weights are assumed to be standardized (that is, they sum up to 1).
3. The total sample size is N . This is the number of rows in $\mathbf{X}$. If you gather $Nw_i = n_i$ copies of the i th row of $\ddot{\mathbf{X}}$, for $i = 1, \dots, q$, then you end up with $\mathbf{X}$.

The preceding quantities are derived from PROC GLMPower syntax as follows:

- Values for $\ddot{\mathbf{X}}$, $\mathbf{Y}^*$, and $\mathbf{w}$ are specified in the exemplary data set (from using the **DATA=** option in the **PROC GLMPower** statement), and the corresponding variables are identified in the **CLASS**, **MODEL**, and **WEIGHT** statements.
- N is specified in the **NTOTAL=** option in the **POWER** statement.

It is useful to express the crossproduct matrix $\mathbf{X}'\mathbf{X}$ in terms of these three parts,

$$\mathbf{X}'\mathbf{X} = N\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}}$$

because this expression factors out the portion (N) that depends on sample size and the portion ($\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}}$) that depends only on the design structure.

A general linear hypothesis for the univariate model has the form

$$H_0: \mathbf{L}\boldsymbol{\beta}\mathbf{M} = \boldsymbol{\theta}_0$$

$$H_A: \mathbf{L}\boldsymbol{\beta}\mathbf{M} \neq \boldsymbol{\theta}_0$$

where $\mathbf{L}$ is an $l \times k$ between-subject contrast matrix with rank r_L , $\mathbf{M}$ is a $p \times m$ within-subject contrast matrix with rank r_M , and $\boldsymbol{\theta}_0$ is an $l \times m$ null contrast matrix (usually just a matrix of zeros).

Note that model effect tests are just between-subject contrasts that use special forms of $\mathbf{L}$, combined with an $\mathbf{M}$ that is the $p \times 1$ mean transformation vector of the dependent variables (a vector of values all equal to $1/p$). Thus, this scheme covers both effect tests (which are specified in the **MODEL** statement and the **EFFECTS=** option in the **POWER** statement) and custom between-subject contrasts (which are specified in the **CONTRAST** statement).

The $\mathbf{M}$ matrix is often referred to as the dependent variable transformation and is specified in the **MANOVA** or **REPEATED** statement.

The model degrees of freedom DF_M are equal to the rank of $\mathbf{X}$, denoted r_X . The error degrees of freedom DF_E are equal to $N - r_X$.

The hypothesis sum of squares SS_H in the univariate model generalizes to the hypothesis SSCP matrix in the multivariate model,

$$\mathbf{H} = (\mathbf{L}\hat{\boldsymbol{\beta}}\mathbf{M} - \boldsymbol{\theta}_0)' (\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}')^{-1} (\mathbf{L}\hat{\boldsymbol{\beta}}\mathbf{M} - \boldsymbol{\theta}_0)$$

The error sum of squares $\hat{\sigma}^2(N - r_X)$ in the univariate model generalizes to the error SSCP matrix in the multivariate model,

$$\mathbf{E} = (N - r_X)\mathbf{M}'\hat{\boldsymbol{\Sigma}}\mathbf{M}$$

where

$$\hat{\boldsymbol{\Sigma}} = (\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\beta}})'(\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\beta}})/(N - r_X)$$

and

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}$$

The population counterpart of $\mathbf{H}/N$ is

$$\mathbf{H}^* = (\mathbf{L}\boldsymbol{\beta}\mathbf{M} - \boldsymbol{\theta}_0)' (\mathbf{L}(\ddot{\mathbf{X}}'\mathbf{W}\ddot{\mathbf{X}})^{-1}\mathbf{L}')^{-1} (\mathbf{L}\boldsymbol{\beta}\mathbf{M} - \boldsymbol{\theta}_0)$$

and the population counterpart of $\mathbf{E}/N$ is

$$\mathbf{E}^* = \mathbf{M}'\boldsymbol{\Sigma}\mathbf{M}$$

The elements of $\boldsymbol{\Sigma}$ are specified in the **MATRIX=** and **STDDEV=** options and identified in the **CORRMAT=**, **CORRS=**, **COVMAT=**, and **SQRTVAR=** options in the **POWER** statement.

The power and sample size computations for all the tests that are supported in the **MTEST=** option in the **POWER** statement are based on $\mathbf{H}^*$ and $\mathbf{E}^*$. The following two subsections cover the computational methods and formulas for the multivariate and univariate tests that are supported in the **MTEST=** and **UEPSDEF=** options in the **POWER** statement.

Multivariate Tests

Power computations for multivariate tests are based on O'Brien and Shieh (1992) (for **METHOD=OBRIENSHIEH**) and Muller and Peterson (1984) (for **METHOD=MULLERPETERSON**).

Let $s = \min(r_L, r_M)$, the smaller of the between-subject and within-subject contrast degrees of freedom. Critical value computations assume that under H_0 , the test statistic F is distributed as $F(r_L r_M, \nu_2)$, where $\nu_2 = (N - r_X) - r_M + 1$ if $s = 1$ but depends on the choice of test if $s > 1$. Power computations assume that under H_A , F is distributed as $F(r_L r_M, \nu_2, \lambda)$, where the noncentrality λ depends on r_L, r_M , the choice of test, and the power computation method.

Formulas for the test statistic F , denominator degrees of freedom ν_2 , and noncentrality λ for all combinations of dimensions, tests, and methods are given in the following subsections.

The power in each case is computed as

$$\text{power} = P(F(r_L r_M, \nu_2, \lambda) \geq F_{1-\alpha}(r_L r_M, \nu_2))$$

Computed power is exact for some cases and approximate for others. Sample size is computed by inverting the power equation.

Let $\Delta = \mathbf{E}^{-1}\mathbf{H}$, and define ϕ as the $s \times 1$ vector of ordered positive eigenvalues of Δ , $\phi = \{\phi_1, \dots, \phi_s\}$, where $\phi_1 \geq \dots \geq \phi_s > 0$. The population equivalent is

$$\begin{aligned} \Delta^* &= \mathbf{E}^{*-1}\mathbf{H}^* \\ &= (\mathbf{M}'\Sigma\mathbf{M})^{-1}(\mathbf{L}\beta\mathbf{M} - \theta_0)' \left(\mathbf{L} \left(\ddot{\mathbf{X}}' \mathbf{W} \ddot{\mathbf{X}} \right)^{-1} \mathbf{L}' \right)^{-1} (\mathbf{L}\beta\mathbf{M} - \theta_0) \end{aligned}$$

where ϕ^* is the $s \times 1$ vector of ordered positive eigenvalues of Δ^* , $\phi^* = \{\phi_1^*, \dots, \phi_s^*\}$ for $\phi_1^* \geq \dots \geq \phi_s^* > 0$.

Case 1: $s = 1$

When $s = 1$, all three multivariate tests (**MTEST=HLT**, **MTEST=PT**, and **MTEST=WILKS**) are equivalent. The test statistic is $F = \phi_1 \nu_2 / (r_L r_M)$, where $\nu_2 = (N - r_X) - r_M + 1$.

When the dependent variable transformation has a single degree of freedom ($r_M = 1$), **METHOD=OBRIENSHIEH** and **METHOD=MULLERPETERSON** are the same, computing exact power by using noncentrality $\lambda = N\Delta^*$. The sample size must satisfy $N \geq r_X + 1$.

When the dependent variable transformation has more than one degree of freedom but the between-subject contrast has a single degree of freedom ($r_M > 1, r_L = 1$), **METHOD=OBRIENSHIEH** computes exact power by using noncentrality $\lambda = N\phi_1^*$, and **METHOD=MULLERPETERSON** computes approximate power by using

$$\lambda = \frac{(N - r_X) - r_M + 1}{(N - r_X)} N\phi_1^*$$

The sample size must satisfy $N \geq r_X + r_M$.

Case 2: $s > 1$

When both the dependent variable transformation and the between-subject contrast have more than one degree of freedom ($s > 1$), **METHOD=OBRIENSHIEH** computes the noncentrality as $\lambda = N\lambda^*$, where λ^* is the primary noncentrality. The form of λ^* depends on the choice of test statistic.

METHOD=MULLERPETERSON computes the noncentrality as $\lambda = v_2 \lambda^{(\text{MP})\star}$, where $\lambda^{(\text{MP})\star}$ has the same form as $\lambda^\star$ except that $\phi^\star$ is replaced by

$$\phi^{(\text{MP})\star} = \frac{N}{(N - r_X)} \phi^\star$$

Computed power is approximate for both methods when $s > 1$.

Hotelling-Lawley Trace (MTEST=HLT) When $s > 1$

If $N > r_X + r_M + 1$, then the denominator degrees of freedom for the Hotelling-Lawley trace are $v_2 = v_{2a}$,

$$v_{2a} = 4 + (r_L r_M + 2)g$$

where

$$g = \frac{(N - r_X)^2 - (N - r_X)(2r_M + 3) + r_M(r_M + 3)}{(N - r_X)(r_L + r_M + 1) - (r_L + 2r_M + r_M^2 - 1)}$$

which is the same as $v_2^{(T_2)}$ in O'Brien and Shieh (1992) and is due to McKeon (1974).

If $N \leq r_X + r_M + 1$, then $v_2 = v_{2b}$,

$$v_{2b} = s((N - r_X) - r_M - 1) + 2$$

which is the same as both $v_2^{(T_1)}$ in O'Brien and Shieh (1992) and v_2 in Muller and Peterson (1984) and is due to Pillai and Samson (1959).

