

SAS/STAT[®] 14.2 User's Guide

The GLM Procedure

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

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

SAS/STAT® 14.2 User's Guide

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

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

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

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

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

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

SAS Institute Inc., SAS Campus Drive, Cary, NC 27513-2414

November 2016

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

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

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

Chapter 47

The GLM Procedure

Contents

Overview: GLM Procedure	3610
PROC GLM Features	3611
PROC GLM Contrasted with Other SAS Procedures	3612
Getting Started: GLM Procedure	3613
PROC GLM for Unbalanced ANOVA	3613
PROC GLM for Quadratic Least Squares Regression	3616
Syntax: GLM Procedure	3622
PROC GLM Statement	3624
ABSORB Statement	3630
BY Statement	3630
CLASS Statement	3631
CODE Statement	3632
CONTRAST Statement	3633
ESTIMATE Statement	3635
FREQ Statement	3636
ID Statement	3637
LSMEANS Statement	3637
MANOVA Statement	3645
MEANS Statement	3650
MODEL Statement	3655
OUTPUT Statement	3660
RANDOM Statement	3662
REPEATED Statement	3664
STORE Statement	3668
TEST Statement	3668
WEIGHT Statement	3669
Details: GLM Procedure	3670
Statistical Assumptions for Using PROC GLM	3670
Specification of Effects	3670
Using PROC GLM Interactively	3673
Parameterization of PROC GLM Models	3673
Hypothesis Testing in PROC GLM	3678
Effect Size Measures for <i>F</i> Tests in GLM	3683
Absorption	3687
Specification of ESTIMATE Expressions	3690
Comparing Groups	3691

Means versus LS-Means	3691
Multiple Comparisons	3693
Simple Effects	3704
Homogeneity of Variance in One-Way Models	3706
Weighted Means	3707
Construction of Least Squares Means	3707
Multivariate Analysis of Variance	3710
Repeated Measures Analysis of Variance	3712
Random-Effects Analysis	3719
Missing Values	3723
Computational Resources	3723
Computational Method	3726
Output Data Sets	3726
Displayed Output	3728
ODS Table Names	3729
ODS Graphics	3732
Examples: GLM Procedure	3735
Example 47.1: Randomized Complete Blocks with Means Comparisons and Contrasts	3735
Example 47.2: Regression with Mileage Data	3739
Example 47.3: Unbalanced ANOVA for Two-Way Design with Interaction	3742
Example 47.4: Analysis of Covariance	3748
Example 47.5: Three-Way Analysis of Variance with Contrasts	3754
Example 47.6: Multivariate Analysis of Variance	3758
Example 47.7: Repeated Measures Analysis of Variance	3765
Example 47.8: Mixed Model Analysis of Variance with the RANDOM Statement	3769
Example 47.9: Analyzing a Doubly Multivariate Repeated Measures Design	3772
Example 47.10: Testing for Equal Group Variances	3777
Example 47.11: Analysis of a Screening Design	3781
References	3785

Overview: GLM Procedure

The GLM procedure uses the method of least squares to fit general linear models. Among the statistical methods available in PROC GLM are regression, analysis of variance, analysis of covariance, multivariate analysis of variance, and partial correlation.

PROC GLM analyzes data within the framework of general linear models. PROC GLM handles models relating one or several continuous dependent variables to one or several independent variables. The independent variables can be either *classification* variables, which divide the observations into discrete groups, or *continuous* variables. Thus, the GLM procedure can be used for many different analyses, including the following:

- simple regression

- multiple regression
- analysis of variance (ANOVA), especially for unbalanced data
- analysis of covariance
- response surface models
- weighted regression
- polynomial regression
- partial correlation
- multivariate analysis of variance (MANOVA)
- repeated measures analysis of variance

PROC GLM Features

The following list summarizes the features in PROC GLM:

- PROC GLM enables you to specify any degree of interaction (crossed effects) and nested effects. It also provides for polynomial, continuous-by-class, and continuous-nesting-class effects.
- Through the concept of estimability, the GLM procedure can provide tests of hypotheses for the effects of a linear model regardless of the number of missing cells or the extent of confounding. PROC GLM displays the sum of squares (SS) associated with each hypothesis tested and, upon request, the form of the estimable functions employed in the test. PROC GLM can produce the general form of all estimable functions.
- The **REPEATED** statement enables you to specify effects in the model that represent repeated measurements on the same experimental unit for the same response, providing both univariate and multivariate tests of hypotheses.
- The **RANDOM** statement enables you to specify random effects in the model; expected mean squares are produced for each Type I, Type II, Type III, Type IV, and contrast mean square used in the analysis. Upon request, F tests that use appropriate mean squares or linear combinations of mean squares as error terms are performed.
- The **ESTIMATE** statement enables you to specify an $\mathbf{L}$ vector for estimating a linear function of the parameters $\mathbf{L}\boldsymbol{\beta}$.
- The **CONTRAST** statement enables you to specify a contrast vector or matrix for testing the hypothesis that $\mathbf{L}\boldsymbol{\beta} = 0$. When specified, the contrasts are also incorporated into analyses that use the **MANOVA** and **REPEATED** statements.
- The **MANOVA** statement enables you to specify both the hypothesis effects and the error effect to use for a multivariate analysis of variance.

- PROC GLM can create an output data set containing the input data set in addition to predicted values, residuals, and other diagnostic measures.
- PROC GLM can be used interactively. After you specify and fit a model, you can execute a variety of statements without recomputing the model parameters or sums of squares.
- For analysis involving multiple dependent variables but not the [MANOVA](#) or [REPEATED](#) statements, a missing value in one dependent variable does not eliminate the observation from the analysis for other dependent variables. PROC GLM automatically groups together those variables that have the same pattern of missing values within the data set or within a BY group. This ensures that the analysis for each dependent variable brings into use all possible observations.
- The GLM procedure automatically produces graphs as part of its ODS output. For general information about ODS Graphics, see the section “[ODS Graphics](#)” on page 3732 and Chapter 21, “[Statistical Graphics Using ODS](#).”

PROC GLM Contrasted with Other SAS Procedures

As described previously, PROC GLM can be used for many different analyses and has many special features not available in other SAS procedures. However, for some types of analyses, other procedures are available. As discussed in the sections “[PROC GLM for Unbalanced ANOVA](#)” on page 3613 and “[PROC GLM for Quadratic Least Squares Regression](#)” on page 3616, sometimes these other procedures are more efficient than PROC GLM. The following procedures perform some of the same analyses as PROC GLM:

ANOVA	performs analysis of variance for balanced designs. The ANOVA procedure is generally more efficient than PROC GLM for these designs.
MIXED	fits mixed linear models by incorporating covariance structures in the model fitting process. Its RANDOM and REPEATED statements are similar to those in PROC GLM but offer different functionalities.
NESTED	performs analysis of variance and estimates variance components for nested random models. The NESTED procedure is generally more efficient than PROC GLM for these models.
NPAR1WAY	performs nonparametric one-way analysis of rank scores. This can also be done using the RANK procedure and PROC GLM.
REG	performs simple linear regression. The REG procedure allows several MODEL statements and gives additional regression diagnostics, especially for detection of collinearity.
RSREG	performs quadratic response surface regression, and canonical and ridge analysis. The RSREG procedure is generally recommended for data from a response surface experiment.

TTEST	compares the means of two groups of observations. Also, tests for equality of variances for the two groups are available. The TTEST procedure is usually more efficient than PROC GLM for this type of data.
VARCOMP	estimates variance components for a general linear model.

Getting Started: GLM Procedure

PROC GLM for Unbalanced ANOVA

Analysis of variance, or ANOVA, typically refers to partitioning the variation in a variable's values into variation between and within several groups or classes of observations. The GLM procedure can perform simple or complicated ANOVA for balanced or unbalanced data.

This example discusses the analysis of variance for the unbalanced 2×2 data shown in [Table 47.1](#). The experimental design is a full factorial, in which each level of one treatment factor occurs at each level of the other treatment factor. Note that there is only one value for the cell with $A='A2'$ and $B='B2'$. Since one cell contains a different number of values from the other cells in the table, this is an unbalanced design.

Table 47.1 Unbalanced Two-Way Data

	A1	A2
B1	12, 14	20, 18
B2	11, 9	17

The following statements read the data into a SAS data set and then invoke PROC GLM to produce the analysis.

```

title 'Analysis of Unbalanced 2-by-2 Factorial';
data exp;
    input A $ B $ Y @@;
    datalines;
A1 B1 12 A1 B1 14      A1 B2 11 A1 B2 9
A2 B1 20 A2 B1 18      A2 B2 17
;

proc glm data=exp;
    class A B;
    model Y=A B A*B;
run;

```

Both treatments are listed in the **CLASS** statement because they are classification variables. $A*B$ denotes the interaction of the A effect and the B effect. The results are shown in [Figure 47.1](#) and [Figure 47.2](#).

Figure 47.1 Class Level Information
Analysis of Unbalanced 2-by-2 Factorial

The GLM Procedure

Class Level Information		
Class	Levels	Values
A	2	A1 A2
B	2	B1 B2

Number of Observations Read	7
Number of Observations Used	7

Figure 47.1 displays information about the classes as well as the number of observations in the data set. Figure 47.2 shows the ANOVA table, simple statistics, and tests of effects.

Figure 47.2 ANOVA Table and Tests of Effects
Analysis of Unbalanced 2-by-2 Factorial

The GLM Procedure

Dependent Variable: Y

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	91.71428571	30.57142857	15.29	0.0253
Error	3	6.00000000	2.00000000		
Corrected Total	6	97.71428571			

R-Square	Coeff Var	Root MSE	Y Mean
0.938596	9.801480	1.414214	14.42857

Source	DF	Type I SS	Mean Square	F Value	Pr > F
A	1	80.04761905	80.04761905	40.02	0.0080
B	1	11.26666667	11.26666667	5.63	0.0982
A*B	1	0.40000000	0.40000000	0.20	0.6850

Source	DF	Type III SS	Mean Square	F Value	Pr > F
A	1	67.60000000	67.60000000	33.80	0.0101
B	1	10.00000000	10.00000000	5.00	0.1114
A*B	1	0.40000000	0.40000000	0.20	0.6850

The degrees of freedom can be used to check your data. The Model degrees of freedom for a 2×2 factorial design with interaction are $(ab - 1)$, where a is the number of levels of A and b is the number of levels of B; in this case, $(2 \times 2 - 1) = 3$. The Corrected Total degrees of freedom are always one less than the number of observations used in the analysis; in this case, $7 - 1 = 6$.

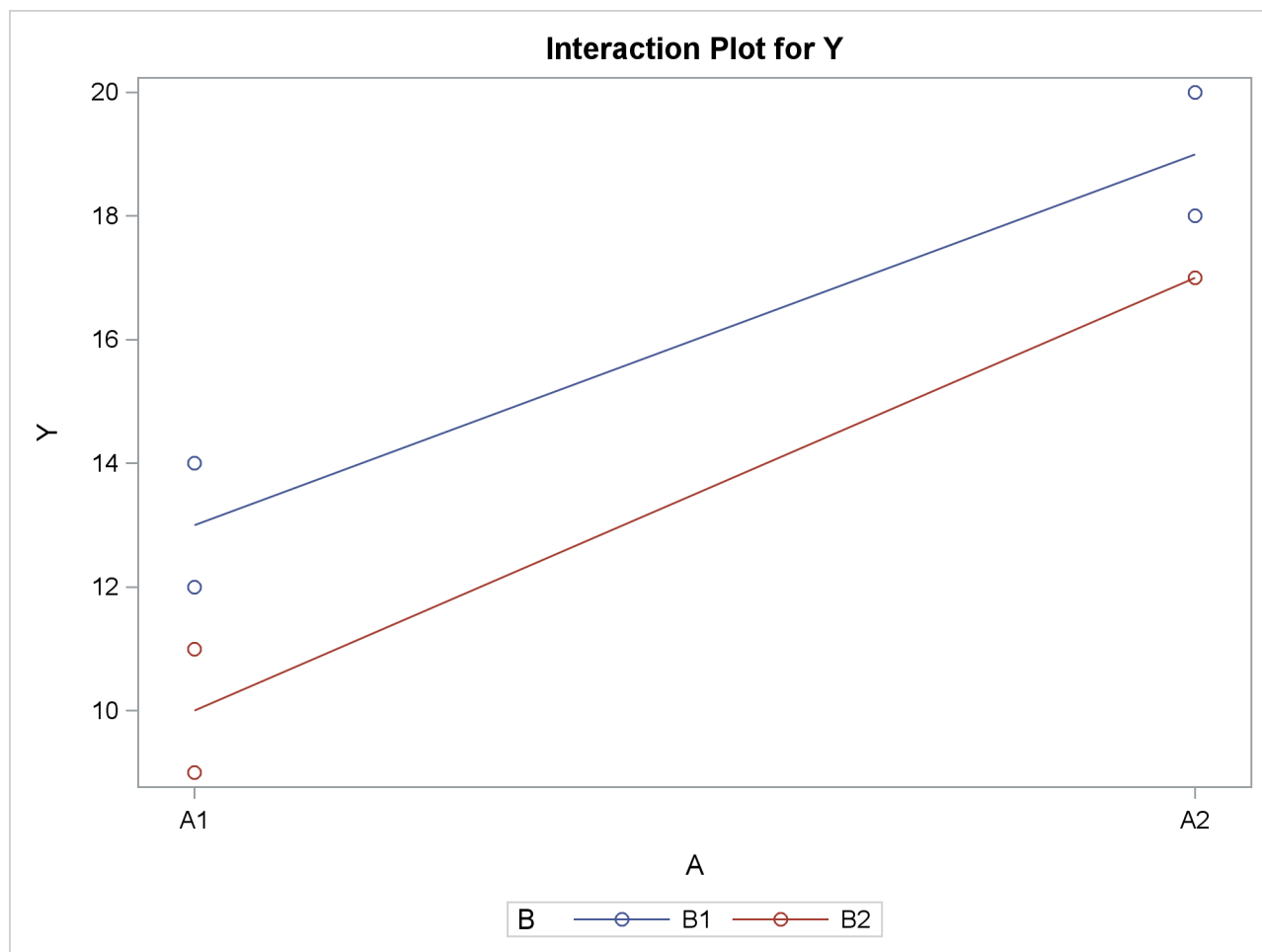
The overall F test is significant ($F = 15.29$, $p = 0.0253$), indicating strong evidence that the means for the four different A×B cells are different. You can further analyze this difference by examining the individual tests for each effect.

Four types of estimable functions of parameters are available for testing hypotheses in PROC GLM. For data with no missing cells, the Type III and Type IV estimable functions are the same and test the same hypotheses that would be tested if the data were balanced. Type I and Type III sums of squares are typically not equal when the data are unbalanced; Type III sums of squares are preferred in testing effects in unbalanced cases because they test a function of the underlying parameters that is independent of the number of observations per treatment combination.

According to a significance level of 5% ($\alpha = 0.05$), the A*B interaction is not significant ($F = 0.20$, $p = 0.6850$). This indicates that the effect of A does not depend on the level of B and vice versa. Therefore, the tests for the individual effects are valid, showing a significant A effect ($F = 33.80$, $p = 0.0101$) but no significant B effect ($F = 5.00$, $p = 0.1114$).

If ODS Graphics is enabled, GLM also displays by default an interaction plot for this analysis. The following statements, which are the same as in the previous analysis but with ODS Graphics enabled, additionally produce [Figure 47.3](#).

```
ods graphics on;
proc glm data=exp;
  class A B;
  model Y=A B A*B;
run;
ods graphics off;
```

Figure 47.3 Plot of Y by A and B

The insignificance of the A*B interaction is reflected in the fact that two lines in Figure 47.3 are nearly parallel. For more information about the graphics that GLM can produce, see the section “ODS Graphics” on page 3732.