The primary noncentrality is

$$\lambda^\star = \sum_{i=1}^s \phi_i^\star$$

The sample size must satisfy

$$N \geq r_X + r_M + 1 - 1/s$$

If $N > r_X + r_M + 1$, then the test statistic is

$$F = \frac{U/v_1}{c/v_{2a}}$$

where

$$\begin{aligned} U &= \text{trace}(\mathbf{E}^{-1} \mathbf{H}) \\ &= \sum_{i=1}^s \phi_i \end{aligned}$$

and

$$c = \frac{2 + (r_L r_M + 2)g}{N - r_X - r_M - 1}$$

If $N \leq r_X + r_M + 1$, then the test statistic is

$$F = \frac{U/v_1}{s/v_2b}$$

Pillai's Trace (MTEST=PT) When $s > 1$

The denominator degrees of freedom for Pillai's trace are

$$v_2 = s((N - r_X) + s - r_M)$$

The primary noncentrality is

$$\lambda^* = s \left(\frac{\sum_{i=1}^s \frac{\phi_i^*}{1+\phi_i^*}}{s - \sum_{i=1}^s \frac{\phi_i^*}{1+\phi_i^*}} \right)$$

The sample size must satisfy

$$N \geq r_X + r_M + 1/s - s$$

The test statistic is

$$F = \frac{V/v_1}{(s - V)/v_2}$$

where

$$\begin{aligned} V &= \text{trace}(\mathbf{H}(\mathbf{H} + \mathbf{E})^{-1}) \\ &= \sum_{i=1}^s \frac{\phi_i}{1 + \phi_i} \end{aligned}$$

Wilks' Lambda (MTEST=WILKS) When $s > 1$

The denominator degrees of freedom for Wilks' lambda are

$$v_2 = t[(N - r_X) - 0.5(r_M - r_L + 1)] - 0.5(r_L r_M - 2)$$

where

$$t = \begin{cases} 1 & \text{if } r_L r_M \leq 3 \\ \left[\frac{(r_L r_M)^2 - 4}{r_L^2 + r_M^2 - 5} \right]^{\frac{1}{2}} & \text{if } r_L r_M \geq 4 \end{cases}$$

The primary noncentrality is

$$\lambda^* = t \left[\left(\prod_{i=1}^s [(1 + \phi_i^*)^{-1}] \right)^{-\frac{1}{t}} - 1 \right]$$

The sample size must satisfy

$$N \geq (1 + 0.5(r_L r_M - 2))/t + r_X + (r_M - r_L + 1)/2$$

The test statistic is

$$F = \frac{(1 - \Lambda^{1/t})/\nu_1}{\Lambda^{1/t}/\nu_2}$$

where

$$\begin{aligned}\Lambda &= \det(\mathbf{E})/\det(\mathbf{H} + \mathbf{E}) \\ &= \prod_{i=1}^s [(1 + \phi_i)^{-1}]\end{aligned}$$

Univariate Tests

Power computations for univariate tests are based on Muller et al. (2007) and Muller and Barton (1989).

The test statistic is

$$F = \frac{\text{trace}(\mathbf{H})/r_L}{\text{trace}(\mathbf{E})/(N - r_X)}$$

Critical value computations assume that under H_0 , F is distributed as $F(\nu_1, \nu_2)$, where ν_1 and ν_2 depend on the choice of test.

The four tests for the univariate approach to repeated measures differ in their assumptions about the sphericity ε of $\mathbf{E}^*$,

$$\varepsilon = \frac{\text{trace}^2(\mathbf{E}^*)}{r_M \text{trace}(\mathbf{E}^{*2})}$$

Power computations assume that under H_A , F is distributed as $F(\nu_1^*, \nu_2^*, \lambda)$.

Formulas for ν_1 and ν_2 for each test and formulas for ν_1^* , ν_2^* , and λ are given in the following subsections.

The power in each case is approximated as

$$\text{power} = P(F(\nu_1^*, \nu_2^*, \lambda) \geq F_{1-\alpha}(\nu_1, \nu_2))$$

Sample size is computed by inverting the power equation.

The sample size must be large enough to yield $\nu_1 > 0$, $\nu_1^* > 0$, $\nu_2 \geq 1$, and $\nu_2^* \geq 1$.

Because these univariate tests are 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 noncentrality parameter λ is set to 0.

Define $\boldsymbol{\phi}^{(E)}$ as the vector of ordered eigenvalues of $\mathbf{E}^*$, $\boldsymbol{\phi}^{(E)} = \{\phi_1^{(E)}, \dots, \phi_{r_M}^{(E)}\}$, where $\phi_1^{(E)} \geq \dots \geq \phi_{r_M}^{(E)}$, and define $\boldsymbol{\gamma}_j^{(E)}$ as the j th eigenvector of $\mathbf{E}^*$. Critical values and power computations are based on the following intermediate parameters:

$$\omega_{*j} = N \left(\boldsymbol{\gamma}_j^{(E)} \right)' \mathbf{H}^* \boldsymbol{\gamma}_j^{(E)} / \phi_j^{(E)}$$

$$S_{t1} = \sum_{j=1}^{r_M} \phi_j^{(E)}$$

$$S_{t2} = \sum_{j=1}^{r_M} \phi_j^{(E)} \omega_{*j}$$

$$S_{t3} = \sum_{j=1}^{r_M} \left(\phi_j^{(E)} \right)^2$$

$$S_{t4} = \sum_{j=1}^{r_M} \left(\phi_j^{(E)} \right)^2 \omega_{*j}$$

$$R_{*1} = \frac{r_L S_{t3} + 2S_{t4}}{r_L S_{t1} + 2S_{t2}}$$

$$R_{*2} = \frac{S_{t3}}{S_{t1}}$$

$$E(t_1) = 2(N - r_X)S_{t3} + (N - r_X)^2 S_{t1}^2$$

$$E(t_2) = (N - r_X)((N - r_X) + 2)S_{t3} + 2(N - r_X) \sum_{j_1=2}^{r_M} \sum_{j_2=1}^{j_1-1} \phi_{j_1}^{(E)} \phi_{j_2}^{(E)}$$

The degrees of freedom and noncentrality in the noncentral F approximation of the test statistic are computed as follows:

$$\nu_1^* = \frac{r_L S_{t1}}{R_{*1}}$$

$$\nu_2^* = \frac{(N - r_X)S_{t1}}{R_{*2}}$$

$$\lambda = \frac{S_{t2}}{R_{*1}}$$

Uncorrected Test

The uncorrected test assumes sphericity $\varepsilon = 1$, in which case the null F distribution is exact, with the following degrees of freedom:

$$\nu_1 = r_L r_M$$

$$\nu_2 = r_M(N - r_X)$$

Greenhouse-Geisser Adjustment (MTEST=UNCORR)

The Greenhouse-Geisser adjustment to the uncorrected test reduces degrees of freedom by the MLE $\hat{\varepsilon}$ of the sphericity,

$$\hat{\varepsilon} = \frac{\text{trace}^2(\mathbf{E})}{r_M \text{trace}(\mathbf{E}^2)}$$

An approximation for the expected value of $\hat{\varepsilon}$ is used to compute the degrees of freedom for the null F distribution,

$$\begin{aligned}\nu_1 &= r_L r_M E(\hat{\varepsilon}) \\ \nu_2 &= r_M (N - r_X) E(\hat{\varepsilon})\end{aligned}$$

where

$$E(\hat{\varepsilon}) = \frac{E(t_1)}{r_M E(t_2)}$$

Huynh-Feldt Adjustments (MTEST=HF)

The Huynh-Feldt adjustment reduces degrees of freedom by a nearly unbiased estimate $\tilde{\varepsilon}$ of the sphericity,

$$\tilde{\varepsilon} = \begin{cases} \frac{N r_M \hat{\varepsilon} - 2}{r_M [(N - r_X) - r_M \hat{\varepsilon}]} & \text{if UEPSDEF=HF} \\ \frac{(N - r_X + 1) r_M \hat{\varepsilon} - 2}{r_M [(N - r_X) - r_M \hat{\varepsilon}]} & \text{if UEPSDEF=HFL} \\ \left(\frac{(v_a - 2)(v_a - 4)}{v_a^2} \right) \left(\frac{(N - r_X + 1) r_M \hat{\varepsilon} - 2}{r_M [(N - r_X) - r_M \hat{\varepsilon}]} \right) & \text{if UEPSDEF=CM} \end{cases}$$

where

$$v_a = (N - r_X - 1) + (N - r_X)(N - r_X - 1)/2$$

The value of $\tilde{\varepsilon}$ is truncated if necessary to be at least $1/r_M$ and at most 1.

An approximation for the expected value of $\tilde{\varepsilon}$ is used to compute the degrees of freedom for the null F distribution,

$$\begin{aligned}\nu_1 &= r_L r_M E_t(\tilde{\varepsilon}) \\ \nu_2 &= r_M (N - r_X) E_t(\tilde{\varepsilon})\end{aligned}$$

where

$$E_t(\tilde{\varepsilon}) = \min(\max(E(\tilde{\varepsilon}), 1/r_M), 1)$$

and

$$E(\tilde{\varepsilon}) = \begin{cases} \frac{N E(t_1) - 2 E(t_2)}{r_M [(N - r_X) E(t_2) - E(t_1)]} & \text{if UEPSDEF=HF} \\ \frac{(N - r_X + 1) E(t_1) - 2 E(t_2)}{r_M [(N - r_X) E(t_2) - E(t_1)]} & \text{if UEPSDEF=HFL} \\ \left(\frac{(v_a - 2)(v_a - 4)}{v_a^2} \right) \left(\frac{(N - r_X + 1) E(t_1) - 2 E(t_2)}{r_M [(N - r_X) E(t_2) - E(t_1)]} \right) & \text{if UEPSDEF=CM} \end{cases}$$

Box Conservative Test (MTEST=BOX)

The Box conservative test assumes the worst case for sphericity, $\varepsilon = 1/r_M$, leading to the following degrees of freedom for the null F distribution:

$$\begin{aligned}\nu_1 &= r_L \\ \nu_2 &= (N - r_X)\end{aligned}$$

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 609 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 608 in Chapter 21, “[Statistical Graphics Using ODS](#).”

If ODS Graphics is not enabled, then PROC GLMPOWER creates traditional graphics.

You can reference every graph produced through ODS Graphics with a name. The names of the graphs that PROC GLMPOWER generates are listed in [Table 48.8](#), along with the required statements and options.

Table 48.8 Graphs Produced by PROC GLMPOWER

ODS Graph Name	Plot Description	Option
PowerPlot	Plot with power and sample size on the axes	PLOT
PowerAbort	Empty plot that shows an error message when a plot could not be produced	PLOT

Examples: GLMPOWER Procedure

Example 48.1: One-Way ANOVA

This example deals with the same situation as in [Example 89.1](#) in Chapter 89, “[The POWER 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 race. 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 [Table 48.9](#). You are least familiar with the LZ1 drink and hence decide to consider a range of reasonable values for that mean.

Table 48.9 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 48.10](#) with appropriate contrast coefficients.

Table 48.10 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 48.9](#). 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; that is, you will use a sample size weighting scheme of 2:1:1:1:1.

Before calling PROC GLMPOWER, you need to create the *exemplary data set* to specify means and weights for the design profiles:

```
data Fluids;
  input Fluid $ LacticAcid1 LacticAcid2 CellWgt;
  datalines;
    Water      35.6      35.6      2
    EZD1       33.7      33.7      1
    EZD2       30.2      30.2      1
    LZ1        29       28       1
    LZ2        25.9     25.9      1
  ;
```

The variable LacticAcid1 represents the cell means scenario with the larger LZ1 mean (29), and LacticAcid2 represents the scenario with the smaller LZ1 mean (28). The variable CellWgt contains the sample size allocation weights.

Use the **DATA=** option in the **PROC GLMPOWER** statement to specify Fluids as the exemplary data set. The following statements perform the sample size analysis:


```

proc glmpower data=Fluids;
  class Fluid;
  model LacticAcid1 LacticAcid2 = Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid  -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"          Fluid  1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"       Fluid  1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"         Fluid  0  0  1 -1 0;
  power
    stddev = 3.75
    alpha  = 0.025
    ntotal = .
    power  = 0.9;
run;

```

The **CLASS** statement identifies Fluid as a classification variable. The **MODEL** statement specifies the model and the two cell means scenarios LacticAcid1 and LacticAcid2. The **WEIGHT** statement identifies CellWgt as the weight variable. The **CONTRAST** statement specifies the contrasts. Since PROC GLMPower by default processes class levels in order of formatted values, the contrast coefficients correspond to the following order: EZD1, EZD2, LZ1, LZ2, Water. (**NOTE:** You could use the **ORDER=DATA** option in the **PROC GLMPower** statement to achieve the same ordering as in Table 48.10 instead.) The **POWER** statement specifies total sample size as the result parameter and provides values for the other analysis parameters (error standard deviation, alpha, and power).

Output 48.1.1 displays the results.

Output 48.1.1 Sample Sizes for One-Way ANOVA Contrasts

The GLMPower Procedure

Fixed Scenario Elements							
Weight Variable		CellWgt					
Alpha		0.025					
Error Standard Deviation		3.75					
Nominal Power		0.9					

Computed N Total							
Index	Dependent	Type	Source	Test DF	Error DF	Actual Power	N Total
1	LacticAcid1	Effect	Fluid	4	25	0.958	30
2	LacticAcid1	Contrast	Water vs. others	1	25	0.947	30
3	LacticAcid1	Contrast	EZD vs. LZ	1	55	0.929	60
4	LacticAcid1	Contrast	EZD1 vs. EZD2	1	169	0.901	174
5	LacticAcid1	Contrast	LZ1 vs. LZ2	1	217	0.902	222
6	LacticAcid2	Effect	Fluid	4	25	0.972	30
7	LacticAcid2	Contrast	Water vs. others	1	19	0.901	24
8	LacticAcid2	Contrast	EZD vs. LZ	1	43	0.922	48
9	LacticAcid2	Contrast	EZD1 vs. EZD2	1	169	0.901	174
10	LacticAcid2	Contrast	LZ1 vs. LZ2	1	475	0.902	480

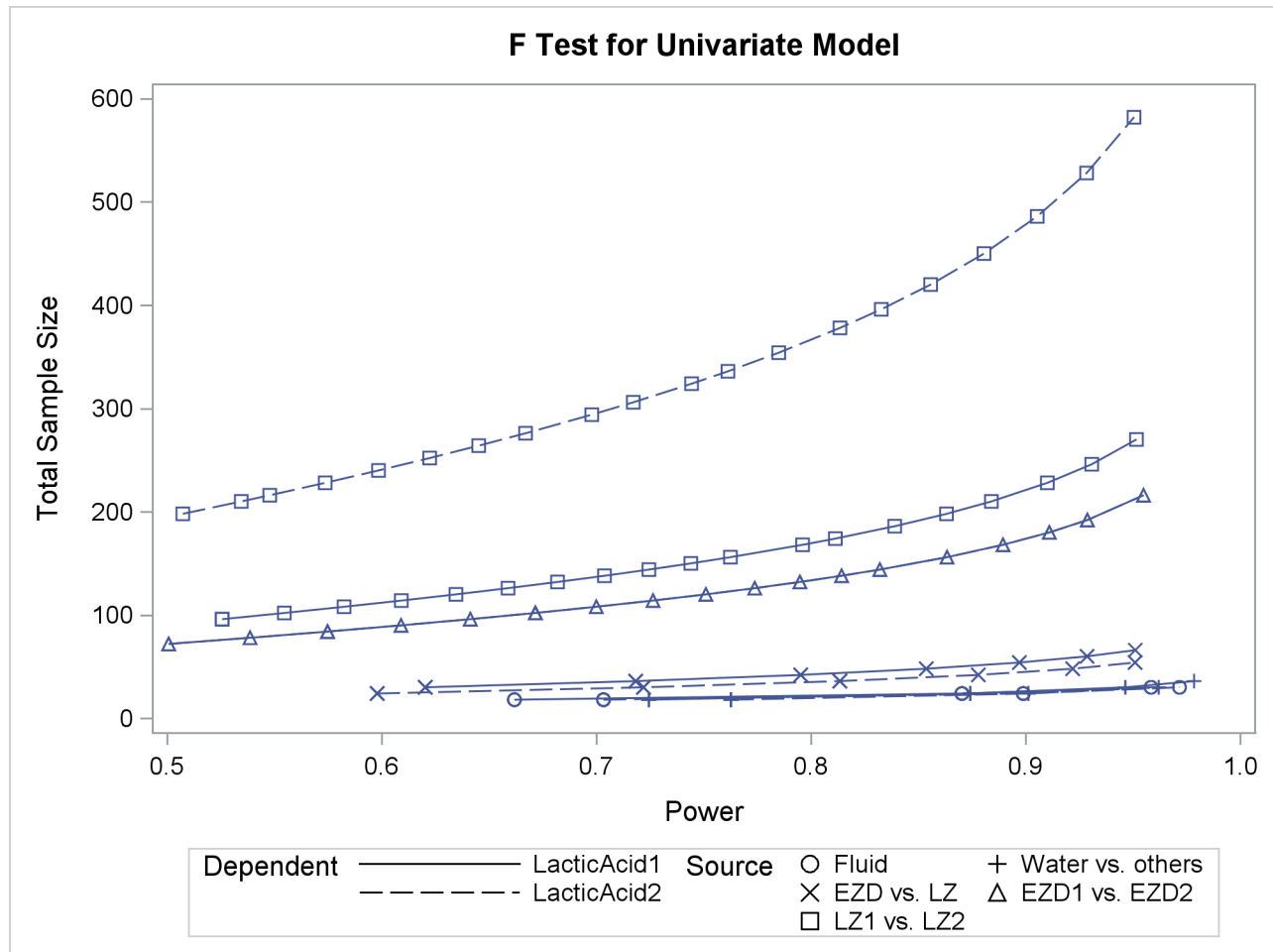
The sample sizes range from 24 for the comparison of water versus electrolytes to 480 for the comparison of LZ1 versus LZ2, both assuming the smaller LZ1 mean. The sample size for the latter comparison is relatively large because the small mean difference of $28 - 25.9 = 2.1$ is hard to detect. PROC GLMPOWER also includes the effect test for Fluid. Note that, in this case, it is equivalent to `TEST=OVERALL_F` in the `ONEWAYANOVA` statement of PROC POWER, since there is only one effect in the model.

The Nominal Power of 0.9 in the “Fixed Scenario Elements” table in [Output 48.1.1](#) 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 `POWER` 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 GLMPOWER` 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:

```
ods graphics on;

proc glmpower data=Fluids plotonly;
  class Fluid;
  model LacticAcid1 LacticAcid2 = Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid  -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"      Fluid   1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"   Fluid   1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"     Fluid    0  0  1 -1 0;
  power
    stddev = 3.75
    alpha  = 0.025
    ntotal = .
    power  = 0.9;
  plot x=power min=.5 max=.95;
run;
```

Output 48.1.2 Plot of Sample Size versus Power for One-Way ANOVA Contrasts

In [Output 48.1.2](#), the line style identifies the cell means scenario, and the plotting symbol identifies the test. The plotting symbol locations identify actual computed powers; the curves are linear interpolations of these points. The plot shows that the required sample size is highest for the test of LZ1 versus LZ2, which was previously found to require the most resources.