PROC GLM for Quadratic Least Squares Regression

In polynomial regression, the values of a dependent variable (also called a response variable) are described or predicted in terms of polynomial terms involving one or more independent or explanatory variables. An example of quadratic regression in PROC GLM follows. These data are taken from Draper and Smith (1966, p. 57). Thirteen specimens of 90/10 Cu-Ni alloys are tested in a corrosion-wheel setup in order to examine corrosion. Each specimen has a certain iron content. The wheel is rotated in salt seawater at 30 ft/sec for 60 days. Weight loss is used to quantify the corrosion. The *fe* variable represents the iron content, and the *loss* variable denotes the weight loss in milligrams/square decimeter/day in the following DATA step.

```
title 'Regression in PROC GLM';
data iron;
  input fe loss @@;
```

```

    datalines;
0.01 127.6   0.48 124.0   0.71 110.8   0.95 103.9
1.19 101.5   0.01 130.1   0.48 122.0   1.44  92.3
0.71 113.1   1.96  83.7   0.01 128.0   1.44  91.4
1.96  86.2
;

```

The SGSCATTER procedure is used in the following statements to request a scatter plot of the response variable versus the independent variable.

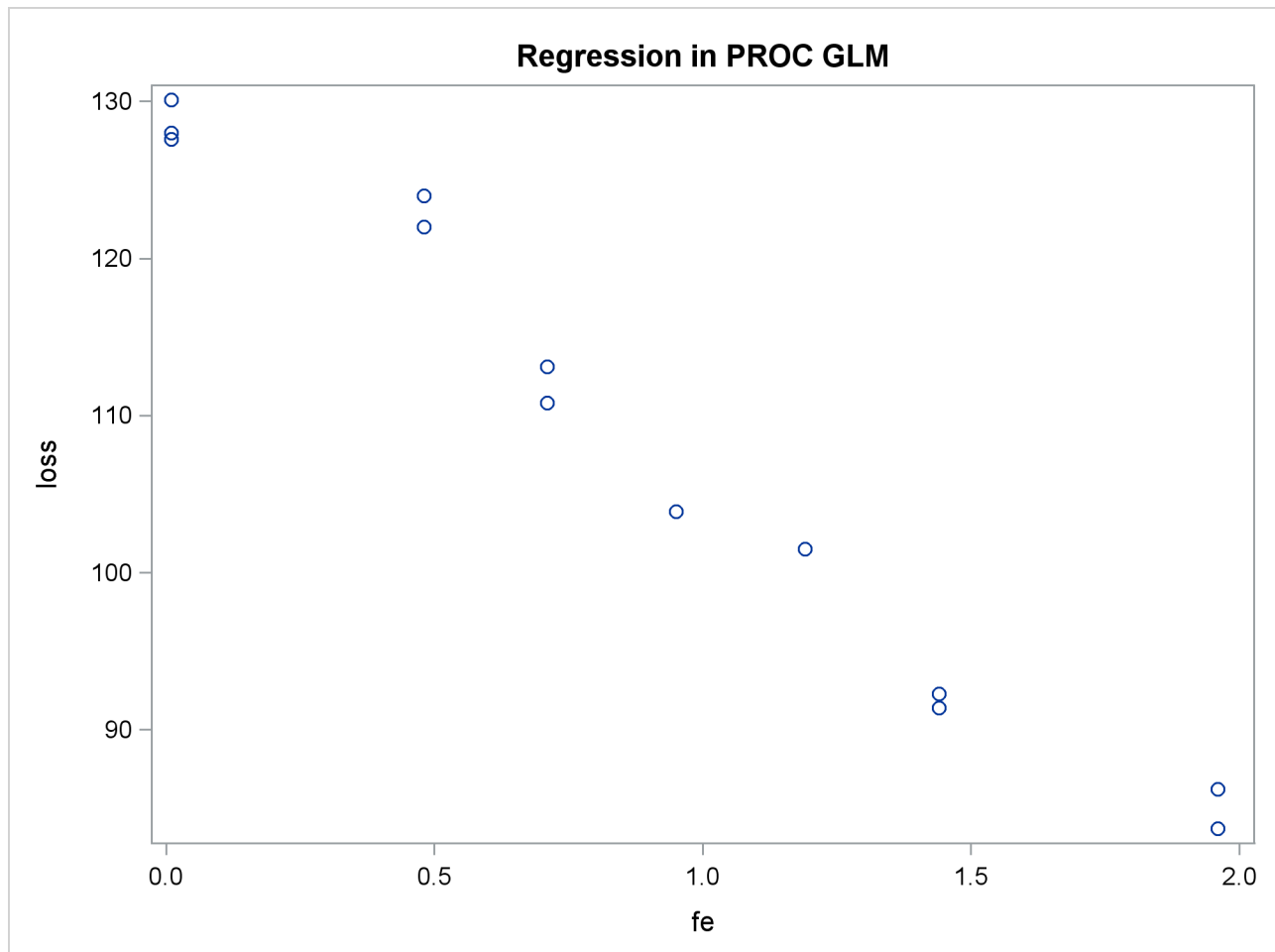
```

proc sgscatter data=iron;
  plot loss*fe;
run;
ods graphics off;

```

The plot in [Figure 47.4](#) displays a strong negative relationship between iron content and corrosion resistance, but it is not clear whether there is curvature in this relationship.

Figure 47.4 Plot of Observed Corrosion Resistance by Iron Content



The following statements fit a quadratic regression model to the data. This enables you to estimate the linear relationship between iron content and corrosion resistance and to test for the presence of a quadratic component. The intercept is automatically fit unless the **NOINT** option is specified.

```
proc glm data=iron;
  model loss=fe fe*fe;
run;
```

The **CLASS** statement is omitted because a regression line is being fitted. Unlike PROC REG, PROC GLM allows polynomial terms in the **MODEL** statement.

PROC GLM first displays preliminary information, shown in Figure 47.5, telling you that the GLM procedure has been invoked and stating the number of observations in the data set. If the model involves classification variables, they are also listed here, along with their levels.

Figure 47.5 Data Information

Regression in PROC GLM

The GLM Procedure

Number of Observations Read	13
Number of Observations Used	13

Figure 47.6 shows the overall ANOVA table and some simple statistics. The degrees of freedom can be used to check that the model is correct and that the data have been read correctly. The Model degrees of freedom for a regression is the number of parameters in the model minus 1. You are fitting a model with three parameters in this case,

$$\text{loss} = \beta_0 + \beta_1 \times (\text{fe}) + \beta_2 \times (\text{fe})^2 + \text{error}$$

so the degrees of freedom are $3 - 1 = 2$. The Corrected Total degrees of freedom are always one less than the number of observations used in the analysis.

Figure 47.6 ANOVA Table

Regression in PROC GLM

The GLM Procedure

Dependent Variable: loss

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	2	3296.530589	1648.265295	164.68	<.0001
Error	10	100.086334	10.008633		
Corrected Total	12	3396.616923			

R-Square	Coeff Var	Root MSE	loss Mean
0.970534	2.907348	3.163642	108.8154

The R square indicates that the model accounts for 97% of the variation in LOSS. The coefficient of variation (Coeff Var), Root MSE (Mean Square for Error), and mean of the dependent variable are also listed.

The overall F test is significant ($F = 164.68$, $p < 0.0001$), indicating that the model as a whole accounts for a significant amount of the variation in LOSS. Thus, it is appropriate to proceed to testing the effects.

Figure 47.7 contains tests of effects and parameter estimates. The latter are displayed by default when the model contains only continuous variables.

Figure 47.7 Tests of Effects and Parameter Estimates

Source	DF	Type I SS	Mean Square	F Value	Pr > F
fe	1	3293.766690	3293.766690	329.09	<.0001
fe*fe	1	2.763899	2.763899	0.28	0.6107

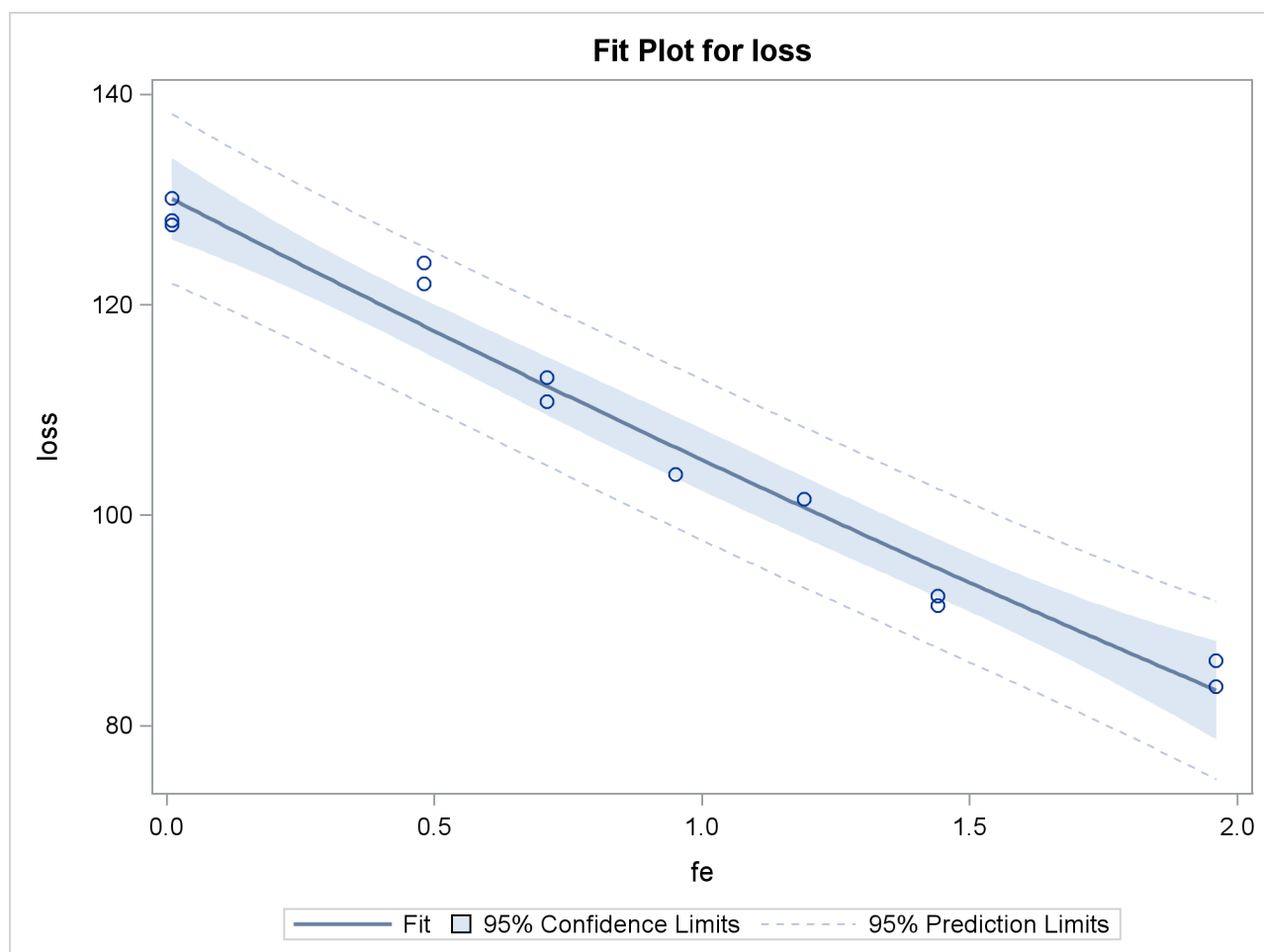
Source	DF	Type III SS	Mean Square	F Value	Pr > F
fe	1	356.7572421	356.7572421	35.64	0.0001
fe*fe	1	2.7638994	2.7638994	0.28	0.6107

Parameter	Estimate	Standard Error	t Value	Pr > t
Intercept	130.3199337	1.77096213	73.59	<.0001
fe	-26.2203900	4.39177557	-5.97	0.0001
fe*fe	1.1552018	2.19828568	0.53	0.6107

The t tests provided are equivalent to the Type III F tests. The quadratic term is not significant ($p = 0.6107$) and thus can be removed from the model; the linear term is significant ($p < 0.0001$). This suggests that there is indeed a straight-line relationship between loss and fe.

Finally, if ODS Graphics is enabled, PROC GLM also displays by default a scatter plot of the original data, as in Figure 47.4, with the quadratic fit overlaid. The following statements, which are the same as the previous analysis but with ODS Graphics enabled, additionally produce Figure 47.8.

```
ods graphics on;
proc glm data=iron;
    model loss=fe fe*fe;
run;
ods graphics off;
```

Figure 47.8 Plot of Observed and Fit Corrosion Resistance by Iron Content, Quadratic Model

The insignificance of the quadratic term in the model is reflected in the fact that the fit is nearly linear.

Fitting the model without the quadratic term provides more accurate estimates for β_0 and β_1 . PROC GLM allows only one **MODEL** statement per invocation of the procedure, so the **PROC GLM** statement must be issued again. The following statements are used to fit the linear model.

```
proc glm data=iron;
  model loss=fe;
run;
```

Figure 47.9 displays the output produced by these statements. The linear term is still significant ($F = 352.27$, $p < 0.0001$). The estimated model is now

$$\text{loss} = 129.79 - 24.02 \times \text{fe}$$

Figure 47.9 Linear Model Output

Regression in PROC GLM

The GLM Procedure

Dependent Variable: loss

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	1	3293.766690	3293.766690	352.27	<.0001
Error	11	102.850233	9.350021		
Corrected Total	12	3396.616923			

R-Square	Coeff Var	Root MSE	loss Mean
0.969720	2.810063	3.057780	108.8154

Source	DF	Type I SS	Mean Square	F Value	Pr > F
fe	1	3293.766690	3293.766690	352.27	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
fe	1	3293.766690	3293.766690	352.27	<.0001

Parameter	Estimate	Standard Error	t Value	Pr > t
Intercept	129.7865993	1.40273671	92.52	<.0001
fe	-24.0198934	1.27976715	-18.77	<.0001

Syntax: GLM Procedure

The following statements are available in the GLM procedure:

```

PROC GLM < options > ;
    CLASS variable < (REF= option) > ... < variable < (REF= option) > > < / global-options > ;
    MODEL dependent-variables = independent-effects < / options > ;
    ABSORB variables ;
    BY variables ;
    CODE < options > ;
    FREQ variable ;
    ID variables ;
    WEIGHT variable ;
    CONTRAST 'label' effect values < ... effect values > < / options > ;
    ESTIMATE 'label' effect values < ... effect values > < / options > ;
    LSMEANS effects < / options > ;
    MANOVA < test-options > < / detail-options > ;
    MEANS effects < / options > ;
    OUTPUT < OUT=SAS-data-set> keyword=names < ... keyword=names > < / option > ;
    RANDOM effects < / options > ;
    REPEATED factor-specification < / options > ;
    STORE < OUT= > item-store-name < / LABEL='label' > ;
    TEST < H=effects > E=effect < / options > ;

```

Although there are numerous statements and options available in PROC GLM, many applications use only a few of them. Often you can find the features you need by looking at an example or by quickly scanning through this section.

To use PROC GLM, the **PROC GLM** and **MODEL** statements are required. You can specify only one **MODEL** statement (in contrast to the REG procedure, for example, which allows several **MODEL** statements in the same PROC REG run). 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, if you use a **CONTRAST** statement in combination with a **MANOVA**, **RANDOM**, **REPEATED**, or **TEST** statement, the **CONTRAST** statement must be entered first in order for the contrast to be included in the **MANOVA**, **RANDOM**, **REPEATED**, or **TEST** analysis.

Table 47.2 summarizes the positional requirements for the statements in the GLM procedure.

Table 47.2 Positional Requirements for PROC GLM Statements

Statement	Must Precede...	Must Follow...
ABSORB	First RUN statement	
BY	First RUN statement	
CLASS	MODEL statement	
CONTRAST	MANOVA, REPEATED, or RANDOM statement	MODEL statement
ESTIMATE		MODEL statement
FREQ	First RUN statement	
ID	First RUN statement	
LSMEANS		MODEL statement
MANOVA		CONTRAST or MODEL statement
MEANS		MODEL statement
MODEL	CONTRAST, ESTIMATE, LSMEANS, or MEANS statement	CLASS statement
OUTPUT		MODEL statement
RANDOM		CONTRAST or MODEL statement
REPEATED		CONTRAST, MODEL, or TEST statement
TEST	MANOVA or REPEATED statement	MODEL statement
WEIGHT	First RUN statement	

Table 47.3 summarizes the function of each statement (other than the PROC statement) in the GLM procedure.