Note that some of the plotted points in [Output 48.1.2](#) 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 **POWER** 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 **POWER** statement, identify power as the result (**POWER=.**), and specify any total sample size value (say, **NTOTAL=100**). Specify the **PLOT** statement with **X=N** to request a plot with sample size on the X axis.

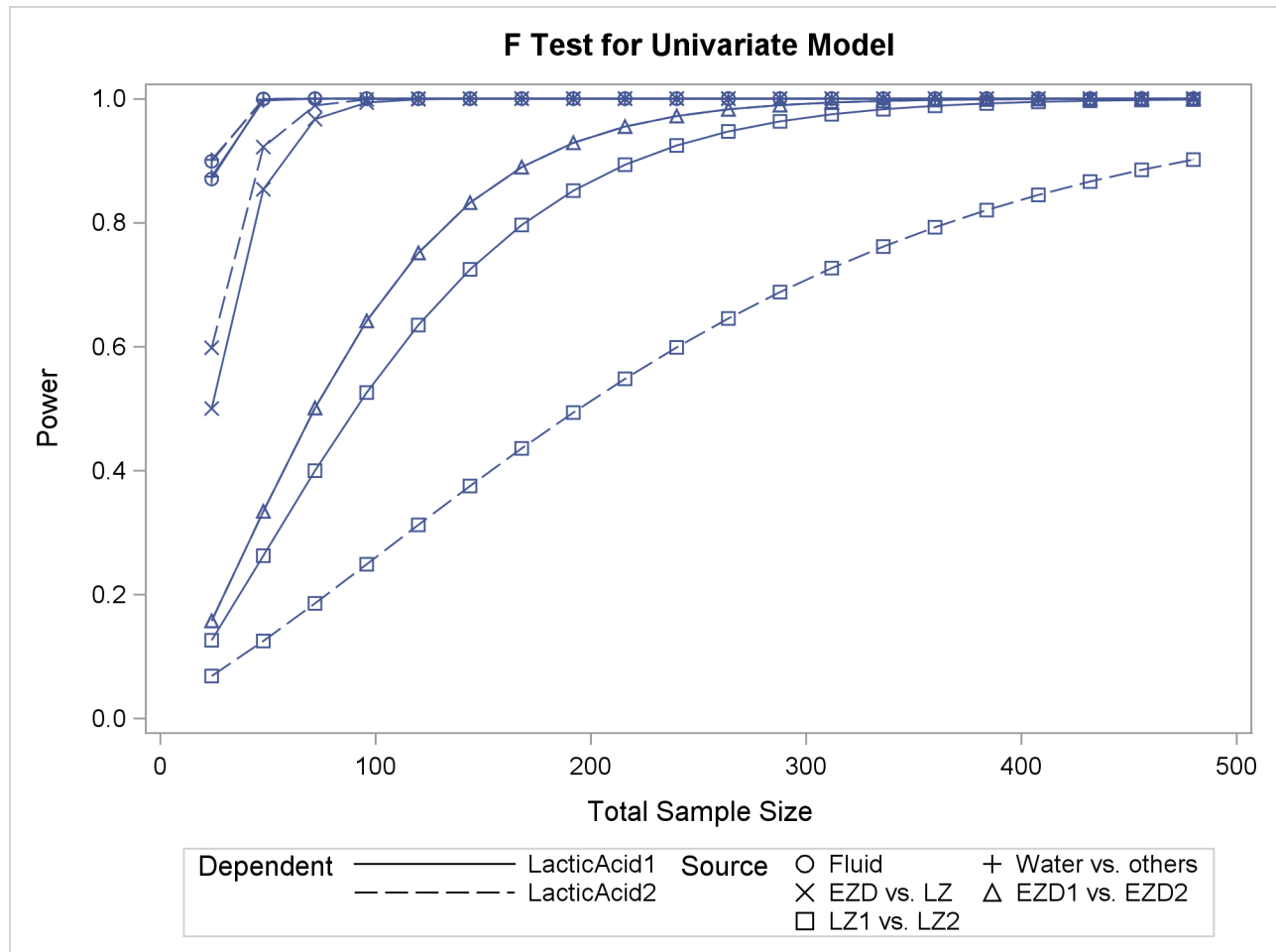
The following statements produce the plot:

```
proc glmpower data=Fluids plotonly;
  class Fluid;
  model LacticAcid1 LacticAcid2 = Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid  -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"          Fluid  1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"       Fluid  1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"         Fluid  0  0  1 -1 0;
  power
    stddev = 3.75
    alpha  = 0.025
    ntotal = 24
    power  = .;
  plot x=n min=24 max=480;
run;

ods graphics off;
```

Note that the value 100 specified with the **NTOTAL=100** 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 GLMPOWER** statement disables nongraphical results. But the **NTOTAL=** option (along with a value) is still needed in the **POWER** statement as a placeholder, to identify the desired parameterization for sample size.

See [Output 48.1.3](#) for the plot.

Output 48.1.3 Plot of Power versus Sample Size for One-Way ANOVA Contrasts

Although [Output 48.1.2](#) and [Output 48.1.3](#) surface essentially the same computations for practical power ranges, they each provide a different quick visual assessment. [Output 48.1.2](#) reveals the range of required sample sizes for powers of interest, and [Output 48.1.3](#) reveals the range of achieved powers for sample sizes of interest.

Example 48.2: Two-Way ANOVA with Covariate

Suppose you can enhance the planned study discussed in [Example 48.1](#) in two ways:

- incorporate results from races at two different altitudes (“high” and “low”)
- measure the body mass index of each runner before the race

This is equivalent to adding a second fixed effect and a continuous covariate to your model.

Since lactic acid buildup is more pronounced at higher altitudes, you will include altitude as a factor in the model along with fluid, extending the one-way ANOVA to a two-way ANOVA. In doing so, you expect to lower the residual standard deviation from about 3.75 to 3.5 (in addition to generalizing the study results). You assume there is negligible interaction between fluid and altitude and plan to use a main-effects-only model. You conjecture that the mean lactic acid buildup follows [Table 48.11](#).

Table 48.11 Mean Lactic Acid Buildup by Fluid and Altitude

Altitude	Water	Fluid			
		EZD1	EZD2	LZ1	LZ2
High	36.9	35.0	31.5	30	27.1
Low	34.3	32.4	28.9	27	24.7

By including a measurement of body mass index as a covariate in the study, you hope to further reduce the error variability. The extent of this reduction in variability is commonly expressed in two alternative ways: (1) the correlation between the covariates and the response or (2) the proportional reduction in total R square incurred by the covariates. You prefer the former and guess that the correlation between body mass index and lactic acid buildup is between 0.2 and 0.3. You specify these estimates with the `NCOVARIATES=` and `CORRXY=` options in the `POWER` statement. The covariate is not included in the `MODEL` statement.

You are interested in the same four fluid comparisons as in [Example 48.1](#), shown in [Table 48.10](#), except this time you want to marginalize over the effect of altitude.

For each of these contrasts, you want to determine the sample size required to achieve a power of 0.9 to detect an effect with magnitude according to [Table 48.11](#). You are not yet attempting to choose a single sample size for the study, but rather checking the range of sample sizes needed by individual contrasts. You plan to test each contrast at $\alpha = 0.025$. You will provide twice as many runners with water as with any of the electrolytes, and you predict that you can study approximately two-thirds as many runners at high altitude than at low altitude. The resulting planned sample size weighting scheme is shown in [Table 48.12](#). Since the scheme is only approximate, you use the `NFRACTIONAL` option in the `POWER` statement to disable the rounding of sample sizes up to integers satisfying the weights exactly.

Table 48.12 Approximate Sample Size Allocation Weights

Altitude	Water	Fluid			
		EZD1	EZD2	LZ1	LZ2
High	4	2	2	2	2
Low	6	3	3	3	3

First, you create the exemplary data set to specify means and weights for the design profiles:

```
data Fluids2;
  input Altitude $ Fluid $ LacticAcid CellWgt;
  datalines;
    High      Water      36.9      4
    High      EZD1       35.0      2
    High      EZD2       31.5      2
    High      LZ1        30        2
    High      LZ2        27.1      2
    Low       Water      34.3      6
    Low       EZD1       32.4      3
    Low       EZD2       28.9      3
    Low       LZ1        27        3
    Low       LZ2        24.7      3
  ;
```

The variables Altitude, Fluid, and LacticAcid specify the factors and cell means in [Table 48.11](#). The variable CellWgt contains the sample size allocation weights in [Table 48.12](#).

Use the **DATA=** option in the **PROC GLMPower** statement to specify Fluids2 as the exemplary data set. The following statements perform the sample size analysis:

```
proc glmpower data=Fluids2;
  class Altitude Fluid;
  model LacticAcid = Altitude Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"      Fluid  1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"   Fluid  1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"     Fluid  0  0  1 -1 0;
  power
    nfractional
    stddev      = 3.5
    ncovariates = 1
    corrxxy     = 0.2 0.3 0
    alpha       = 0.025
    ntotal      = .
    power       = 0.9;
run;
```

The **CLASS** statement identifies Altitude and Fluid as classification variables. The **MODEL** statement specifies the model, and the **WEIGHT** statement identifies CellWgt as the weight variable. The **CONTRAST** statement specifies the contrasts in [Table 48.10](#). As in [Example 48.1](#), the order of the contrast coefficients corresponds to the formatted class levels (EZD1, EZD2, LZ1, LZ2, Water). The **POWER** statement specifies total sample size as the result parameter and provides values for the other analysis parameters. The **NCOVARIATES=** option specifies the single covariate (body mass index), and the **CORRXY=** option specifies the two scenarios for its correlation with lactic acid buildup (0.2 and 0.3). [Output 48.2.1](#) displays the results.