Table 47.3 Statements in the GLM Procedure

Statement	Description
ABSORB	Absorbs classification effects in a model
BY	Specifies variables to define subgroups for the analysis
CLASS	Declares classification variables
CODE	Requests that the procedure write SAS DATA step code to a file or catalog entry for computing predicted values according to the fitted model
CONTRAST	Constructs and tests linear functions of the parameters
ESTIMATE	Estimates linear functions of the parameters
FREQ	Specifies a frequency variable
ID	Identifies observations on output
LSMEANS	Computes least squares (marginal) means
MANOVA	Performs a multivariate analysis of variance

Table 47.3 *continued*

Statement	Description
MEANS	Computes and optionally compares arithmetic means
MODEL	Defines the model to be fit
OUTPUT	Requests an output data set containing diagnostics for each observation
RANDOM	Declares certain effects to be random and computes expected mean squares
REPEATED	Performs multivariate and univariate repeated measures analysis of variance
STORE	Requests that the procedure save the context and results of the statistical analysis into an item store
TEST	Constructs tests that use the sums of squares for effects and the error term you specify
WEIGHT	Specifies a variable for weighting observations

The rest of this section provides detailed syntax information for each of these statements, beginning with the **PROC GLM** statement. The remaining statements are covered in alphabetical order.

The **STORE** and **CODE** statements are also used by many other procedures. A summary description of functionality and syntax for these statements is also shown after the **PROC GLM** statement in alphabetical order, but you can find full documentation about them in the section “**STORE Statement**” on page 509 in Chapter 19, “**Shared Concepts and Topics**.”

PROC GLM Statement

PROC GLM <options> ;

The **PROC GLM** statement invokes the GLM procedure. [Table 47.4](#) summarizes the *options* available in the **PROC GLM** statement.

Table 47.4 PROC GLM Statement Options

Option	Description
ALPHA=	Specifies the level of significance for confidence intervals
DATA=	Names the SAS data set used by the GLM procedure
MANOVA	Requests the multivariate mode of eliminating observations with missing values
MULTIPASS	Requests that the input data set be reread when necessary, instead of using a utility file
NAMELEN=	Specifies the length of effect names
NOPRINT	Suppresses the normal display of results
ORDER=	Specifies the order in which to sort classification variables
OUTSTAT=	Names an output data set for information and statistics on each model effect
PLOTS	Controls the plots produced through ODS Graphics

You can specify the following *options* in the PROC GLM statement.

ALPHA=*p*

specifies the level of significance p for $100(1 - p)\%$ confidence intervals. The value must be between 0 and 1; the default value of $p = 0.05$ results in 95% intervals. This value is used as the default confidence level for limits computed by the following *options*.

Statement	Options
LSMEANS	CL
MEANS	CLM CLDIFF
MODEL	CLI CLM CLPARM
OUTPUT	UCL= LCL= UCLM= LCLM=

You can override the default in each of these cases by specifying the ALPHA= option for each statement individually.

DATA=SAS-data-set

names the SAS data set used by the GLM procedure. By default, PROC GLM uses the most recently created SAS data set.

MANOVA

requests the multivariate mode of eliminating observations with missing values. If any of the dependent variables have missing values, the procedure eliminates that observation from the analysis. The MANOVA option is useful if you use PROC GLM in interactive mode and plan to perform a multivariate analysis.

MULTIPASS

requests that PROC GLM reread the input data set when necessary, instead of writing the necessary values of dependent variables to a utility file. This option decreases disk space usage at the expense of increased execution times, and is useful only in rare situations where disk space is at an absolute premium.

NAMELEN=*n*

specifies the length of effect names in tables and output data sets to be n characters long, where n is a value between 20 and 200 characters. The default length is 20 characters.

NOPRINT

suppresses the normal display of results. The NOPRINT option is useful when you want only to create one or more output data sets with the procedure. Note that this option temporarily disables the Output Delivery System (ODS); see Chapter 20, “[Using the Output Delivery System](#),” for more information.

ORDER=DATA | FORMATTED | FREQ | INTERNAL

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

This ordering determines which parameters in the model correspond to each level in the data, so the ORDER= option can be useful when you specify the [CONTRAST](#) or [ESTIMATE](#) 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 ORDER=	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*.

OUTSTAT=SAS-data-set

names an output data set that contains sums of squares, degrees of freedom, F statistics, and probability levels for each effect in the model, as well as for each [CONTRAST](#) that uses the overall residual or error mean square (MSE) as the denominator in constructing the F statistic. If you use the [CANONICAL](#) option in the [MANOVA](#) statement and do not use an [M=](#) specification in the [MANOVA](#) statement, the data set also contains results of the canonical analysis.

See the section “[Output Data Sets](#)” on page 3726 for more information.

PLOTS < (global-plot-options) > < =plot-request < (options) > >

PLOTS < (global-plot-options) > < =(plot-request < (options) > < ... plot-request < (options) > > >

controls the plots produced through ODS Graphics. When you specify only one *plot-request*, you can omit the parentheses from around the *plot-request*. For example:

```
PLOTS=NONE
PLOTS=(DIAGNOSTICS RESIDUALS)
PLOTS(UNPACK)=RESIDUALS
PLOT=MEANPLOT(CLBAND)
```

ODS Graphics must be enabled before plots can be requested. For example:

```
ods graphics on;
proc glm data=iron;
    model loss=fe fe*fe;
run;
ods graphics off;
```

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

If ODS Graphics is enabled but you do not specify the PLOTS= option, then PROC GLM produces a default set of plots, which might be different for different models, as discussed in the following.

- If you specify a one-way analysis of variance model, with just one **CLASS** variable, the GLM procedure produces a grouped box plot of the response values versus the **CLASS** levels. For an example of the box plot, see the section “[One-Way Layout with Means Comparisons](#)” on page 972 in Chapter 26, “[The ANOVA Procedure](#).”
- If you specify a two-way analysis of variance model, with just two **CLASS** variables, the GLM procedure produces an interaction plot of the response values, with horizontal position representing one **CLASS** variable and marker style representing the other; and with predicted response values connected by lines representing the two-way analysis. For an example of the interaction plot, see the section “[PROC GLM for Unbalanced ANOVA](#)” on page 3613.
- If you specify a model with a single continuous predictor, the GLM procedure produces a fit plot of the response values versus the covariate values, with a curve representing the fitted relationship and a band representing the confidence limits for individual mean values. For an example of the fit plot, see the section “[PROC GLM for Quadratic Least Squares Regression](#)” on page 3616.
- If you specify a model with two continuous predictors and no **CLASS** variables, the GLM procedure produces a contour fit plot, overlaying a scatter plot of the data and a contour plot of the predicted surface.
- If you specify an analysis of covariance model, with one or two **CLASS** variables and one continuous variable, the GLM procedure produces an analysis of covariance plot of the response values versus the covariate values, with lines representing the fitted relationship within each classification level. For an example of the analysis of covariance plot, see [Example 47.4](#).
- If you specify an **LSMEANS** statement with the **PDIF** option, the GLM procedure produces a plot appropriate for the type of LS-means comparison. For **PDIF**=ALL (which is the default if you specify only **PDIF**), the procedure produces a diffogram, which displays all pairwise LS-means differences and their significance. The display is also known as a “mean-mean scatter plot” (Hsu 1996). For **PDIF**=CONTROL, the procedure produces a display of each noncontrol LS-mean compared to the control LS-mean, with two-sided confidence intervals for the comparison. For **PDIF**=CONTROLL and **PDIF**=CONTROLLU a similar display is produced, but with one-sided confidence intervals. Finally, for the **PDIF**=ANOM option, the procedure produces an “analysis of means” plot, comparing each LS-mean to the average LS-mean.
- If you specify a **MEANS** statement, the GLM procedure produces a grouped box plot of the response values versus the effect for which means are being calculated.

The *global-plot-options* include the following:

MAXPOINTS=NONE | *number*

specifies that plots with elements that require processing of more than *number* points be suppressed. The default is MAXPOINTS=5000. This limit is ignored if you specify MAXPOINTS=NONE.

ONLY

suppresses the default plots. Only plots specifically requested are displayed.

UNPACKPANEL

UNPACK

suppresses paneling. By default, multiple plots can appear in some output panels. Specify UNPACKPANEL to get each plot in a separate panel. You can specify PLOTS(UNPACKPANEL) to just unpack the default plots. You can also specify UNPACKPANEL as a suboption with DIAGNOSTICS and RESIDUALS.

The following individual plots and plot options are available. If you specify only one *plot*, then you can omit the parentheses.

ALL

produces all appropriate plots. You can specify other options with ALL; for example, to request all plots and unpack just the residuals, specify: PLOTS=(ALL RESIDUALS(UNPACK)).

ANCOVAPLOT<(CLM CLI LIMITS)>

modifies the analysis of covariance plot produced by default when you have an analysis of covariance model, with one or two **CLASS** variables and one continuous variable. By default the plot does not show confidence limits around the predicted values. The PLOTS=ANCOVAPLOT(CLM) option adds limits for the expected predicted values, and PLOTS=ANCOVAPLOT(CLI) adds limits for new predictions. Use PLOTS=ANCOVAPLOT(LIMITS) to add both kinds of limits.

ANOMPLOT

requests an analysis of means display, in which least squares means are compared against an average least squares mean (Ott 1967; Nelson 1982, 1991, 1993). LS-mean ANOM plots are produced only if you also specify **PDIFF=ANOM** or **ADJUST=NELSON** in the **LSMEANS** statement, and in this case they are produced by default.

BOXPLOT<(NPANELPOS=*n*)>

modifies the plot produced by default for the model effect in a one-way analysis of variance model, or for an effect specified in the **MEANS** statement. Suppose the effect has m levels. By default, or if you specify PLOTS=BOXPLOT(NPANELPOS=0), all m levels of the effect are displayed in a single plot. Specifying a nonzero value of n will result in P panels, where P is the integer part of $m/n + 1$. If $n > 0$, then the levels will be approximately balanced across the P panels; whereas if $n < 0$, precisely $|n|$ levels will be displayed on each panel except possibly the last.

CONTOURFIT<(OBS=*obs-options*)>

modifies the contour fit plot produced by default when you have a model involving only two continuous predictors. The plot displays a contour plot of the predicted surface overlaid with a scatter plot of the observed data. You can use the following *obs-options* to control how the observations are displayed:

GRADIENT

specifies that observations are displayed as circles colored by the observed response. The same color gradient is used to display the fitted surface and the observations. Observations where the predicted response is close to the observed response have similar colors: the greater the contrast between the color of an observation and the surface, the larger the residual is at that point.

NONE

suppresses the observations.

OUTLINE

specifies that observations are displayed as circles with a border but with a completely transparent fill.

OUTLINEGRADIENT

is the same as `OBS=GRADIENT` except that a border is shown around each observation. This option is useful to identify the location of observations where the residuals are small, since at these points the color of the observations and the color of the surface are indistinguishable. `OBS=OUTLINEGRADIENT` is the default if you do not specify any *obs-options*.

CONTROLPLOT

requests a display in which least squares means are compared against a reference level. LS-mean control plots are produced only when you specify `PDIFF=CONTROL` or `ADJUST=DUNNETT` in the `LSMEANS` statement, and in this case they are produced by default.

DIAGNOSTICS<(LABEL UNPACK)>

requests that a panel of summary diagnostics for the fit be displayed. The panel displays scatter plots of residuals, studentized residuals, and observed responses by predicted values; studentized residuals by leverage; Cook's D by observation; a Q-Q plot of residuals; a residual histogram; and a residual-fit spread plot. The `LABEL` option displays labels on observations satisfying $RSTUDENT > 2$, $LEVERAGE > 2p/n$, and on the Cook's D plot, $COOKSD > 4/n$, where n is the number of observations used in fitting the model, and p is the number of parameters in the model. The label is the first `ID` variable if the `ID` statement is specified; otherwise, it is the observation number. The `UNPACK` option unpanels the diagnostic display and produces the series of individual plots that form the paneled display.

DIFFPLOT<(ABS NOABS CENTER NOLINES)>

modifies the plot produced by an `LSMEANS` statement with the `PDIFF=ALL` option (or just `PDIFF`, since `ALL` is the default argument). The `ABS` and `NOABS` options determine the positioning of the line segments in the plot. When the `ABS` option is in effect, and this is the default, all line segments are shown on the same side of the reference line. The `NOABS` option separates comparisons according to the sign of the difference. The `CENTER` option marks the center point for each comparison. This point corresponds to the intersection of two least squares means. The `NOLINES` option suppresses the display of the line segments that represent the confidence bounds for the differences of the least squares means. The `NOLINES` option implies the `CENTER` option. The default is to draw line segments in the upper portion of the plot area without marking the center point.

FITPLOT<(NOCLM NOCLI NOLIMITS)>

modifies the fit plot produced by default when you have a model with a single continuous predictor. By default the plot includes confidence limits for both the expected predicted values and individual new predictions. The `PLOTS=FITPLOT(NOCLM)` option removes the limits on the expected values and the `PLOTS=FITPLOT(NOCLI)` option removes the limits on new predictions. The `PLOTS=FITPLOT(NOLIMITS)` option removes both kinds of confidence limits.

INTPLOT<(CLM CLI LIMITS)>

modifies the interaction plot produced by default when you have a two-way analysis of variance model, with just two `CLASS` variables. By default the plot does not show confidence limits around the predicted values. The `PLOTS=INTPLOT(CLM)` option adds limits for the expected predicted values and `PLOTS=INTPLOT(CLI)` adds limits for new predictions. Use `PLOTS=INTPLOT(LIMITS)` to add both kinds of limits.

MEANPLOT<(CL CLBAND CONNECT ASCENDING DESCENDING)>

modifies the grouped box plot produced by an **MEANS** statement. Upper and lower confidence limits are plotted when the CL option is used. When the CLBAND option is in effect, confidence limits are shown as bands and the means are connected. By default, means are not joined by lines. You can achieve that effect with the CONNECT option. Means are displayed in the same order as they appear in the “Means” table. You can change that order for plotting with the ASCENDING and DESCENDING options.

NONE

specifies that no graphics be displayed.

RESIDUALS<(SMOOTH UNPACK)>

requests that scatter plots of the residuals against each continuous covariate be displayed. The SMOOTH option overlays a Loess smooth on each residual plot. Note that if a **WEIGHT** variable is specified, then it is not used to weight the smoother. See Chapter 72, “The LOESS Procedure,” for more information. The UNPACK option unpanels the residual display and produces a series of individual plots that form the paneled display.

ABSORB Statement

ABSORB *variables* ;

Absorption is a computational technique that provides a large reduction in time and memory requirements for certain types of models. The *variables* are one or more variables in the input data set.

For a main-effect variable that does not participate in interactions, you can absorb the effect by naming it in an ABSORB statement. This means that the effect can be adjusted out before the construction and solution of the rest of the model. This is particularly useful when the effect has a large number of levels.

Several variables can be specified, in which case each one is assumed to be nested in the preceding variable in the ABSORB statement.

NOTE: When you use the ABSORB statement, the data set (or each BY group, if a **BY** statement appears) must be sorted by the variables in the ABSORB statement. The GLM procedure cannot produce predicted values or least squares means (LS-means) or create an output data set of diagnostic values if an ABSORB statement is used. If the ABSORB statement is used, it must appear before the first RUN statement; otherwise, it is ignored.

When you use an ABSORB statement and also use the **INT** option in the **MODEL** statement, the procedure ignores the option but computes the uncorrected total sum of squares (SS) instead of the corrected total sums of squares.

See the section “[Absorption](#)” on page 3687 for more information.

BY Statement

BY *variables* ;

You can specify a BY statement with PROC GLM 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 GLM 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).

Since sorting the data changes the order in which PROC GLM reads observations, the sort order for the levels of the classification variables might be affected if you also specify **ORDER=DATA** in the **PROC GLM** statement. This, in turn, affects specifications in the **CONTRAST** and **ESTIMATE** statements.

If you specify the BY statement, it must appear before the first RUN statement; otherwise, it is ignored. When you use a BY statement, the interactive features of PROC GLM are disabled.

When both the BY and **ABSORB** statements are used, observations must be sorted first by the variables in the BY statement, and then by the variables in the **ABSORB** statement.

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 *variable* <(REF= option)> ... <*variable* <(REF= option)>> </global-options> ;

The CLASS statement names the classification variables to be used in the model. Typical classification variables are Treatment, Sex, Race, Group, and Replication. 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.

NOTE: Prior to SAS 9, class levels were determined by using no more than the first 16 characters of the formatted values. To revert to this previous behavior, you can use the TRUNCATE option in the CLASS statement.

In any case, you can use formats to group values into levels. See the discussion of the FORMAT procedure in the *Base SAS Procedures Guide* and the discussions of the FORMAT statement and SAS formats in *SAS Formats and Informats: Reference*. You can adjust the order of CLASS variable levels with the **ORDER=** option in the **PROC GLM** statement.

The GLM procedure displays a table that summarizes the CLASS variables and their levels. You can use this table to check the ordering of levels and, hence, of the corresponding parameters for main effects. If you need to check the ordering of parameters for interaction effects, use the E option in the **MODEL**,

CONTRAST, ESTIMATE, or LSMEANS statement. See the section “Parameterization of PROC GLM Models” on page 3673 for more information.

You can specify the following REF= option to indicate how the levels of an individual classification variable are to be ordered by enclosing it in parentheses after the variable name:

REF= 'level' | FIRST | LAST

specifies a level of the classification variable to be put at the end of the list of levels. This level thus corresponds to the reference level in the usual interpretation of the estimates with PROC GLM's singular parameterization. You can specify the *level* of the variable to use as the reference level; specify a value that corresponds to the formatted value of the variable if a format is assigned. Alternatively, you can specify REF=FIRST to designate that the first ordered level serve as the reference, or REF=LAST to designate that the last ordered level serve as the reference. To specify that REF=FIRST or REF=LAST be used for all classification variables, use the **REF= global-option** after the slash (/) in the CLASS statement.

You can specify the following *global-options* in the CLASS statement after a slash (/):

REF=FIRST | LAST

specifies a level of all classification variables to be put at the end of the list of levels. This level thus corresponds to the reference level in the usual interpretation of the estimates with PROC GLM's singular parameterization. Specify REF=FIRST to designate that the first ordered level for each classification variable serve as the reference. Specify REF=LAST to designate that the last ordered level serve as the reference. This option applies to all the variables specified in the CLASS statement. To specify different reference levels for different classification variables, use **REF=** options for individual variables.

TRUNCATE

specifies that class levels be determined by using only up to the first 16 characters of the formatted values of CLASS variables. When formatted values are longer than 16 characters, you can use this option to revert to the levels as determined in releases prior to SAS 9.

CODE Statement

CODE < options > ;

The CODE statement writes SAS DATA step code for computing predicted values of the fitted model either to a file or to a catalog entry. This code can then be included in a DATA step to score new data.

Table 47.5 summarizes the *options* available in the CODE statement.

Table 47.5 CODE Statement Options

Option	Description
CATALOG=	Names the catalog entry where the generated code is saved
DUMMIES	Retains the dummy variables in the data set
ERROR	Computes the error function
FILE=	Names the file where the generated code is saved
FORMAT=	Specifies the numeric format for the regression coefficients
GROUP=	Specifies the group identifier for array names and statement labels
IMPUTE	Imputes predicted values for observations with missing or invalid covariates
LINESIZE=	Specifies the line size of the generated code
LOOKUP=	Specifies the algorithm for looking up CLASS levels
RESIDUAL	Computes residuals

For details about the syntax of the CODE statement, see the section “[CODE Statement](#)” on page 393 in Chapter 19, “[Shared Concepts and Topics](#).”

CONTRAST Statement

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

The CONTRAST statement enables you to perform custom hypothesis tests by specifying an **L** vector or matrix for testing the univariate hypothesis $\mathbf{L}\boldsymbol{\beta} = 0$ or the multivariate hypothesis $\mathbf{LBM} = 0$. Thus, 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 3673. All of the elements of the **L** vector might be given, or if only certain portions of the **L** vector are given, the remaining elements are constructed by PROC GLM from the context (in a manner similar to rule 4 discussed in the section “[Construction of Least Squares Means](#)” on page 3707).

There is no limit to the number of CONTRAST statements you can specify, but they must appear after the **MODEL** statement. In addition, if you use a CONTRAST statement and a **MANOVA**, **REPEATED**, or **TEST** statement, appropriate tests for contrasts are carried out as part of the **MANOVA**, **REPEATED**, or **TEST** analysis. If you use a CONTRAST statement and a **RANDOM** statement, the expected mean square of the contrast is displayed. As a result of these additional analyses, the CONTRAST statement must appear before the **MANOVA**, **REPEATED**, **RANDOM**, or **TEST** statement.

In the CONTRAST statement,

<i>label</i>	identifies the contrast on the output. A label is required for every contrast specified. Labels must be enclosed in quotes.
<i>effect</i>	identifies an effect that appears in the MODEL statement, or the INTERCEPT effect. The INTERCEPT effect can be used when an intercept is fitted in the model. You do not need to include all effects that are in the MODEL statement.
<i>values</i>	are constants that are elements of the L vector associated with the effect.

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

E

displays the entire **L** vector. This option is useful in confirming the ordering of parameters for specifying **L**.

E=effect

specifies an error term, which must be one of the effects in the model. The procedure uses this effect as the denominator in *F* tests in univariate analysis. In addition, if you use a **MANOVA** or **REPEATED** statement, the procedure uses the effect specified by the **E=** option as the basis of the **E** matrix. By default, the procedure uses the overall residual or error mean square (MSE) as an error term.

ETYPE=n

specifies the type (1, 2, 3, or 4, corresponding to a Type I, II, III, or IV test, respectively) of the **E=** effect. If the **E=** option is specified and the **ETYPE=** option is not, the procedure uses the highest type computed in the analysis.

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 perform custom hypothesis tests. If the hypothesis is testable in the univariate case, $\text{SS}(H_0: \mathbf{L}\boldsymbol{\beta} = 0)$ is computed as

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

where $\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y}$. This is the sum of squares displayed on the analysis-of-variance table.

For multivariate testable hypotheses, the usual multivariate tests are performed 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.

For example, for the model

```
proc glm;
  class A B;
  model Y=A B;
run;
```

with **A** at 5 levels and **B** at 2 levels, the parameter vector is

$$(\mu \ \alpha_1 \ \alpha_2 \ \alpha_3 \ \alpha_4 \ \alpha_5 \ \beta_1 \ \beta_2)$$

To test the hypothesis that the pooled A linear and A quadratic effect is zero, you can use the following **L** matrix:

$$\mathbf{L} = \begin{bmatrix} 0 & -2 & -1 & 0 & 1 & 2 & 0 & 0 \\ 0 & 2 & -1 & -2 & -1 & 2 & 0 & 0 \end{bmatrix}$$

The corresponding CONTRAST statement is

```
contrast 'A LINEAR & QUADRATIC'
      a -2 -1  0  1  2,
      a  2 -1 -2 -1  2;
```

If the first level of A is a control level and you want a test of control versus others, you can use this statement:

```
contrast 'CONTROL VS OTHERS'  a -1 0.25 0.25 0.25 0.25;
```

See the following discussion of the [ESTIMATE](#) statement and the section “[Specification of ESTIMATE Expressions](#)” on page 3690 for rules on specification, construction, distribution, and estimability in the CONTRAST statement.

ESTIMATE Statement

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

The ESTIMATE statement enables you to estimate linear functions of the parameters by multiplying the vector **L** by the parameter estimate vector **b**, resulting in **Lb**. All of the elements of the **L** vector might be given, or, if only certain portions of the **L** vector are given, the remaining elements are constructed by PROC GLM from the context (in a manner similar to rule 4 discussed in the section “[Construction of Least Squares Means](#)” on page 3707).

The linear function is checked for estimability. The estimate **Lb**, where $\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y}$, is displayed along with its associated standard error, $\sqrt{\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}'s^2}$, and *t* test. If you specify the [CLPARM](#) option in the [MODEL](#) statement (see page 3657), confidence limits for the true value are also displayed.

There is no limit to the number of ESTIMATE statements that you can specify, but they must appear after the [MODEL](#) statement. In the ESTIMATE statement,

<i>label</i>	identifies the estimate on the output. A label is required for every contrast specified. Labels must be enclosed in quotes.
<i>effect</i>	identifies an effect that appears in the MODEL statement, or the INTERCEPT effect. The INTERCEPT effect can be used as an effect when an intercept is fitted in the model. You do not need to include all effects that are in the MODEL statement.
<i>values</i>	are constants that are the elements of the L vector associated with the preceding effect. For example,

```
estimate 'A1 VS A2' A 1 -1;
```

forms an estimate that is the difference between the parameters estimated for the first and second levels of the [CLASS](#) variable A.

You can specify the following *options* in the ESTIMATE statement after a slash (/):

DIVISOR=*number*

specifies a value by which to divide all coefficients so that fractional coefficients can be entered as integer numerators. For example, you can use

```
estimate '1/3(A1+A2) - 2/3A3' a 1 1 -2 / divisor=3;
```

instead of

```
estimate '1/3(A1+A2) - 2/3A3' a 0.33333 0.33333 -0.66667;
```

E

displays the entire **L** vector. This option is useful in confirming the ordering of parameters for specifying **L**.

SINGULAR=*number*

tunes the estimability checking. If $\text{ABS}(\mathbf{L} - \mathbf{LH}) > C \times \text{number}$, then the **L** vector 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.

See also the section “[Specification of ESTIMATE Expressions](#)” on page 3690.

FREQ Statement

FREQ *variable* ;

The FREQ statement names a variable that provides frequencies for each observation in the [DATA=](#) data set. Specifically, if n is the value of the FREQ variable for a given observation, then that observation is used n times.

The analysis produced using a FREQ statement reflects the expanded number of observations. For example, means and total degrees of freedom reflect the expanded number of observations. You can produce the same analysis (without the FREQ statement) by first creating a new data set that contains the expanded number of observations. For example, if the value of the FREQ variable is 5 for the first observation, the first 5 observations in the new data set are identical. Each observation in the old data set is replicated n_i times in the new data set, where n_i is the value of the FREQ variable for that observation.

If the value of the FREQ variable is missing or is less than 1, the observation is not used in the analysis. If the value is not an integer, only the integer portion is used.

If you specify the FREQ statement, it must appear before the first RUN statement or it is ignored.

ID Statement

ID *variables* ;

When predicted values are requested as a **MODEL** statement option, values of the variables given in the ID statement are displayed beside each observed, predicted, and residual value for identification. Although there are no restrictions on the length of ID variables, PROC GLM might truncate the number of values listed in order to display them on one line. The GLM procedure displays a maximum of five ID variables.

If you specify the ID statement, it must appear before the first RUN statement or it is ignored.

LSMEANS Statement

LSMEANS *effects* < / *options* > ;

Table 47.6 summarizes the *options* available in the LSMEANS statement.

Table 47.6 LSMEANS Statement Options

Option	Description
ADJUST=	Requests a multiple comparison adjustment
ALPHA=	Specifies the level of significance
AT	Enables you to modify the values of the covariates
BYLEVEL	Processes the OM data set by each level of the LS-mean effect
CL	Requests confidence limits
COV	Includes variances and covariances of the LS-means
E	Displays the coefficients of the linear functions
E=	Specifies an effect in the model to use as an error term
ETYPE=	Specifies the type of the E= effect
LINES	Uses connecting lines to indicate nonsignificantly different subsets of LS-means
NOPRINT	Suppresses the normal display of results
OBSMARGINS	Specifies a potentially different weighting scheme
OUT=	Creates an output data set
PDIF	Requests that <i>p</i> -values for differences
PLOT=	Requests graphics related to least squares means
SLICE=	Specifies effects within which to test for differences
SINGULAR=	Tunes the estimability checking
STDERR	Produces the standard error
TDIF	Produces the <i>t</i> values

Least squares means (LS-means) are computed for each *effect* listed in the LSMEANS statement. You can specify only classification effects in the LSMEANS statement—that is, effects that contain only classification variables. You can also specify options to perform multiple comparisons. In contrast to the **MEANS** statement, the LSMEANS statement performs multiple comparisons on interactions as well as main effects.

LS-means are *predicted population margins*; that is, they estimate the marginal means over a balanced population. In a sense, LS-means are to unbalanced designs as class and subclass arithmetic means are to balanced designs. Each LS-mean is computed as $L'b$ for a certain column vector L , where b is the vector of parameter estimates—that is, the solution of the normal equations. For further information, see the section “Construction of Least Squares Means” on page 3707.

Multiple effects can be specified in one LSMEANS statement, or multiple LSMEANS statements can be used, but they must all appear after the **MODEL** statement. For example:

```
proc glm;
  class A B;
  model Y=A B A*B;
  lsmeans A B A*B;
run;
```

LS-means are displayed for each level of the A, B, and A*B effects.

You can specify the following *options* in the LSMEANS statement after a slash (/):

ADJUST=BON

ADJUST=DUNNETT

ADJUST=NELSON

ADJUST=SCHEFFE

ADJUST=SIDAK

ADJUST=SIMULATE <(simoptions)>

ADJUST=SMM | GT2

ADJUST=TUKEY

ADJUST=T

requests a multiple comparison adjustment for the p -values and confidence limits for the differences of LS-means. The ADJUST= option modifies the results of the **TDIFF** and **PDIFF** options; thus, if you omit the **TDIFF** or **PDIFF** option then the ADJUST= option implies **PDIFF**. By default, PROC GLM analyzes all pairwise differences. If you specify ADJUST=DUNNETT, PROC GLM analyzes all differences with a control level. If you specify the ADJUST=NELSON option, PROC GLM analyzes all differences with the average LS-mean. The default is ADJUST=T, which really signifies no adjustment for multiple comparisons.

The BON (Bonferroni) and SIDAK adjustments involve correction factors described in the section “Multiple Comparisons” on page 3693 and in Chapter 80, “The **MULTTEST** Procedure.” When you specify ADJUST=TUKEY and your data are unbalanced, PROC GLM uses the approximation described in Kramer (1956) and identifies the adjustment as “Tukey-Kramer” in the results. Similarly, when you specify either ADJUST=DUNNETT or the ADJUST=NELSON option and the LS-means are correlated, PROC GLM uses the factor-analytic covariance approximation described in Hsu (1992) and identifies the adjustment in the results as “Dunnnett-Hsu” or “Nelson-Hsu,” respectively. The preceding references also describe the SCHEFFE and SMM adjustments.

The SIMULATE adjustment computes the adjusted p -values from the simulated distribution of the maximum or maximum absolute value of a multivariate t random vector. The simulation estimates q , the true $(1 - \alpha)$ quantile, where $1 - \alpha$ is the confidence coefficient. The default α is the value of the **ALPHA=** option in the **PROC GLM** statement or 0.05 if that option is not specified. You can change this value with the **ALPHA=** option in the LSMEANS statement.

The number of samples for the SIMULATE adjustment is set so that the tail area for the simulated q is within a certain *accuracy radius* γ of $1 - \alpha$ with an *accuracy confidence* of $100(1 - \epsilon)\%$. In equation form,

$$P(|F(\hat{q}) - (1 - \alpha)| \leq \gamma) = 1 - \epsilon$$

where $\hat{q}$ is the simulated q and F is the true distribution function of the maximum; see Edwards and Berry (1987) for details. By default, $\gamma = 0.005$ and $\epsilon = 0.01$, so that the tail area of $\hat{q}$ is within 0.005 of 0.95 with 99% confidence.

You can specify the following *simoptions* in parentheses after the ADJUST=SIMULATE option.

ACC=value

specifies the target accuracy radius γ of a $100(1 - \epsilon)\%$ confidence interval for the true probability content of the estimated $(1 - \alpha)$ quantile. The default value is ACC=0.005. Note that, if you also specify the CVADJUST *simoption*, then the actual accuracy radius will probably be substantially less than this target.