Output 48.2.1 Sample Sizes for Two-Way ANOVA Contrasts**The GLMPOWER Procedure**

Fixed Scenario Elements									
Dependent Variable		LacticAcid							
Weight Variable		CellWgt							
Alpha		0.025							
Number of Covariates		1							
Std Dev Without Covariate Adjustment		3.5							
Nominal Power		0.9							

Computed Ceiling N Total									
Index	Type	Source	Adj				Fractional N Total	Actual Power	Ceiling N Total
			Corr XY	Std Dev	Test DF	Error DF			
1	Effect	Altitude	0.2	3.43	1	84	90.418451	0.902	91
2	Effect	Altitude	0.3	3.34	1	79	85.862649	0.901	86
3	Effect	Altitude	0.0	3.50	1	88	94.063984	0.903	95
4	Effect	Fluid	0.2	3.43	4	16	22.446173	0.912	23
5	Effect	Fluid	0.3	3.34	4	15	21.687544	0.908	22
6	Effect	Fluid	0.0	3.50	4	17	23.055716	0.919	24
7	Contrast	Water vs. others	0.2	3.43	1	15	21.720195	0.905	22
8	Contrast	Water vs. others	0.3	3.34	1	14	20.848805	0.903	21
9	Contrast	Water vs. others	0.0	3.50	1	16	22.422381	0.910	23
10	Contrast	EZD vs. LZ	0.2	3.43	1	35	41.657424	0.903	42
11	Contrast	EZD vs. LZ	0.3	3.34	1	33	39.674037	0.903	40
12	Contrast	EZD vs. LZ	0.0	3.50	1	37	43.246415	0.906	44
13	Contrast	EZD1 vs. EZD2	0.2	3.43	1	139	145.613657	0.901	146
14	Contrast	EZD1 vs. EZD2	0.3	3.34	1	132	138.173983	0.902	139
15	Contrast	EZD1 vs. EZD2	0.0	3.50	1	145	151.565917	0.901	152
16	Contrast	LZ1 vs. LZ2	0.2	3.43	1	268	274.055008	0.901	275
17	Contrast	LZ1 vs. LZ2	0.3	3.34	1	253	259.919126	0.900	260
18	Contrast	LZ1 vs. LZ2	0.0	3.50	1	279	285.363976	0.901	286

The sample sizes in [Output 48.2.1](#) range from 21 for the comparison of water versus electrolytes (assuming a correlation of 0.3 between body mass and lactic acid buildup) to 275 for the comparison of LZ1 versus LZ2 (assuming a correlation of 0.2). PROC GLMPOWER also includes the effect tests for Altitude and Fluid. Note that the required sample sizes for this study are lower than those for the study in [Example 48.1](#).

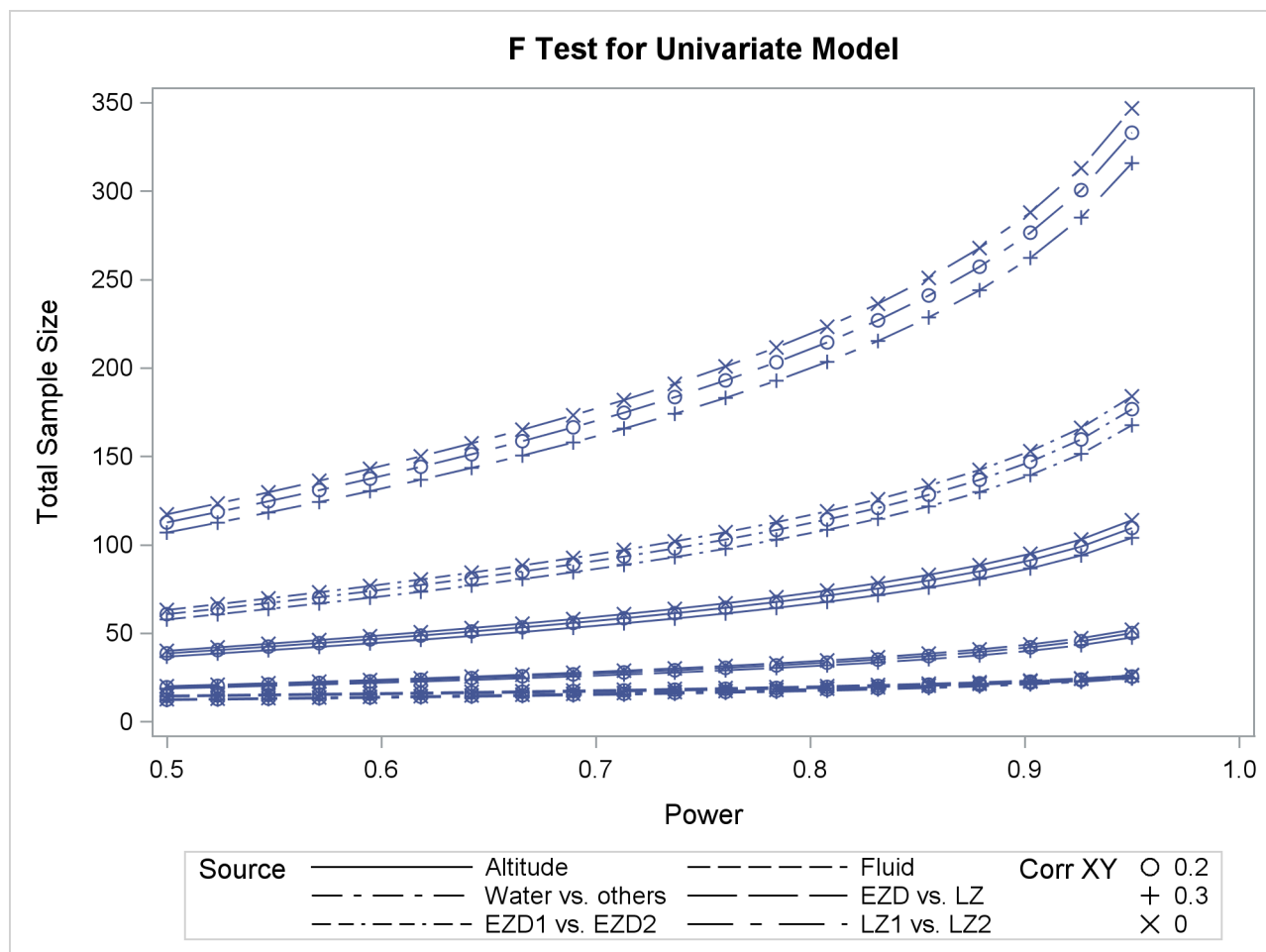
Note that the error standard deviation has been reduced from 3.5 to 3.43 (when correlation is 0.2) or 3.34 (when correlation is 0.3) in the approximation of the effect of the body mass index covariate. The error degrees of freedom has also been automatically adjusted, lowered by 1 (the number of covariates).

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 GLMPower** statement to disable the nongraphical results. Next, specify the **PLOT** statement with **X=POWER** to request a plot with power on the X axis. Sample size is automatically placed on the Y axis. Use the **MIN=** and **MAX=** options in the **PLOT** statement to specify the power range. The following statements produce the plot:

```
ods graphics on;

proc glmpower data=Fluids2 plotonly;
  class Altitude Fluid;
  model LacticAcid = Altitude Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid  -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"          Fluid  1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"       Fluid  1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"         Fluid  0  0  1 -1 0;
  power
    nfractional
    stddev      = 3.5
    ncovariates = 1
    corrxxy     = 0.2 0.3 0
    alpha       = 0.025
    ntotal      = .
    power       = 0.9;
  plot x=power min=.5 max=.95;
run;
```

See [Output 48.2.2](#) for the resulting plot.

Output 48.2.2 Plot of Sample Size versus Power for Two-Way ANOVA Contrasts

In [Output 48.1.2](#), the line style identifies the test, and the plotting symbol identifies the scenario for the correlation between covariate and response. The plotting symbol locations identify actual computed powers; the curves are linear interpolations of these points. As in [Example 48.1](#), the required sample size is highest for the test of LZ1 versus LZ2.

Finally, suppose you want to plot the power for the range of sample sizes you will likely consider for the study (the range of 21 to 275 that achieves 0.9 power for different comparisons). In the **POWER** statement, identify power as the result (**POWER=.**), and specify **NTOTAL=21**. Specify the **PLOT** statement with **X=N** to request a plot with sample size on the X axis.