CVADJUST

specifies that the quantile should be estimated by the control variate adjustment method of Hsu and Nelson (1998) instead of simply as the quantile of the simulated sample. Specifying the CVADJUST option typically has the effect of significantly reducing the accuracy radius γ of a $100 \times (1 - \epsilon)\%$ confidence interval for the true probability content of the estimated $(1 - \alpha)$ quantile. The control-variate-adjusted quantile estimate takes roughly twice as long to compute, but it is typically much more accurate than the sample quantile.

EPS=value

specifies the value ϵ for a $100 \times (1 - \epsilon)\%$ confidence interval for the true probability content of the estimated $(1 - \alpha)$ quantile. The default value for the accuracy confidence is 99%, corresponding to EPS=0.01.

NSAMP= n

specifies the sample size for the simulation. By default, n is set based on the values of the target accuracy radius γ and accuracy confidence $100 \times (1 - \epsilon)\%$ for an interval for the true probability content of the estimated $(1 - \alpha)$ quantile. With the default values for γ , ϵ , and α (0.005, 0.01, and 0.05, respectively), NSAMP=12604 by default.

REPORT

specifies that a report on the simulation should be displayed, including a listing of the parameters, such as γ , ϵ , and α , as well as an analysis of various methods for estimating or approximating the quantile.

SEED=number

specifies an integer used to start the pseudo-random number generator for the simulation. If you do not specify a seed, or specify a value less than or equal to zero, the seed is by default generated from reading the time of day from the computer's clock.

THREADS

specifies that the computational work for the simulation be divided into parallel threads, where the number of threads is the value of the SAS system option CPUCOUNT=. For large simulations (as specified directly using the NSAMP= *simoption* or indirectly using the ACC= or EPS=

simoptions), parallel processing can markedly speed up the computation of adjusted p -values and confidence intervals. However, because the parallel processing has different pseudo-random number streams, the precise results are different from the default ones, which are computed in sequence rather than in parallel. This option overrides the SAS system option `THREADS | NOTTHREADS`.

NOTTHREADS

specifies that the computational work for the simulation be performed in sequence rather than in parallel. `NOTTHREADS` is the default. This option overrides the SAS system option `THREADS | NOTTHREADS`.

ALPHA= p

specifies the level of significance p for $100(1 - p)\%$ confidence intervals. This option is useful only if you also specify the `CL` option, and, optionally, the `PDIFF` option. By default, p is equal to the value of the `ALPHA=` option in the `PROC GLM` statement or 0.05 if that option is not specified. This value is used to set the endpoints for confidence intervals for the individual means as well as for differences between means.

AT *variable = value*

AT (*variable-list*)= (*value-list*)

AT MEANS

enables you to modify the values of the covariates used in computing LS-means. By default, all covariate effects are set equal to their mean values for computation of standard LS-means. The `AT` option enables you to set the covariates to whatever values you consider interesting. For more information, see the section “[Setting Covariate Values](#)” on page 3708.

BYLEVEL

requests that `PROC GLM` process the `OM` data set by each level of the LS-mean effect in question. For more details, see the entry for the `OM` option in this section.

CL

requests confidence limits for the individual LS-means. If you specify the `PDIFF` option, confidence limits for differences between means are produced as well. You can control the confidence level with the `ALPHA=` option. Note that, if you specify an `ADJUST=` option, the confidence limits for the differences are adjusted for multiple inference but the confidence intervals for individual means are **not** adjusted.

COV

includes variances and covariances of the LS-means in the output data set specified in the `OUT=` option in the `LSMEANS` statement. Note that this is the covariance matrix for the LS-means themselves, not the covariance matrix for the differences between the LS-means, which is used in the `PDIFF` computations. If you omit the `OUT=` option, the `COV` option has no effect. When you specify the `COV` option, you can specify only one effect in the `LSMEANS` statement.

E

displays the coefficients of the linear functions used to compute the LS-means.

E=*effect*

specifies an effect in the model to use as an error term. The procedure uses the mean square for the *effect* as the error mean square when calculating estimated standard errors (requested with the `STDERR` option) and probabilities (requested with the `STDERR`, `PDIFF`, or `TDIFF` option). Unless

you specify **STDERR**, **PDIFF** or **TDIFF**, the **E=** option is ignored. By default, if you specify the **STDERR**, **PDIFF**, or **TDIFF** option and do not specify the **E=** option, the procedure uses the error mean square for calculating standard errors and probabilities.

ETYPE=*n*

specifies the type (1, 2, 3, or 4, corresponding to a Type I, II, III, or IV test, respectively) of the **E=** effect. If you specify the **E=** option but not the **ETYPE=** option, the highest type computed in the analysis is used. If you omit the **E=** option, the **ETYPE=** option has no effect.

LINES

presents results of comparisons between all pairs of means (specified by the **PDIFF=ALL** option) by listing the means in descending order and indicating nonsignificant subsets by line segments beside the corresponding means. When all differences have the same variance, these comparison lines are guaranteed to accurately reflect the inferences based on the corresponding tests, made by comparing the respective *p*-values to the value of the **ALPHA=** option (0.05 by default). However, equal variances are rarely the case for differences between LS-means. If the variances are not all the same, then the comparison lines might be conservative, in the sense that if you base your inferences on the lines alone, you will detect fewer significant differences than the tests indicate. If there are any such differences, a note is appended to the table that lists the pairs of means that are inferred to be significantly different by the tests but not by the comparison lines. Note, however, that in many cases, even though the variances are unbalanced, they are near enough that the comparison lines in fact accurately reflect the test inferences.

NOPRINT

suppresses the normal display of results from the LSMEANS statement. This option is useful when an output data set is created with the **OUT=** option in the LSMEANS statement.

OBSMARGINS

OM

specifies a potentially different weighting scheme for computing LS-means coefficients. The standard LS-means have equal coefficients across classification effects; however, the **OM** option changes these coefficients to be proportional to those found in the input data set. For more information, see the section “[Changing the Weighting Scheme](#)” on page 3709.

The **BYLEVEL** option modifies the observed-margins LS-means. Instead of computing the margins across the entire data set, the procedure computes separate margins for each level of the LS-mean effect in question. The resulting LS-means are actually equal to raw means in this case. If you specify the **BYLEVEL** option, it disables the **AT** option.

OUT=SAS-data-set

creates an output data set that contains the values, standard errors, and, optionally, the covariances (see the **COV** option) of the LS-means.

For more information, see the section “[Output Data Sets](#)” on page 3726.

PDIFF< =*difftype*>

requests that *p*-values for differences of the LS-means be produced. The optional *difftype* specifies which differences to display. Possible values for *difftype* are ALL, CONTROL, CONTROLL, CONTROLU, and ANOM. The ALL value requests all pairwise differences, and it is the default. The CONTROL value requests the differences with a control that, by default, is the first level of each of the specified LS-mean effects. The ANOM value requests differences between each LS-mean and

the average LS-mean, as in the *analysis of means* (Ott 1967). The average is computed as a weighted mean of the LS-means, the weights being inversely proportional to the variances. Note that the ANOM procedure in SAS/QC software implements both tables and graphics for the analysis of means with a variety of response types. For one-way designs, the PDIFF=ANOM computations are equivalent to the results of PROC ANOM. See the section “[Analysis of Means: Comparing Each Treatments to the Average](#)” on page 3700 for more details.

To specify which levels of the effects are the controls, list the quoted formatted values in parentheses after the keyword CONTROL. For example, if the effects A, B, and C are CLASS variables, each having two levels, '1' and '2', the following LSMEANS statement specifies the '1' '2' level of A*B and the '2' '1' level of B*C as controls:

```
lsmeans A*B B*C / pdiff=control('1' '2', '2' '1');
```

For multiple-effect situations such as this one, the ordering of the list is significant, and you should check the output to make sure that the controls are correct.

Two-tailed tests and confidence limits are associated with the CONTROL *difftype*. For one-tailed results, use either the CONTROLL or CONTROLU *difftype*.

- **CONTROLL** tests whether the noncontrol levels are less than the control; you declare a noncontrol level to be significantly less than the control if the associated upper confidence limit for the noncontrol level minus the control is less than zero, and you ignore the associated lower confidence limits (which are set to minus infinity).
- **CONTROLU** tests whether the noncontrol levels are greater than the control; you declare a noncontrol level to be significantly greater than the control if the associated lower confidence limit for the noncontrol level minus the control is greater than zero, and you ignore the associated upper confidence limits (which are set to infinity).

The default multiple comparisons adjustment for each *difftype* is shown in the following table.

<i>difftype</i>	Default ADJUST=
Not specified	T
ALL	TUKEY
CONTROL	
CONTROLL	DUNNETT
CONTROLU	
ANOM	NELSON

If no *difftype* is specified, the default for the ADJUST= option is T (that is, no adjustment); for PDIFF=ALL, ADJUST=TUKEY is the default; for PDIFF=CONTROL, PDIFF=CONTROLL, or PDIFF=CONTROLU, the default value for the ADJUST= option is DUNNETT. For PDIFF=ANOM, ADJUST=NELSON is the default. If there is a conflict between the PDIFF= and ADJUST= options, the ADJUST= option takes precedence.

For example, in order to compute one-sided confidence limits for differences with a control, adjusted according to Dunnett’s procedure, the following statements are equivalent:

```
lsmeans Treatment / pdiff=control1 c1;
lsmeans Treatment / pdiff=control1 c1 adjust=dunnett;
```

PLOT | PLOTS<=*plot-request*<(options)>>

PLOT | PLOTS<=*(plot-request*<(options)> <...*plot-request*<(options)> >)>

requests that graphics related to least squares means be produced via ODS Graphics, provided that ODS Graphics is enabled and the *plot-request* does not conflict with other options in the LSMEANS statement. For general information about ODS Graphics, see Chapter 21, “[Statistical Graphics Using ODS](#).”

The available options and suboptions are as follows:

ALL

requests that the default plots that correspond to this LSMEANS statement be produced. The default plot depends on the options in the statement.

ANOMPLOT

ANOM

requests an analysis-of-means display in which least squares means are compared to an average least squares mean. Least squares mean ANOM plots are produced only for those model effects that are listed in LSMEANS statements and have options that do not contradict with the display. For example, the following statements produce analysis-of-mean plots for effects A and C:

```
lsmeans A / diff=anom plot=anom;
lsmeans B / diff      plot=anom;
lsmeans C /           plot=anom;
```

The **PDIF** option in the second LSMEANS statement implies all pairwise differences.

CONTROLPLOT

CONTROL

requests a display in which least squares means are visually compared against a reference level. These plots are produced only for statements with options that are compatible with control differences. For example, the following statements produce control plots for effects A and C:

```
lsmeans A / diff=control('1') plot=control;
lsmeans B / diff              plot=control;
lsmeans C                    plot=control;
```

The **PDIF** option in the second LSMEANS statement implies all pairwise differences.

DIFFPLOT<(diffplot-options)>

DIFFOGRAM<(diffplot-options)>

requests a display of all pairwise least squares mean differences and their significance. The display is also known as a “mean-mean scatter plot” when it is based on arithmetic means (Hsu 1996; Hsu and Peruggia 1994). For each comparison a line segment, centered at the LS-means in the pair, is drawn. The length of the segment corresponds to the projected width of a confidence interval for the least squares mean difference. Segments that fail to cross the 45-degree reference line correspond to significant least squares mean differences.

LS-mean difference plots are produced only for statements with options that are compatible with the display. For example, the following statements request differences against a control level for the A effect, all pairwise differences for the B effect, and the least squares means for the C effect:

```
lsmeans A / diff=control('1') plot=diffplot;
lsmeans B / diff                plot=diffplot;
lsmeans C                      plot=diffplot;
```

The **PDIF=** type in the first statement is incompatible with a display of all pairwise differences.

You can specify the following *diffplot-options*:

ABS

determines the positioning of the line segments in the plot. This is the default *diffplot-options*. When the ABS option is in effect, all line segments are shown on the same side of the reference line.

NOABS

determines the positioning of the line segments in the plot. The NOABS option separates comparisons according to the sign of the difference.

CENTER

marks the center point for each comparison. This point corresponds to the intersection of two least squares means.

NOLINES

suppresses the display of the line segments that represent the confidence bounds for the differences of the least squares means. The NOLINES option implies the CENTER option. The default is to draw line segments in the upper portion of the plot area without marking the center point.

MEANPLOT<(meanplot-options)>

requests displays of the least squares means.

The following *meanplot-options* control the display of the least squares means.

ASCENDING

displays the least squares means in ascending order. This option has no effect if means are displayed in separate plots.

CL

displays upper and lower confidence limits for the least squares means. By default, 95% limits are drawn. You can change the confidence level with the **ALPHA=** option. Confidence limits are drawn by default if the **CL** option is specified in the LSMEANS statement.

CLBAND

displays confidence limits as bands. This option implies the JOIN option.

DESCENDING

displays the least squares means in descending order. This option has no effect if means are displayed in separate plots.

ILINK

requests that means (and confidence limits) be displayed on the inverse linked scale.

JOIN**CONNECT**

connects the least squares means with lines. This option is implied by the CLBAND option.

NONE

requests that no plots be produced.

When LS-mean calculations are adjusted for multiplicity by using the **ADJUST=** option, the plots are adjusted accordingly.

SLICE=*fixed-effect* | (*fixed-effects*)

specifies effects within which to test for differences between interaction LS-mean effects. This can produce what are known as tests of simple effects (Winer 1971). For example, suppose that **A*B** is significant and you want to test for the effect of **A** within each level of **B**. The appropriate LSMEANS statement is

```
lsmeans A*B / slice=B;
```

This statement tests for the simple main effects of **A** for **B**, which are calculated by extracting the appropriate rows from the coefficient matrix for the **A*B** LS-means and using them to form an *F* test as performed by the **CONTRAST** statement.

SINGULAR=*number*

tunes the estimability checking. If $ABS(\mathbf{L} - \mathbf{LH}) > C \times \text{number}$ for any row, then **L** is declared nonestimable. **H** is the $(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}$ matrix, and *C* is $ABS(\mathbf{L})$ 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.

STDERR

produces the standard error of the LS-means and the probability level for the hypothesis $H_0: \text{LS-mean} = 0$.

TDIFF

produces the *t* values for all hypotheses $H_0: \text{LS-mean}(i) = \text{LS-mean}(j)$ and the corresponding probabilities.

MANOVA Statement

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

If the **MODEL** statement includes more than one dependent variable, you can perform multivariate analysis of variance with the MANOVA statement. The *test-options* define which effects to test, while the *detail-options* specify how to execute the tests and what results to display. Table 47.7 summarizes the *options* available in the MANOVA statement.

Table 47.7 MANOVA Statement Options

Option	Description
Test Options	
H=	Specifies hypothesis effects
E=	Specifies the error effect
M=	Specifies a transformation matrix for the dependent variables
MNAMES=	Provides names for the transformed variables
PREFIX=	Alternatively identifies the transformed variables
Detail Options	
CANONICAL	Displays a canonical analysis of the H and E matrices
ETYPE=	Specifies the type of the E matrix
HTYPE=	Specifies the type of the H matrix
MSTAT=	Specifies the method of evaluating the multivariate test statistics
ORTH	Orthogonalizes the rows of the transformation matrix
PRINTE	Displays the error SSCP matrix E
PRINTH	Displays the hypothesis SSCP matrix H
SUMMARY	Produces analysis-of-variance tables for each dependent variable

When a MANOVA statement appears before the first RUN statement, PROC GLM enters a multivariate mode with respect to the handling of missing values; in addition to observations with missing independent variables, observations with *any* missing dependent variables are excluded from the analysis. If you want to use this mode of handling missing values and do not need any multivariate analyses, specify the **MANOVA** option in the **PROC GLM** statement.

If you use both the **CONTRAST** and MANOVA statements, the MANOVA statement must appear after the **CONTRAST** statement.

Test Options

The following *options* can be specified in the MANOVA statement as *test-options* in order to define which multivariate tests to perform.

H=effects | INTERCEPT | _ALL_

specifies effects in the preceding model to use as hypothesis matrices. For each **H** matrix (the SSCP matrix associated with an effect), the **H=** specification displays the characteristic roots and vectors of $\mathbf{E}^{-1}\mathbf{H}$ (where **E** is the matrix associated with the error effect), along with the Hotelling-Lawley trace, Pillai's trace, Wilks' lambda, and Roy's greatest root. By default, these statistics are tested with approximations based on the *F* distribution. To test them with exact (but computationally intensive) calculations, use the **MSTAT=EXACT** option.

Use the keyword **INTERCEPT** to produce tests for the intercept. To produce tests for all effects listed in the **MODEL** statement, use the keyword **_ALL_** in place of a list of effects.

For background and further details, see the section “**Multivariate Analysis of Variance**” on page 3710.

E=effect

specifies the error effect. If you omit the E= specification, the GLM procedure uses the error SSCP (residual) matrix from the analysis.

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

specifies a transformation matrix for the dependent variables listed in the **MODEL** statement. The equations in the M= specification are of the form

$$\begin{aligned} c_1 \times \text{dependent-variable} &\pm c_2 \times \text{dependent-variable} \\ &\dots \pm c_n \times \text{dependent-variable} \end{aligned}$$

where the c_i values are coefficients for 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. For sample syntax, see the section “**Examples**” on page 3649.

Alternatively, you can input the transformation matrix directly by entering the elements of the matrix with commas separating the rows and parentheses surrounding the matrix. When this alternate form of input is used, 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 analysis requested in the MANOVA statement is carried out for the variables defined by the equations in the specification, not the original dependent variables. If you omit the M= option, the analysis is performed for the original dependent variables in the **MODEL** statement.

If an M= specification is included without either the **MNAMES=** or **PREFIX=** option, the variables are labeled MVAR1, MVAR2, and so forth, by default. For further information, see the section “**Multivariate Analysis of Variance**” on page 3710.

MNAMES=names

provides names for the variables defined by the equations in the **M=** specification. Names in the list correspond to the **M=** equations or to the rows of the **M** matrix (as it is entered).

PREFIX=name

is an alternative means of identifying the transformed variables defined by the **M=** specification. For example, if you specify **PREFIX=DIFF**, the transformed variables are labeled DIFF1, DIFF2, and so forth.

Detail Options

You can specify the following *options* in the MANOVA statement after a slash (/) as *detail-options*.

CANONICAL

displays a canonical analysis of the **H** and **E** matrices (transformed by the **M** matrix, if specified) instead of the default display of characteristic roots and vectors.

ETYPE=*n*

specifies the type (1, 2, 3, or 4, corresponding to a Type I, II, III, or IV test, respectively) of the **E** matrix, the SSCP matrix associated with the **E=** effect. You need this option if you use the **E=** specification to specify an error effect other than residual error and you want to specify the type of sums of squares used for the effect. If you specify **ETYPE=*n***, the corresponding test must have been performed in the **MODEL** statement, either by options **SS*n***, **En**, or the default Type I and Type III tests. By default, the procedure uses an **ETYPE=** value corresponding to the highest type (largest *n*) used in the analysis.

HTYPE=*n*

specifies the type (1, 2, 3, or 4, corresponding to a Type I, II, III, or IV test, respectively) of the **H** matrix. See the **ETYPE=** option for more details.

MSTAT=FAPPROX | EXACT

specifies the method of evaluating the multivariate test statistics. The default is **MSTAT=FAPPROX**, which specifies that the multivariate tests are evaluated using the usual approximations based on the *F* distribution, as discussed in the section “Multivariate Tests” in Chapter 4, “[Introduction to Regression Procedures](#).” Alternatively, you can specify **MSTAT=EXACT** to compute exact *p*-values for three of the four tests (Wilks’ lambda, the Hotelling-Lawley trace, and Roy’s greatest root) and an improved *F* approximation for the fourth (Pillai’s trace). While **MSTAT=EXACT** provides better control of the significance probability for the tests, especially for Roy’s greatest root, computations for the exact *p*-values can be appreciably more demanding, and are in fact infeasible for large problems (many dependent variables). Thus, although **MSTAT=EXACT** is more accurate for most data, it is not the default method. For more information about the results of **MSTAT=EXACT**, see the section “[Multivariate Analysis of Variance](#)” on page 3710.

ORTH

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

PRINTE

displays the error SSCP matrix **E**. If the **E** matrix is the error SSCP (residual) matrix from the analysis, the partial correlations of the dependent variables given the independent variables are also produced.

For example, the statement

```
manova / printe;
```

displays the error SSCP matrix and the partial correlation matrix computed from the error SSCP matrix.

PRINTH

displays the hypothesis SSCP matrix **H** associated with each effect specified by the **H=** specification.

SUMMARY

produces analysis-of-variance tables for each dependent variable. When no **M** matrix is specified, a table is displayed for each original dependent variable from the **MODEL** statement; with an **M** matrix other than the identity, a table is displayed for each transformed variable defined by the **M** matrix.

Examples

The following statements provide several examples of using a **MANOVA** statement.

```
proc glm;
  class A B;
  model Y1-Y5=A B(A) / nouni;
  manova h=A e=B(A) / printh printe htype=1 etype=1;
  manova h=B(A) / printe;
  manova h=A e=B(A) m=Y1-Y2, Y2-Y3, Y3-Y4, Y4-Y5
    prefix=diff;
  manova h=A e=B(A) m=(1 -1 0 0 0,
                        0 1 -1 0 0,
                        0 0 1 -1 0,
                        0 0 0 1 -1) prefix=diff;
run;
```

Since this **MODEL** statement requests no options for type of sums of squares, the procedure uses Type I and Type III sums of squares. The first **MANOVA** statement specifies **A** as the hypothesis effect and **B(A)** as the error effect. As a result of the **PRINTH** option, the procedure displays the hypothesis SSCP matrix associated with the **A** effect; and, as a result of the **PRINTE** option, the procedure displays the error SSCP matrix associated with the **B(A)** effect. The option **HTYPE=1** specifies a Type I **H** matrix, and the option **ETYPE=1** specifies a Type I **E** matrix.

The second **MANOVA** statement specifies **B(A)** as the hypothesis effect. Since no error effect is specified, PROC GLM uses the error SSCP matrix from the analysis as the **E** matrix. The **PRINTE** option displays this **E** matrix. Since the **E** matrix is the error SSCP matrix from the analysis, the partial correlation matrix computed from this matrix is also produced.

The third **MANOVA** statement requests the same analysis as the first **MANOVA** statement, but the analysis is carried out for variables transformed to be successive differences between the original dependent variables. The option **PREFIX=DIFF** labels the transformed variables as **DIFF1**, **DIFF2**, **DIFF3**, and **DIFF4**.

Finally, the fourth **MANOVA** statement has the identical effect as the third, but it uses an alternative form of the **M=** specification. Instead of specifying a set of equations, the fourth **MANOVA** statement specifies rows of a matrix of coefficients for the five dependent variables.

As a second example of the use of the **M=** specification, consider the following:

```
proc glm;
  class group;
  model dose1-dose4=group / nouni;
  manova h = group
    m = -3*dose1 - dose2 + dose3 + 3*dose4,
        dose1 - dose2 - dose3 + dose4,
        -dose1 + 3*dose2 - 3*dose3 + dose4
    mnames = Linear Quadratic Cubic
    / printe;
run;
```

The **M=** specification gives a transformation of the dependent variables **dose1** through **dose4** into orthogonal polynomial components, and the **MNAMES=** option labels the transformed variables **LINEAR**, **QUADRATIC**, and **CUBIC**, respectively. Since the **PRINTE** option is specified and the default residual matrix is used as an error term, the partial correlation matrix of the orthogonal polynomial components is also produced.

MEANS Statement

MEANS *effects* < / *options* > ;

Within each group corresponding to each effect specified in the MEANS statement, PROC GLM computes the arithmetic means and standard deviations of all continuous variables in the model (both dependent and independent). You can specify only classification effects in the MEANS statement—that is, effects that contain only classification variables.

Note that the arithmetic means are not adjusted for other effects in the model; for adjusted means, see the section “[LSMEANS Statement](#)” on page 3637.

If you use a [WEIGHT](#) statement, PROC GLM computes weighted means; see the section “[Weighted Means](#)” on page 3707.

You can also specify options to perform multiple comparisons. However, the MEANS statement performs multiple comparisons only for main-effect means; for multiple comparisons of interaction means, see the section “[LSMEANS Statement](#)” on page 3637.

You can use any number of MEANS statements, provided that they appear after the [MODEL](#) statement. For example, suppose A and B each have two levels. Then, if you use the statements

```
proc glm;
  class A B;
  model Y=A B A*B;
  means A B / tukey;
  means A*B;
run;
```

the means, standard deviations, and Tukey’s multiple comparisons tests are displayed for each level of the main effects A and B, and just the means and standard deviations are displayed for each of the four combinations of levels for A*B. Since multiple comparisons tests apply only to main effects, the single MEANS statement

```
means A B A*B / tukey;
```

produces the same results.

PROC GLM does not compute means for interaction effects containing continuous variables. Thus, if you have the model

```
class A;
model Y=A X A*X;
```

then the effects X and A*X cannot be used in the MEANS statement. However, if you specify the effect A in the means statement

```
means A;
```

then PROC GLM, by default, displays within-A arithmetic means of both Y and X. You can use the [DEPONLY](#) option to display means of only the dependent variables.

means A / deponly;

If you use a **WEIGHT** statement, PROC GLM computes weighted means and estimates their variance as inversely proportional to the corresponding sum of weights (see the section “**Weighted Means**” on page 3707). However, note that the statistical interpretation of multiple comparison tests for weighted means is not well understood. See the section “**Multiple Comparisons**” on page 3693 for formulas. **Table 47.8** summarizes the *options* available in the MEANS statement.

Table 47.8 MEANS Statement Options

Option	Description
Modify output	
DEONLY	Displays only means for the dependent variables
Perform multiple comparison tests	
BON	Performs Bonferroni <i>t</i> tests
DUNCAN	Performs Duncan’s multiple range test
DUNNETT	Performs Dunnett’s two-tailed <i>t</i> test
DUNNETTL	Performs Dunnett’s lower one-tailed <i>t</i> test
DUNNETTU	Performs Dunnett’s upper one-tailed <i>t</i> test
GABRIEL	Performs Gabriel’s multiple-comparison procedure
REGWQ	Performs the Ryan-Einot-Gabriel-Welsch multiple range test
SCHEFFE	Performs Scheffé’s multiple-comparison procedure
SIDAK	Performs pairwise <i>t</i> tests on differences between means
SMM or GT2	Performs pairwise comparisons based on the studentized maximum modulus and Sidak’s uncorrelated- <i>t</i> inequality
SNK	Performs the Student-Newman-Keuls multiple range test
T or LSD	Performs pairwise <i>t</i> tests
TUKEY	Performs Tukey’s studentized range test (HSD)
WALLER	Performs the Waller-Duncan <i>k</i> -ratio <i>t</i> test
Specify additional details for multiple comparison tests	
ALPHA=	Specifies the level of significance
CLDIFF	Presents confidence intervals for all pairwise differences between means
CLM	Presents results as intervals for the mean of each level of the variables
E=	Specifies the error mean square used in the multiple comparisons
ETYPE=	Specifies the type of mean square for the error effect
HTYPE=	Specifies the MS type for the hypothesis MS
KRATIO=	Specifies the Type 1/Type 2 error seriousness ratio
LINES	Lists the means in descending order and indicating nonsignificant subsets by line segments
NOSORT	Prevents the means from being sorted into descending order
Test for homogeneity of variances	
HOVTEST	Requests a homogeneity of variance test
Compensate for heterogeneous variances	
WELCH	Requests the variance-weighted one-way ANOVA of Welch (1951)

The *options* available in the MEANS statement are described in the following list.

ALPHA=

ALPHA= p specifies the level of significance for comparisons among the means. By default, p is equal to the value of the ALPHA= option in the PROC GLM statement or 0.05 if that option is not specified. You can specify any value greater than 0 and less than 1.

BON

performs Bonferroni t tests of differences between means for all main-effect means in the MEANS statement. See the CLDIFF and LINES options for a discussion of how the procedure displays results.

CLDIFF

presents results of the BON, GABRIEL, SCHEFFE, SIDAK, SMM, GT2, T, LSD, and TUKEY options as confidence intervals for all pairwise differences between means, and the results of the DUNNETT, DUNNETTU, and DUNNETTL options as confidence intervals for differences with the control. The CLDIFF option is the default for unequal cell sizes unless the DUNCAN, REGWQ, SNK, or WALLER option is specified.

CLM

presents results of the BON, GABRIEL, SCHEFFE, SIDAK, SMM, T, and LSD options as intervals for the mean of each level of the variables specified in the MEANS statement. For all options except GABRIEL, the intervals are confidence intervals for the true means. For the GABRIEL option, they are *comparison intervals* for comparing means pairwise: in this case, if the intervals corresponding to two means overlap, then the difference between them is insignificant according to Gabriel's method.

DEPONLY

displays only means for the dependent variables. By default, PROC GLM produces means for all continuous variables, including continuous independent variables.

DUNCAN

performs Duncan's multiple range test on all main-effect means given in the MEANS statement. See the LINES option for a discussion of how the procedure displays results.

DUNNETT <(formatted-control-values)>

performs Dunnett's two-tailed t test, testing if any treatments are significantly different from a single control for all main-effect means in the MEANS statement.

To specify which level of the effect is the control, enclose the formatted value in quotes and parentheses after the *keyword*. If more than one effect is specified in the MEANS statement, you can use a list of control values within the parentheses. By default, the first level of the effect is used as the control. For example:

```
means A / dunnett('CONTROL');
```

where CONTROL is the formatted control value of A. As another example:

```
means A B C / dunnett('CNTLA' 'CNTLB' 'CNTLC');
```

where CNTLA, CNTLB, and CNTLC are the formatted control values for A, B, and C, respectively.

DUNNETTL < (*formatted-control-value*) >

performs Dunnett's one-tailed t test, testing if any treatment is significantly less than the control. Control level information is specified as described for the [DUNETT](#) option.

DUNETTU < (*formatted-control-value*) >

performs Dunnett's one-tailed t test, testing if any treatment is significantly greater than the control. Control level information is specified as described for the [DUNETT](#) option.

E=effect

specifies the error mean square used in the multiple comparisons. By default, PROC GLM uses the overall residual or error mean square (MS). The effect specified with the E= option must be a term in the model; otherwise, the procedure uses the residual MS.

ETYPE=n

specifies the type of mean square for the error effect. When you specify E=*effect*, you might need to indicate which type (1, 2, 3, or 4) of MS is to be used. The n value must be one of the types specified in or implied by the [MODEL](#) statement. The default MS type is the highest type used in the analysis.

GABRIEL

performs Gabriel's multiple-comparison procedure on all main-effect means in the MEANS statement. See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

GT2

See the [SMM](#) option.

HOVTEST**HOVTEST=BARTLETT****HOVTEST=BF****HOVTEST=LEVENE** < (**TYPE= ABS | SQUARE**) >**HOVTEST=OBRIEN** < (**W=number**) >

requests a homogeneity of variance test for the groups defined by the MEANS effect. You can optionally specify a particular test; if you do not specify a test, Levene's test (Levene 1960) with TYPE=SQUARE is computed. Note that this option is ignored unless your [MODEL](#) statement specifies a simple one-way model.

The HOVTEST=BARTLETT option specifies Bartlett's test (Bartlett 1937), a modification of the normal-theory likelihood ratio test.

The HOVTEST=BF option specifies Brown and Forsythe's variation of Levene's test (Brown and Forsythe 1974).

The HOVTEST=LEVENE option specifies Levene's test (Levene 1960), which is widely considered to be the standard homogeneity of variance test. You can use the TYPE= option in parentheses to specify whether to use the absolute residuals (TYPE=ABS) or the squared residuals (TYPE=SQUARE) in Levene's test. TYPE=SQUARE is the default.

The HOVTEST=OBRIEN option specifies O'Brien's test (O'Brien 1979), which is basically a modification of HOVTEST=LEVENE(TYPE=SQUARE). You can use the W= option in parentheses to tune the variable to match the suspected kurtosis of the underlying distribution. By default, W=0.5, as suggested by O'Brien (1979, 1981).