The following statements produce the plot:

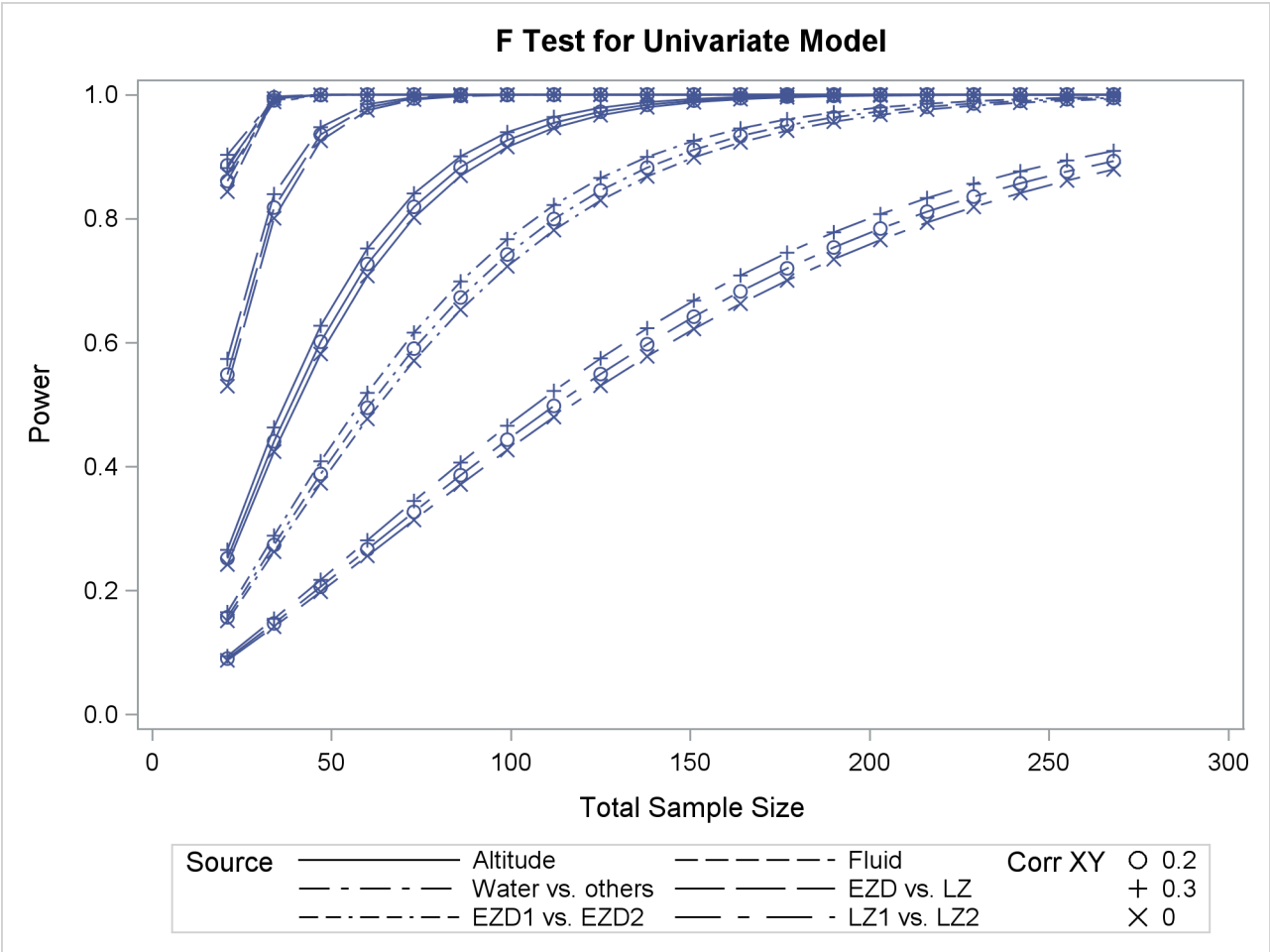
```
proc glmpower data=Fluids2 plotonly;
  class Altitude Fluid;
  model LacticAcid = Altitude Fluid;
  weight CellWgt;
  contrast "Water vs. others" Fluid  -1 -1 -1 -1 4;
  contrast "EZD vs. LZ"          Fluid  1  1 -1 -1 0;
  contrast "EZD1 vs. EZD2"       Fluid  1 -1  0  0 0;
  contrast "LZ1 vs. LZ2"         Fluid  0  0  1 -1 0;
  power
    nfractional
    stddev      = 3.5
    ncovariates = 1
    corrxxy     = 0.2 0.3 0
    alpha       = 0.025
    ntotal      = 21
    power       = .;
  plot x=n min=21 max=275;
run;

ods graphics off;
```

The **MAX=275** option in the **PLOT** statement sets the maximum sample size value. The **MIN=** option automatically defaults to the value of 21 from the **NTOTAL=** option in the **POWER** statement.

See [Output 48.2.3](#) for the plot.

Output 48.2.3 Plot of Power versus Sample Size for Two-Way ANOVA Contrasts



Although [Output 48.2.2](#) and [Output 48.2.3](#) surface essentially the same computations for practical power ranges, they each provide a different quick visual assessment. [Output 48.2.2](#) reveals the range of required sample sizes for powers of interest, and [Output 48.2.3](#) reveals the range of powers achieved for sample sizes of interest.

Example 48.3: Repeated Measures ANOVA

Logan, Baron, and Kohout (1995) and Guo et al. (2013) study the effect of a dental intervention on the memory of pain after root canal therapy. The intervention is a sensory focus strategy, in which patients are instructed to pay attention only to the physical sensations in their mouth during the root canal procedure.

Suppose you are interested in the long-term effects of this sensory focus intervention, because avoidance behavior has been shown to build along with memory of pain. You are planning a study to compare sensory focus to standard of care over a period of a year, asking patients to self-report their memory of pain immediately after the procedure and then again at 1 week, 6 months, and 12 months. You use a scale from 0 (no pain remembered) to 5 (maximum pain remembered).

The between-subject factor in your model is treatment, with two levels (sensory focus versus standard of care), and you allocate each treatment equally for a balanced design. The within-subject factor is time, with four levels (0, 1, 26, and 52 weeks).

You want to determine the number of patients who are needed in order to achieve a power of 0.9 at significance level $\alpha = 0.01$ for the test of the interaction between time and treatment, where the contrast over time contains all pairwise comparisons. You also want to generate a plot of power versus sample size that covers the power range of 0.05 to 0.99.

The default Hotelling-Lawley F test is appropriate for this study, especially because it is the same as the Wald test in PROC MIXED with the DDFM=KR Kenward-Roger degrees-of-freedom method and an unstructured covariance model.

You conjecture that the mean memory of pain for each treatment follows the information in [Table 48.13](#).

Table 48.13 Mean Memory of Pain by Treatment

Treatment	Time Since Root Canal Therapy			
	Later on Same Day	1 Week	6 Months	12 Months
Sensory Focus	2.40	2.38	2.05	1.90
Standard of Care	2.40	2.39	2.36	2.30

The following statements create a data set named Pain that is to contain these means over treatment and time:

```
data Pain;
  input Treatment $ PainMem0 PainMem1Wk PainMem6Mo PainMem12Mo;
  datalines;
    SensoryFocus      2.40  2.38  2.05  1.90
    StandardOfCare    2.40  2.39  2.36  2.30
  ;
```

The variable Treatment specifies the two treatments. The four variables PainMem0, PainMem1Wk, PainMem6Mo, and PainMem12Mo specify the mean memory of pain scores in [Table 48.13](#).

To characterize the variability, you must specify a set of parameters that defines the entire covariance matrix of the residuals. You conjecture that the error standard deviation is the same at all four time points, with a value somewhere between 0.92 and 1.04, and you account for your uncertainty by including both the lower and upper ends of this range in the sample size analysis. You believe that the correlation has a linear exponent autoregressive (LEAR) structure, with a correlation of about 0.6 between measurements one week apart and a decay rate of about 0.8 over one-week intervals. The correlation matrix that contains these LEAR parameters, rounded to three decimal places, is shown in [Table 48.14](#).

Table 48.14 Conjectured Correlation Matrix

	0	1	26	52
0	1	0.6	0.491	0.399
1	0.6	1	0.495	0.402
26	0.491	0.495	1	0.491
52	0.399	0.402	0.491	1

Use the **DATA=** option in the **PROC GLMPOWER** statement to specify **Pain** as the exemplary data set. Specify the between- and within-subject factors and the model by using the **CLASS**, **MODEL**, and **REPEATED** statements just as you would in **PROC GLM** for the repeated measures data analysis. Use the **POWER** statement to indicate sample size as the result parameter and specify the other analysis parameters, and use the **PLOT** statement to generate the power curves. The following statements perform the sample size analysis:

```
ods graphics on;

proc glmpower data=Pain;
  class Treatment;
  model PainMem0 PainMem1Wk PainMem6Mo PainMem12Mo = Treatment;
  repeated Time contrast;
  power
    mtest = hlt
    alpha = 0.01
    power = .9
    ntotal = .
    stddev = 0.92 1.04
    matrix ("PainCorr") = lear(0.6, 0.8, 4, 0 1 26 52)
    corrmat = "PainCorr";
  plot y=power min=0.05 max=0.99 yopts=(ref=0.9)
    vary (linestyle by stddev, symbol by dependent source);
run;
ods graphics off;
```

The **STDDEV=** option specifies the two scenarios for the common residual standard deviation, 0.92 and 1.04. The **MATRIX=** option defines the **LEAR** correlation structure, and the **CORRMAT=** option specifies it as the correlation matrix of the residuals. The **Y=POWER** option in the **PLOT** statement requests a plot that has power on the Y axis. (The result parameter—in this case, total sample size—is always plotted on the other axis.) The **MIN=** and **MAX=** options in the **PLOT** statement specify the power range. The **YOPTS=(REF=)** option adds a reference line at the target power value of 0.9. The **VARY** option specifies that the line style vary by the residual standard deviation and that the plotting symbol vary by the combination of within-subject and between-subject effects. The **ODS GRAPHICS ON** statement enables ODS Graphics.

Output 48.3.1 shows the output, and Output 48.3.2 shows the plot.