See the section “Homogeneity of Variance in One-Way Models” on page 3706 for more details on these methods. [Example 47.10](#) illustrates the use of the [HOVTEST](#) and [WELCH](#) options in the MEANS statement in testing for equal group variances and adjusting for unequal group variances in a one-way ANOVA.

HTYPE=*n*

specifies the MS type for the hypothesis MS. The HTYPE= option is needed only when the [WALLER](#) option is specified. The default HTYPE= value is the highest type used in the model.

KRATIO=*value*

specifies the Type 1/Type 2 error seriousness ratio for the Waller-Duncan test. Reasonable values for the KRATIO= option are 50, 100, 500, which roughly correspond for the two-level case to [ALPHA](#) levels of 0.1, 0.05, and 0.01, respectively. By default, the procedure uses the value of 100.

LINES

presents results of the [BON](#), [DUNCAN](#), [GABRIEL](#), [REGWQ](#), [SCHEFFE](#), [SIDAK](#), [SMM](#), [GT2](#), [SNK](#), [T](#), [LSD](#), [TUKEY](#), and [WALLER](#) options by listing the means in descending order and indicating nonsignificant subsets by line segments beside the corresponding means. The LINES option is appropriate for equal cell sizes, for which it is the default. The LINES option is also the default if the [DUNCAN](#), [REGWQ](#), [SNK](#), or [WALLER](#) option is specified, or if there are only two cells of unequal size. The LINES option cannot be used in combination with the [DUNNETT](#), [DUNNETTL](#), or [DUNNETTU](#) option. In addition, the procedure has a restriction that no more than 24 overlapping groups of means can exist. If a mean belongs to more than 24 groups, the procedure issues an error message. You can either reduce the number of levels of the variable or use a multiple comparison test that allows the [CLDIFF](#) option rather than the LINES option.

NOTE: If the cell sizes are unequal, the harmonic mean of the cell sizes is used to compute the critical ranges. This approach is reasonable if the cell sizes are not too different, but it can lead to liberal tests if the cell sizes are highly disparate. In this case, you should not use the LINES option for displaying multiple comparisons results; use the [TUKEY](#) and [CLDIFF](#) options instead.

LSD

See the [T](#) option.

NOSORT

prevents the means from being sorted into descending order when the [CLDIFF](#) or [CLM](#) option is specified.

REGWQ

performs the Ryan-Einot-Gabriel-Welsch multiple range test on all main-effect means in the MEANS statement. See the [LINES](#) option for a discussion of how the procedure displays results.

SCHEFFE

performs Scheffé’s multiple-comparison procedure on all main-effect means in the MEANS statement. See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

SIDAK

performs pairwise *t* tests on differences between means with levels adjusted according to Sidak’s inequality for all main-effect means in the MEANS statement. See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

SMM**GT2**

performs pairwise comparisons based on the studentized maximum modulus and Sidak's uncorrelated- t inequality, yielding Hochberg's GT2 method when sample sizes are unequal, for all main-effect means in the MEANS statement. See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

SNK

performs the Student-Newman-Keuls multiple range test on all main-effect means in the MEANS statement. See the [LINES](#) option for discussions of how the procedure displays results.

T**LSD**

performs pairwise t tests, equivalent to Fisher's least significant difference test in the case of equal cell sizes, for all main-effect means in the MEANS statement. See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

TUKEY

performs Tukey's studentized range test (HSD) on all main-effect means in the MEANS statement. (When the group sizes are different, this is the Tukey-Kramer test.) See the [CLDIFF](#) and [LINES](#) options for discussions of how the procedure displays results.

WALLER

performs the Waller-Duncan k -ratio t test on all main-effect means in the MEANS statement. See the [KRATIO=](#) and [HTYPE=](#) options for information about controlling details of the test, and the [LINES](#) option for a discussion of how the procedure displays results.

WELCH

requests the variance-weighted one-way ANOVA of Welch (1951). This alternative to the usual analysis of variance for a one-way model is robust to the assumption of equal within-group variances. This option is ignored unless your [MODEL](#) statement specifies a simple one-way model.

Note that using the WELCH option merely produces one additional table consisting of Welch's ANOVA. It does not affect all of the other tests displayed by the GLM procedure, which still require the assumption of equal variance for exact validity.

See the section "[Homogeneity of Variance in One-Way Models](#)" on page 3706 for more details on Welch's ANOVA. [Example 47.10](#) illustrates the use of the [HOVTEST](#) and WELCH options in the MEANS statement in testing for equal group variances and adjusting for unequal group variances in a one-way ANOVA.

MODEL Statement

MODEL *dependent-variables = independent-effects* *</ options >* ;

The MODEL statement names the dependent variables and independent effects. The syntax of effects is described in the section "[Specification of Effects](#)" on page 3670. For any model effect involving classification variables (interactions as well as main effects), 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 (in contrast to the REG procedure, for example, which allows several MODEL statements in the same PROC REG run).

Table 47.9 summarizes the *options* available in the MODEL statement.

Table 47.9 MODEL Statement Options

Option	Description
Produce effect size information	
EFFECTSIZE	Adds measures of effect size to each analysis of variance table
Produce tests for the intercept	
INTERCEPT	Produces the hypothesis tests associated with the intercept
Omit the intercept parameter from model	
NOINT	Omits the intercept parameter from the model
Produce parameter estimates	
SOLUTION	Produces parameter estimates
Produce tolerance analysis	
TOLERANCE	Displays the tolerances used in the SWEEP routine
Suppress univariate tests and output	
NOUNI	Suppresses the display of univariate statistics
Display estimable functions	
E	Displays the general form of all estimable functions
E1	Displays the Type I estimable functions and computes sums of squares
E2	Displays the Type II estimable functions and computes sums of squares
E3	Displays the Type III estimable functions and computes sums of squares
E4	Displays the Type IV estimable functions computes sums of squares
ALIASING	Specifies that the estimable functions should be displayed as an <i>aliasing structure</i>
Control hypothesis tests performed	
SS1	Displays Type I sum of squares
SS2	Displays Type II sum of squares
SS3	Displays Type III sum of squares
SS4	Displays Type IV sum of squares
Produce confidence intervals	
ALPHA=	Specifies the level of significance
CLI	Produces confidence limits for individual predicted values
CLM	Produces confidence limits for a mean predicted value
CLPARM	Produces confidence limits for the parameter estimates
Display predicted and residual values	
P	Displays observed, predicted, and residual values
Display intermediate calculations	
INVERSE	Displays the augmented inverse (or generalized inverse) $X'X$ matrix
XPX	Displays the augmented $X'X$ crossproducts matrix
Tune sensitivity	
SINGULAR=	Tunes the sensitivity of the regression routine to linear dependencies

Table 47.9 *continued*

Option	Description
ZETA=	Tunes the sensitivity of the check for estimability for Type III and Type IV functions

The *options* available in the MODEL statement are described in the following list.

ALIASING

specifies that the estimable functions should be displayed as an *aliasing structure*, for which each row says which linear combination of the parameters is estimated by each estimable function; also, this option adds a column of the same information to the table of parameter estimates, giving for each parameter the expected value of the estimate associated with that parameter. This option is most useful in fractional factorial experiments that can be analyzed without a **CLASS** statement.

ALPHA=*p*

specifies the level of significance *p* for $100(1 - p)\%$ confidence intervals. By default, *p* is equal to the value of the **ALPHA=** option in the **PROC GLM** statement, or 0.05 if that option is not specified. You can use values between 0 and 1.

CLI

produces confidence limits for individual predicted values for each observation. The CLI option is ignored if the **CLM** option is also specified.

CLM

produces confidence limits for a mean predicted value for each observation.

CLPARM

produces confidence limits for the parameter estimates (if the **SOLUTION** option is also specified) and for the results of all **ESTIMATE** statements.

E

displays the general form of all estimable functions. This is useful for determining the order of parameters when you are writing **CONTRAST** and **ESTIMATE** statements.

E1

displays the Type I estimable functions for each effect in the model and computes the corresponding sums of squares.

E2

displays the Type II estimable functions for each effect in the model and computes the corresponding sums of squares.

E3

displays the Type III estimable functions for each effect in the model and computes the corresponding sums of squares.

E4

displays the Type IV estimable functions for each effect in the model and computes the corresponding sums of squares.

EFFECTSIZE

adds measures of effect size to each analysis of variance table displayed by the procedure, except for those displayed by the **TEST** option in the **RANDOM** statement, by **CONTRAST** and **TEST** statements with the **E=** option, and by **MANOVA** and **REPEATED** statements. The effect size measures include the intraclass correlation and both estimates and confidence intervals for the noncentrality for the *F* test, as well as for the semipartial and partial correlation ratios. For more information about the computation and interpretation of these measures, see the section “Effect Size Measures for *F* Tests in GLM” on page 3683.

INTERCEPT**INT**

produces the hypothesis tests associated with the intercept as an effect in the model. By default, the procedure includes the intercept in the model but does not display associated tests of hypotheses. Except for producing the uncorrected total sum of squares instead of the corrected total sum of squares, the **INT** option is ignored when you use an **ABSORB** statement.

INVERSE**I**

displays the augmented inverse (or generalized inverse) $X'X$ matrix:

$$\begin{bmatrix} (X'X)^- & (X'X)^-X'y \\ y'X(X'X)^- & y'y - y'X(X'X)^-X'y \end{bmatrix}$$

The upper-left corner is the generalized inverse of $X'X$, the upper-right corner is the parameter estimates, and the lower-right corner is the error sum of squares.

NOINT

omits the intercept parameter from the model. The **NOINT** option is ignored when you use an **ABSORB** statement.

NOUNI

suppresses the display of univariate statistics. You typically use the **NOUNI** option with a multivariate or repeated measures analysis of variance when you do not need the standard univariate results. The **NOUNI** option in a **MODEL** statement does not affect the univariate output produced by the **REPEATED** statement.

P

displays observed, predicted, and residual values for each observation that does not contain missing values for independent variables. The Durbin-Watson statistic is also displayed when the **P** option is specified. The **PRESS** statistic is also produced if either the **CLM** or **CLI** option is specified.

SINGULAR=number

tunes the sensitivity of the regression routine to linear dependencies in the design. If a diagonal pivot element is less than $C \times \text{number}$ as PROC GLM sweeps the $X'X$ matrix, the associated design column is declared to be linearly dependent with previous columns, and the associated parameter is zeroed.

The C value adjusts the check to the relative scale of the variable. The C value is equal to the corrected sum of squares for the variable, unless the corrected sum of squares is 0, in which case C is 1. If you specify the **NOINT** option but not the **ABSORB** statement, PROC GLM uses the uncorrected sum of squares instead.

The default value of the **SINGULAR=** option, which is approximately 10^{-9} on most machines, might be too small, but this value is necessary in order to handle the high-degree polynomials used in the literature to compare regression routines.

SOLUTION

produces a solution to the normal equations (parameter estimates). PROC GLM displays a solution by default when your model involves no classification variables, so you need this option only if you want to see the solution for models with classification effects.

SS1

displays the sum of squares associated with Type I estimable functions for each effect. These are also displayed by default.

SS2

displays the sum of squares associated with Type II estimable functions for each effect.

SS3

displays the sum of squares associated with Type III estimable functions for each effect. These are also displayed by default.

SS4

displays the sum of squares associated with Type IV estimable functions for each effect.

TOLERANCE

displays the tolerances used in the SWEEP routine. The tolerances are of the form C/USS or C/CSS , as described in the discussion of the **SINGULAR=** option. The tolerance value for the intercept is not divided by its uncorrected sum of squares.

XPX

displays the augmented $X'X$ crossproducts matrix:

$$\begin{bmatrix} X'X & X'y \\ y'X & y'y \end{bmatrix}$$

ZETA=value

tunes the sensitivity of the check for estimability for Type III and Type IV functions. Any element in the estimable function basis with an absolute value less than the **ZETA=** option is set to zero. The default value for the **ZETA=** option is 10^{-8} .

Although it is possible to generate data for which this absolute check can be defeated, the check suffices in most practical examples. Additional research is needed in order to make this check relative rather than absolute.

OUTPUT Statement

OUTPUT < **OUT**=SAS-data-set> *keyword=names* < . . . *keyword=names*> < *option*> ;

The OUTPUT statement creates a new SAS data set that saves diagnostic measures calculated after fitting the model. At least one specification of the form *keyword=names* is required.

All the variables in the original data set are included in the new data set, along with variables created in the OUTPUT statement. These new variables contain the values of a variety of diagnostic measures that are calculated for each observation in the data set. If you want to create a SAS data set in a permanent library, you must specify a two-level name. For more information about permanent libraries and SAS data sets, see *SAS Language Reference: Concepts*.

Details on the specifications in the OUTPUT statement follow.

keyword=names

specifies the statistics to include in the output data set and provides names to the new variables that contain the statistics. Specify a *keyword* for each desired statistic (see the following list of *keywords*), an equal sign, and the variable or variables to contain the statistic.

In the output data set, the first variable listed after a *keyword* in the OUTPUT statement contains that statistic for the first dependent variable listed in the **MODEL** statement; the second variable contains the statistic for the second dependent variable in the **MODEL** statement, and so on. The list of variables following the equal sign can be shorter than the list of dependent variables in the **MODEL** statement. In this case, the procedure creates the new names in order of the dependent variables in the **MODEL** statement. See the section “[Examples](#)” on page 3662.

The *keywords* allowed and the statistics they represent are as follows:

COOKD

Cook’s *D* influence statistic

COVRATIO

standard influence of observation on covariance of parameter estimates

DFFITS

standard influence of observation on predicted value

H

leverage, $h_i = x_i(\mathbf{X}'\mathbf{X})^{-1}x_i'$

LCL

lower bound of a $100(1 - p)\%$ confidence interval for an individual prediction. The *p*-level is equal to the value of the **ALPHA=** option in the OUTPUT statement or, if this option is not specified, to the **ALPHA=** option in the **PROC GLM** statement. If neither of these options is set, then $p = 0.05$ by default, resulting in the lower bound for a 95% confidence interval. The interval also depends on the variance of the error, as well as the variance of the parameter estimates. For the corresponding upper bound, see the **UCL** keyword.

LCLM

lower bound of a $100(1 - p)\%$ confidence interval for the expected value (mean) of the predicted value. The p -level is equal to the value of the **ALPHA=** option in the OUTPUT statement or, if this option is not specified, to the **ALPHA=** option in the **PROC GLM** statement. If neither of these options is set, then $p = 0.05$ by default, resulting in the lower bound for a 95% confidence interval. For the corresponding upper bound, see the **UCLM** keyword.

PREDICTED | P

predicted values

PRESS

residual for the i th observation that results from dropping it and predicting it on the basis of all other observations. This is the residual divided by $(1 - h_i)$, where h_i is the **leverage**, defined previously.

RESIDUAL | R

residuals, calculated as **ACTUAL** – **PREDICTED**

RSTUDENT

a studentized residual with the current observation deleted

STDI

standard error of the individual predicted value

STDP

standard error of the mean predicted value

STDR

standard error of the residual

STUDENT

studentized residuals, the residual divided by its standard error

UCL

upper bound of a $100(1 - p)\%$ confidence interval for an individual prediction. The p -level is equal to the value of the **ALPHA=** option in the OUTPUT statement or, if this option is not specified, to the **ALPHA=** option in the **PROC GLM** statement. If neither of these options is set, then $p = 0.05$ by default, resulting in the upper bound for a 95% confidence interval. The interval also depends on the variance of the error, as well as the variance of the parameter estimates. For the corresponding lower bound, see the **LCL** keyword.

UCLM

upper bound of a $100(1 - p)\%$ confidence interval for the expected value (mean) of the predicted value. The p -level is equal to the value of the **ALPHA=** option in the OUTPUT statement or, if this option is not specified, to the **ALPHA=** option in the **PROC GLM** statement. If neither of these options is set, then $p = 0.05$ by default, resulting in the upper bound for a 95% confidence interval. For the corresponding lower bound, see the **LCLM** keyword.

OUT=SAS-data-set

gives the name of the new data set. By default, the procedure uses the `DATA $n$` convention to name the new data set.