Output 48.3.1 Sample Size Analysis for Repeated Measures

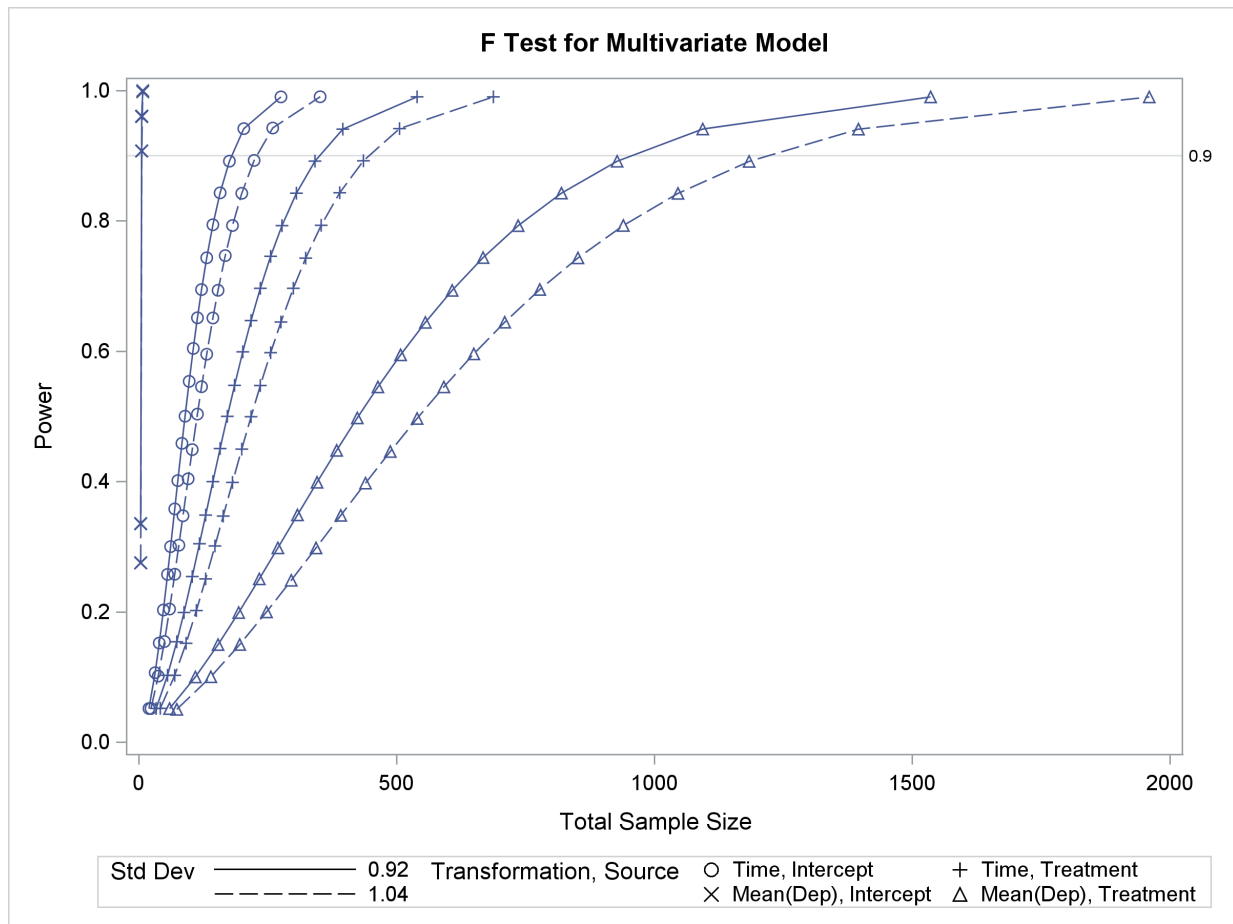
The GLMPOWER Procedure
F Test for Multivariate Model

Fixed Scenario Elements								
Wilks/HLT/PT Method			O'Brien-Shieh					
F Test			Hotelling-Lawley Trace					
Alpha			0.01					
Correlation Matrix			PainCorr					
Nominal Power			0.9					

Computed N Total								
Index	Transformation	Source	Std Dev	Effect	Num DF	Den DF	Actual Power	N Total
1	Time	Intercept	0.92	Time	3	176	0.900	180
2	Time	Intercept	1.04	Time	3	226	0.903	230
3	Time	Treatment	0.92	Time*Treatment	3	346	0.901	350
4	Time	Treatment	1.04	Time*Treatment	3	442	0.901	446
5	Mean(Dep)	Intercept	0.92	Intercept	1	4	0.960	6
6	Mean(Dep)	Intercept	1.04	Intercept	1	4	0.907	6
7	Mean(Dep)	Treatment	0.92	Treatment	1	950	0.900	952
8	Mean(Dep)	Treatment	1.04	Treatment	1	1214	0.900	1216

Output 48.3.1 reveals that the required sample size to achieve a power of 0.9 for the test of the Time*Treatment interaction is 350 for the error standard deviation of 0.92 and 446 for the error standard deviation of 1.04.

Output 48.3.2 Plot of Power versus Sample Size for Repeated Measures Analysis



References

- 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>.
- Chi, Y.-Y., Gribbin, M. J., Lamers, Y., Gregory, J. F., III, and Muller, K. E. (2012). "Global Hypothesis Testing for High-Dimensional Repeated Measures Outcomes." *Statistics in Medicine* 31:724–742.
- Guo, Y., Logan, H. L., Glueck, D. H., and Muller, K. E. (2013). "Selecting a Sample Size for Studies with Repeated Measures." *BMC Medical Research Methodology* 13:100.

- Hocking, R. R. (1985). *The Analysis of Linear Models*. Monterey, CA: Brooks/Cole.
- Huynh, H., and Feldt, L. S. (1976). "Estimation of the Box Correction for Degrees of Freedom from Sample Data in the Randomized Block and Split Plot Designs." *Journal of Educational Statistics* 1:69–82.
- Lecoutre, B. (1991). "A Correction for the Epsilon Approximate Test with Repeated Measures Design with Two or More Independent Groups." *Journal of Educational Statistics* 16:371–372.
- Lenth, R. V. (2001). "Some Practical Guidelines for Effective Sample Size Determination." *American Statistician* 55:187–193.
- Logan, H. L., Baron, R. S., and Kohout, F. (1995). "Sensory Focus as Therapeutic Treatments for Acute Pain." *Psychosomatic Medicine* 57:475–484.
- McKeon, J. J. (1974). " F Approximations to the Distribution of Hotelling's T_0^2 ." *Biometrika* 61:381–383.
- Muller, K. E., and Barton, C. N. (1989). "Approximate Power for Repeated-Measures ANOVA Lacking Sphericity." *Journal of the American Statistical Association* 84:549–555. Also see "Correction to 'Approximate Power for Repeated-Measures ANOVA Lacking Sphericity'," *Journal of the American Statistical Association* 86 (1991): 255–256.
- Muller, K. E., and Benignus, V. A. (1992). "Increasing Scientific Power with Statistical Power." *Neurotoxicology and Teratology* 14:211–219.
- Muller, K. E., Edwards, L. J., Simpson, S. L., and Taylor, D. J. (2007). "Statistical Tests with Accurate Size and Power for Balanced Linear Mixed Models." *Statistics in Medicine* 26:3639–3660.
- Muller, K. E., and Peterson, B. L. (1984). "Practical Methods for Computing Power in Testing the Multivariate General Linear Hypothesis." *Computational Statistics and Data Analysis* 2:143–158.
- 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.
- O'Brien, R. G., and Shieh, G. (1992). "Pragmatic, Unifying Algorithm Gives Power Probabilities for Common F Tests of the Multivariate General Linear Hypothesis." Poster presented at the American Statistical Association Meetings, Statistical Computing Section, Boston.
- Pillai, K. C. S., and Samson, P., Jr. (1959). "On Hotelling's Generalization of T^2 ." *Biometrika* 46:160–168.
- Simpson, S. L., Edwards, L. J., Muller, K. E., Sen, P. K., and Styner, M. A. (2010). "A Linear Exponent AR(1) Family of Correlation Structures." *Statistics in Medicine* 29:1825–1838.

Subject Index

- actual power
 - GLMPOWER procedure, 3767, 3769, 3784
- alpha level
 - GLMPOWER procedure, 3757
- analysis of covariance
 - power and sample size (GLMPOWER), 3788
- analysis of variance
 - multivariate (GLMPOWER), 3749, 3763
 - power and sample size (GLMPOWER), 3739, 3781, 3788
- between-subject factors
 - repeated measures, 3738, 3745, 3748, 3773
- ceiling sample size
 - GLMPOWER procedure, 3769
- classification variables
 - GLMPOWER procedure, 3748, 3751
- contrasts
 - power and sample size (GLMPOWER), 3743, 3748, 3771, 3773, 3781, 3788
 - repeated measures (GLMPOWER), 3765
- correlation
 - GLMPOWER procedure, 3757, 3758, 3763
- covariance
 - GLMPOWER procedure, 3758, 3763
- covariates
 - GLMPOWER procedure, 3751, 3758, 3761–3763, 3772, 3788
- effect
 - GLMPOWER procedure, 3769
- exemplary data set
 - power and sample size (GLMPOWER), 3738, 3740, 3747, 3751, 3766, 3767, 3782
- fractional sample size
 - GLMPOWER procedure, 3769
- GLMPOWER procedure
 - actual alpha, 3769
 - actual power, 3767, 3769, 3784
 - alpha level, 3757
 - analysis of variance, 3739, 3781, 3788
 - ceiling sample size, 3769
 - compared to other procedures, 3739
 - computational methods, 3770
 - contrasts, 3743, 3748, 3765, 3771, 3773, 3781, 3788
 - correlation, 3757, 3758, 3763
 - covariance, 3758, 3763
 - covariates, class and continuous, 3751, 3758, 3761–3763, 3772, 3788
 - displayed output, 3769
 - effect, 3769
 - exemplary data set, 3738, 3740, 3747, 3751, 3766, 3767, 3782
 - fractional sample size, 3769
 - graphics, 3781
 - introductory example, 3739
 - keyword-lists, 3766
 - Kronecker product, 3759
 - linear exponent autoregressive (LEAR), 3759
 - multivariate analysis of variance, 3749, 3763, 3794
 - name-lists, 3766
 - nominal power, 3767, 3769, 3784
 - number-lists, 3766
 - ODS graph names, 3781
 - ODS Graphics, 3781
 - ODS table names, 3770
 - ordering of effects, 3747
 - plots, 3739, 3745, 3747, 3751
 - positional requirements for statements, 3745
 - repeated measures, 3749, 3763, 3794
 - repeated measures examples, 3766
 - sample size adjustment, 3767
 - standard deviation, 3758, 3762, 3763
 - statistical graphics, 3781
 - summary of statements, 3746
 - transformations for MANOVA, 3750
 - transformations for repeated measures, 3765
 - value lists, 3766
 - variance, 3758, 3762
- graphics
 - GLMPOWER procedure, 3781
- keyword-lists
 - GLMPOWER procedure, 3766
- Kronecker product
 - GLMPOWER procedure, 3759
- LEAR, *see* linear exponent autoregressive (LEAR)
- linear exponent autoregressive (LEAR)
 - GLMPOWER procedure, 3759
- MANOVA, *see* multivariate analysis of variance
- multivariate analysis of variance

- GLMPOWER procedure, [3749](#), [3763](#)
 - power and sample size (GLMPOWER), [3794](#)
- name-lists
 - GLMPOWER procedure, [3766](#)
- nominal power
 - GLMPOWER procedure, [3767](#), [3769](#), [3784](#)
- number-lists
 - GLMPOWER procedure, [3766](#)
- ODS graph names
 - GLMPOWER procedure, [3781](#)
- ODS Graphics
 - GLMPOWER procedure, [3781](#)
- orthogonalizing transformation matrix
 - GLMPOWER procedure, [3751](#)
- plots
 - power and sample size (GLMPOWER), [3739](#), [3745](#), [3747](#), [3751](#)
- power
 - See GLMPOWER procedure, [3737](#)
- POWER procedure
 - compared to other procedures, [3739](#)
- prospective power, [3738](#)
- repeated measures
 - contrasts (GLMPOWER), [3765](#)
 - examples (GLMPOWER), [3766](#)
 - GLMPOWER procedure, [3749](#), [3763](#)
 - power and sample size (GLMPOWER), [3794](#)
- retrospective power, [3738](#)
- sample size
 - See GLMPOWER procedure, [3737](#)
- sample size adjustment
 - GLMPOWER procedure, [3767](#)
- standard deviation
 - GLMPOWER procedure, [3758](#), [3762](#), [3763](#)
- statistical graphics
 - GLMPOWER procedure, [3781](#)
- transformation matrix
 - orthogonalizing, [3751](#)
- transformations
 - for multivariate ANOVA, [3750](#)
- transformations for repeated measures
 - GLMPOWER procedure, [3765](#)
- value lists
 - GLMPOWER procedure, [3766](#)
- variance
 - GLMPOWER procedure, [3758](#), [3762](#)
- within-subject factors
 - repeated measures, [3739](#), [3746](#), [3764](#), [3773](#)

Syntax Index

ALPHA= option
POWER statement (GLMPOWER), 3757

BY statement
GLMPOWER procedure, 3747

CLASS statement
GLMPOWER procedure, 3748

CONTRAST option
REPEATED statement (GLMPOWER), 3765

CONTRAST statement
GLMPOWER procedure, 3748

CORRMAT option
POWER statement (GLMPOWER), 3757

CORRS option
POWER statement (GLMPOWER), 3757

CORRXY= option
POWER statement (GLMPOWER), 3758

COVMAT option
POWER statement (GLMPOWER), 3758

DATA= option
PROC GLMPOWER statement, 3747

DEPENDENT option
POWER statement (GLMPOWER), 3758

DESCRIPTION= option
PLOT statement (GLMPOWER), 3755

EFFECTS option
POWER statement (GLMPOWER), 3758

factor specification
REPEATED statement (GLMPOWER), 3764

GLMPOWER procedure
syntax, 3745

GLMPOWER procedure, BY statement, 3747

GLMPOWER procedure, CLASS statement, 3748

GLMPOWER procedure, CONTRAST statement,
3748

SINGULAR= option, 3749

GLMPOWER procedure, MANOVA statement, 3749
M= option, 3750

ORTH option, 3751

GLMPOWER procedure, MODEL statement, 3751

GLMPOWER procedure, PLOT statement, 3751

DESCRIPTION= option, 3755

INTERPOL= option, 3753

KEY= option, 3753

MARKERS= option, 3753

MAX= option, 3754

MIN= option, 3754

NAME= option, 3755

NPOINTS= option, 3754

STEP= option, 3754

VARY option, 3754

X= option, 3754

XOPTS= option, 3755

Y= option, 3755

YOPTS= option, 3755

GLMPOWER procedure, POWER statement, 3756

ALPHA= option, 3757

CORRMAT option, 3757

CORRS option, 3757

CORRXY= option, 3758

COVMAT option, 3758

DEPENDENT option, 3758

EFFECTS option, 3758

MATRIX option, 3758

METHOD= option, 3760

MTEST= option, 3760

NCOVARIATES= option, 3761

NFRACTIONAL option, 3761

NTOTAL= option, 3761

OUTPUTORDER= option, 3762

POWER= option, 3762

PROPVARREDUCTION= option, 3762

SQRTVAR option, 3762

STDDEV= option, 3763

UEPSDEF= option, 3763

GLMPOWER procedure, PROC GLMPOWER
statement, 3746

DATA= option, 3747

ORDER= option, 3747

PLOTONLY= option, 3747

GLMPOWER procedure, REPEATED statement, 3763

CONTRAST option, 3765

factor specification, 3764

HELMERT option, 3765

IDENTITY option, 3765

MEAN option, 3765

POLYNOMIAL option, 3766

PROFILE option, 3766

GLMPOWER procedure, WEIGHT statement, 3766

HELMERT option

REPEATED statement (GLMPOWER), 3765

- IDENTITY option
 - REPEATED statement (GLMPOWER), 3765
- INTERPOL= option
 - PLOT statement (GLMPOWER), 3753
- KEY= option
 - PLOT statement (GLMPOWER), 3753
- M= option
 - MANOVA statement (GLMPOWER), 3750
- MANOVA statement
 - GLMPOWER procedure, 3749
- MARKERS= option
 - PLOT statement (GLMPOWER), 3753
- MATRIX option
 - POWER statement (GLMPOWER), 3758
- MAX= option
 - PLOT statement (GLMPOWER), 3754
- MEAN option
 - REPEATED statement (GLMPOWER), 3765
- METHOD= option
 - POWER statement (GLMPOWER), 3760
- MIN= option
 - PLOT statement (GLMPOWER), 3754
- MODEL statement
 - GLMPOWER procedure, 3751
- MTEST= option
 - POWER statement (GLMPOWER), 3760
- NAME= option
 - PLOT statement (GLMPOWER), 3755
- NCOVARIATES= option
 - POWER statement (GLMPOWER), 3761
- NFRACTIONAL option
 - POWER statement (GLMPOWER), 3761
- NPOINTS= option
 - PLOT statement (GLMPOWER), 3754
- NTOTAL= option
 - POWER statement (GLMPOWER), 3761
- ORDER= option
 - PROC GLMPOWER statement, 3747
- ORTH option
 - MANOVA statement (GLMPOWER), 3751
- OUTPUTORDER= option
 - POWER statement (GLMPOWER), 3762
- PLOT statement
 - GLMPOWER procedure, 3751
- PLOTONLY= option
 - PROC GLMPOWER statement, 3747
- POLYNOMIAL option
 - REPEATED statement (GLMPOWER), 3766
- POWER statement
 - GLMPOWER procedure, 3756
- POWER= option
 - POWER statement (GLMPOWER), 3762
- PROC GLMPOWER statement, *see* GLMPOWER procedure
- PROFILE option
 - REPEATED statement (GLMPOWER), 3766
- PROPVARREDUCTION= option
 - POWER statement (GLMPOWER), 3762
- REPEATED statement
 - GLMPOWER procedure, 3763
- SINGULAR= option
 - CONTRAST statement (GLMPOWER), 3749
- SQRTVAR option
 - POWER statement (GLMPOWER), 3762
- STDDEV= option
 - POWER statement (GLMPOWER), 3763
- STEP= option
 - PLOT statement (GLMPOWER), 3754
- UEPSDEF= option
 - POWER statement (GLMPOWER), 3763
- VARY option
 - PLOT statement (GLMPOWER), 3754
- WEIGHT statement
 - GLMPOWER procedure, 3766
- X= option
 - PLOT statement (GLMPOWER), 3754
- XOPTS= option
 - PLOT statement (GLMPOWER), 3755
- Y= option
 - PLOT statement (GLMPOWER), 3755
- YOPTS= option
 - PLOT statement (GLMPOWER), 3755