The following *option* is available in the OUTPUT statement and is specified after a slash (/):

ALPHA= p

specifies the level of significance p for $100(1 - p)\%$ confidence intervals. By default, p is equal to the value of the ALPHA= option in the PROC GLM statement or 0.05 if that option is not specified. You can use values between 0 and 1.

See Chapter 4, “Introduction to Regression Procedures,” and the section “Influence Statistics” on page 8048 in Chapter 99, “The REG Procedure,” for details on the calculation of these statistics.

Examples

The following statements show the syntax for creating an output data set with a single dependent variable.

```
proc glm;
  class a b;
  model y=a b a*b;
  output out=new p=yhat r=resid stdr=eresid;
run;
```

These statements create an output data set named new. In addition to all the variables from the original data set, new contains the variable yhat, with values that are predicted values of the dependent variable y; the variable resid, with values that are the residual values of y; and the variable eresid, with values that are the standard errors of the residuals.

The following statements show a situation with five dependent variables.

```
proc glm;
  by group;
  class a;
  model y1-y5=a x(a);
  output out=pout predicted=py1-py5;
run;
```

The data set pout contains five new variables, py1 through py5. The values of py1 are the predicted values of y1; the values of py2 are the predicted values of y2; and so on.

For more information about the data set produced by the OUTPUT statement, see the section “Output Data Sets” on page 3726.

RANDOM Statement

RANDOM *effects* </ *options* > ;

When some model effects are random (that is, assumed to be sampled from a normal population of effects), you can specify these effects in the RANDOM statement in order to compute the expected values of mean squares for various model effects and contrasts and, optionally, to perform random effects analysis of variance tests. You can use as many RANDOM statements as you want, provided that they appear after the MODEL

statement. If you use a **CONTRAST** statement with a **RANDOM** statement and you want to obtain the expected mean squares for the contrast hypothesis, you must enter the **CONTRAST** statement before the **RANDOM** statement.

NOTE: PROC GLM uses only the information pertaining to expected mean squares when you specify the **TEST** option in the **RANDOM** statement and, even then, only in the extra *F* tests produced by the **RANDOM** statement. Other features in the GLM procedure—including the results of the **LSMEANS** and **ESTIMATE** statements—assume that all effects are fixed, so that all tests and estimability checks for these statements are based on a fixed-effects model, even when you use a **RANDOM** statement. Therefore, you should use the **MIXED** procedure to compute tests involving these features that take the random effects into account; see the section “**PROC GLM versus PROC MIXED for Random-Effects Analysis**” on page 3720 and Chapter 78, “**The MIXED Procedure**,” for more information.

When you use the **RANDOM** statement, by default the GLM procedure produces the Type III expected mean squares for model effects and for contrasts specified before the **RANDOM** statement in the program statements. In order to obtain expected values for other types of mean squares, you need to specify which types of mean squares are of interest in the **MODEL** statement. See the section “**Computing Type I, II, and IV Expected Mean Squares**” on page 3722 for more information.

The list of effects in the **RANDOM** statement should contain one or more of the pure classification effects specified in the **MODEL** statement (that is, main effects, crossed effects, or nested effects involving only classification variables). The coefficients corresponding to each effect specified are assumed to be normally and independently distributed with common variance. Levels in different effects are assumed to be independent.

You can specify the following *options* in the **RANDOM** statement after a slash (/):

Q

displays all quadratic forms in the fixed effects that appear in the expected mean squares. For some designs, such as large mixed-level factorials, the **Q** option might generate a substantial amount of output.

TEST

performs hypothesis tests for each effect specified in the model, using appropriate error terms as determined by the expected mean squares.

CAUTION: PROC GLM does not automatically declare interactions to be random when the effects in the interaction are declared random. For example,

```
random a b / test;
```

does not produce the same expected mean squares or tests as

```
random a b a*b / test;
```

To ensure correct tests, you need to list all random interactions and random main effects in the **RANDOM** statement.

See the section “**Random-Effects Analysis**” on page 3719 for more information about the calculation of expected mean squares and tests and on the similarities and differences between the GLM and **MIXED** procedures. See Chapter 5, “**Introduction to Analysis of Variance Procedures**,” and Chapter 78, “**The MIXED Procedure**,” for more information about random effects.

REPEATED Statement

REPEATED *factor-specification* < / options > ;

When values of the dependent variables in the **MODEL** statement represent repeated measurements on the same experimental unit, the **REPEATED** statement enables you to test hypotheses about the measurement factors (often called *within-subject factors*) as well as the interactions of within-subject factors with independent variables in the **MODEL** statement (often called *between-subject factors*). The **REPEATED** statement provides multivariate and univariate tests as well as hypothesis tests for a variety of single-degree-of-freedom contrasts. There is no limit to the number of within-subject factors that can be specified.

The **REPEATED** statement is typically used for handling repeated measures designs with one repeated response variable. Usually, the variables on the left-hand side of the equation in the **MODEL** statement represent one repeated response variable. This does not mean that only one factor can be listed in the **REPEATED** statement. For example, one repeated response variable (hemoglobin count) might be measured 12 times (implying variables Y1 to Y12 on the left-hand side of the equal sign in the **MODEL** statement), with the associated within-subject factors treatment and time (implying two factors listed in the **REPEATED** statement). See the section “[Examples](#)” on page 3667 for an example of how PROC GLM handles this case.

Designs with two or more repeated response variables can, however, be handled with the **IDENTITY transformation**; see the description of this transformation in the following section, and see [Example 47.9](#) for an example of analyzing a doubly multivariate repeated measures design.

When a **REPEATED** statement appears, the GLM procedure enters a multivariate mode of handling missing values. If any values for variables corresponding to each combination of the within-subject factors are missing, the observation is excluded from the analysis.

If you use a **CONTRAST** or **TEST** statement with a **REPEATED** statement, you must enter the **CONTRAST** or **TEST** statement before the **REPEATED** statement.

The simplest form of the **REPEATED** statement requires only a *factor-name*. With two 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*), a *transformation* that defines single-degree-of-freedom contrasts, and *options* for additional analyses and output. 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.

These terms are described in the following section, “Syntax Details.”

Syntax Details

Table 47.10 summarizes the *options* available in the **REPEATED** statement.

Table 47.10 REPEATED Statement Options

Option	Description
CANONICAL	Performs a canonical analysis of the H and E matrices
HTYPE=	Specifies the type of the H matrix used for analysis
MEAN	Generates the overall arithmetic means of the within-subject variables
MSTAT=	Specifies the method of evaluating the test statistics

Table 47.10 *continued*

Option	Description
NOM	Displays only univariate analyses results
NOU	Displays only multivariate analyses results
PRINTE	Displays E and partial correlation matrices for multivariate tests, and prints sphericity tests
PRINTH	Displays the H matrices for multivariate tests
PRINTM	Displays the transformation matrices that define tested contrasts
PRINTRV	Displays characteristic roots and vectors for multivariate tests
SUMMARY	Produces analysis-of-variance tables for univariate tests
UEPSDEF=	Specifies the <i>F</i> test adjustment for univariate tests

You can specify the following terms in the REPEATED statement.

factor-specification

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

<i>factor-name</i>	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 specifying more than one factor, list the dependent variables in the MODEL statement so that the within-subject factors defined in the REPEATED statement are nested; that is, the first factor defined in the REPEATED statement should be the one with values that change least frequently.
<i>levels</i>	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, <i>levels</i> is optional. When more than one within-subject factor is defined, however, <i>levels</i> 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.
(<i>level-values</i>)	gives values that correspond to levels of a repeated-measures factor. These values are used to label output and as spacings for constructing orthogonal polynomial contrasts if you specify a POLYNOMIAL transformation. The number of values specified must correspond to the number of levels for that factor in the REPEATED statement. Enclose the <i>level-values</i> in parentheses.

The following *transformation* keywords define single-degree-of-freedom contrasts for factors specified in the REPEATED statement. Since the number of contrasts 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; this transformation is not composed of contrasts and has the same degrees of freedom as the factor has levels. By default, the procedure uses the **CONTRAST** transformation.

CONTRAST<(ordinal-reference-level)>

generates contrasts between levels of the factor and a reference level. By default, the procedure 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 specified. For example, to generate contrasts between the first level of a factor and the other levels, use

```
contrast (1)
```

HELMERT

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

IDENTITY

generates an identity transformation corresponding to the associated factor. This transformation is *not* composed of contrasts; it has n degrees of freedom for an n -level factor, instead of $n - 1$. This can be used for doubly multivariate repeated measures.

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. Without a reference level, the contrast involving the last level is omitted. See the CONTRAST transformation for an example.

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.

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

CANONICAL

performs a canonical analysis of the **H** and **E** matrices corresponding to the transformed variables specified in the REPEATED statement.

HTYPE= n

specifies the type of the **H** matrix used in the multivariate tests and the type of sums of squares used in the univariate tests. See the **HTYPE**= option in the specifications for the **MANOVA** statement for further details.

MEAN

generates the overall arithmetic means of the within-subject variables.

MSTAT=FAPPROX | EXACT

specifies the method of evaluating the test statistics for the multivariate analysis. The default is **MSTAT**=FAPPROX, which specifies that the multivariate tests are evaluated using the usual approximations based on the F distribution, as discussed in the section “Multivariate Tests” in Chapter 4, “[Introduction to Regression Procedures](#).” Alternatively, you can specify **MSTAT**=EXACT to compute exact p -values for three of the four tests (Wilks’ lambda, the Hotelling-Lawley trace, and Roy’s greatest root) and an improved F approximation for the fourth (Pillai’s trace). While **MSTAT**=EXACT

provides better control of the significance probability for the tests, especially for Roy's greatest root, computations for the exact p -values can be appreciably more demanding, and are in fact infeasible for large problems (many dependent variables). Thus, although `MSTAT=EXACT` is more accurate for most data, it is not the default method. For more information about the results of `MSTAT=EXACT`, see the section "[Multivariate Analysis of Variance](#)" on page 3710.

NOM

displays only the results of the univariate analyses.

NOU

displays only the results of the multivariate analyses.

PRINTE

displays the **E** matrix for each combination of within-subject factors, as well as partial correlation matrices for both the original dependent variables and the variables defined by the transformations specified in the `REPEATED` statement. In addition, the `PRINTE` option provides sphericity tests for each set of transformed variables. If the requested transformations are not orthogonal, the `PRINTE` option also provides a sphericity test for a set of orthogonal contrasts.

PRINTH

displays the **H** (SSCP) matrix associated with each multivariate test.

PRINTM

displays the transformation matrices that define the contrasts in the analysis. `PROC GLM` always displays the **M** matrix so that the transformed variables are defined by the rows, not the columns, of the displayed **M** matrix. In other words, `PROC GLM` actually displays M' .

PRINTRV

displays the characteristic roots and vectors for each multivariate test.

SUMMARY

produces analysis-of-variance tables for each contrast defined by the within-subject factors. Along with tests for the effects of the independent variables specified in the `MODEL` statement, a term labeled **MEAN** tests the hypothesis that the overall mean of the contrast is zero.

UEPSDEF=unbiased-epsilon-definition

specifies the type of adjustment for the univariate F test that is displayed in addition to the Greenhouse-Geisser adjustment. The default is `UEPSDEF=HFL`, corresponding 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 (the only available method in previous releases of SAS/STAT software), and `UEPSDEF=CM`, the adjustment of Chi et al. (2012). See the section "[Hypothesis Testing in Repeated Measures Analysis](#)" on page 3714 for details about these adjustments.

Examples

When specifying more than one factor, list the dependent variables in the `MODEL` statement so that the within-subject factors defined in the `REPEATED` statement are nested; that is, the first factor defined in the `REPEATED` statement should be the one with values that change least frequently. For example, assume that three treatments are administered at each of four times, for a total of twelve dependent variables on each experimental unit. If the variables are listed in the `MODEL` statement as **Y1** through **Y12**, then the `REPEATED` statement in

```
proc glm;
  class group;
  model Y1-Y12=group / nouni;
  repeated trt 3, time 4;
run;
```

implies the following structure:

	Dependent Variables											
	Y1	Y2	Y3	Y4	Y5	Y6	Y7	Y8	Y9	Y10	Y11	Y12
Value of trt	1	1	1	1	2	2	2	2	3	3	3	3
Value of time	1	2	3	4	1	2	3	4	1	2	3	4

The REPEATED statement always produces a table like the preceding one. For more information, see the section “[Repeated Measures Analysis of Variance](#)” on page 3712.

STORE Statement

STORE <OUT=>*item-store-name* </ LABEL='label'> ;

The STORE statement requests that the procedure save the context and results of the statistical analysis. The resulting item store has a binary file format that cannot be modified. The contents of the item store can be processed with the PLM procedure. For details about the syntax of the STORE statement, see the section “[STORE Statement](#)” on page 509 in Chapter 19, “[Shared Concepts and Topics](#).”

TEST Statement

TEST H=*effects* E=*effect* </ options> ;

By default, for each sum of squares in the analysis an F value is computed that uses the residual MS as an error term. Use a TEST statement to request additional F tests that use other effects as error terms. You need a TEST statement when a nonstandard error structure (as in a split-plot design) exists. Note, however, that this might not be appropriate if the design is unbalanced, since in most unbalanced designs with nonstandard error structures, mean squares are not necessarily independent with equal expectations under the null hypothesis.

CAUTION: The GLM procedure does not check any of the assumptions underlying the F statistic. When you specify a TEST statement, you assume sole responsibility for the validity of the F statistic produced. To help validate a test, you can use the [RANDOM](#) statement and inspect the expected mean squares, or you can use the [TEST](#) option of the [RANDOM](#) statement.

You can use as many TEST statements as you want, provided that they appear after the [MODEL](#) statement.

You can specify the following terms in the TEST statement.

H=effects

specifies which effects in the preceding model are to be used as hypothesis (numerator) effects. The H= specification is required.

E=effect

specifies one, and only one, effect to use as the error (denominator) term. The E= specification is required.

By default, the sum of squares type for all hypothesis sum of squares and error sum of squares is the highest type computed in the model. If the hypothesis type or error type is to be another type that was computed in the model, you should specify one or both of the following *options* after a slash (/).

ETYPE=*n*

specifies the type of sum of squares to use for the error term. The type must be a type computed in the model (*n*=1, 2, 3, or 4).

HTYPE=*n*

specifies the type of sum of squares to use for the hypothesis. The type must be a type computed in the model (*n*=1, 2, 3, or 4).

This example illustrates the TEST statement with a split-plot model:

```
proc glm;
  class a b c;
  model y=a b(a) c a*c b*c(a);
  test h=a e=b(a) / htype=1 etype=1;
  test h=c a*c e=b*c(a) / htype=1 etype=1;
run;
```

WEIGHT Statement

WEIGHT *variable* ;

When a WEIGHT statement is used, a weighted residual sum of squares

$$\sum_i w_i (y_i - \hat{y}_i)^2$$

is minimized, where w_i is the value of the variable specified in the WEIGHT statement, y_i is the observed value of the response variable, and $\hat{y}_i$ is the predicted value of the response variable.

If you specify the WEIGHT statement, it must appear before the first RUN statement or it is ignored.

An observation is used in the analysis only if the value of the WEIGHT statement variable is nonmissing and greater than zero.

The WEIGHT statement has no effect on degrees of freedom or number of observations, but it is used by the **MEANS** statement when calculating means and performing multiple comparison tests (as described in the section “**MEANS Statement**” on page 3650).

The normal equations used when a WEIGHT statement is present are

$$\mathbf{X}'\mathbf{W}\mathbf{X}\boldsymbol{\beta} = \mathbf{X}'\mathbf{W}\mathbf{Y}$$

where $\mathbf{W}$ is a diagonal matrix consisting of the values of the variable specified in the WEIGHT statement.

If the weights for the observations are proportional to the reciprocals of the error variances, then the weighted least squares estimates are best linear unbiased estimators (BLUE).

Details: GLM Procedure

Statistical Assumptions for Using PROC GLM

The basic statistical assumption underlying the least squares approach to general linear modeling is that the observed values of each dependent variable can be written as the sum of two parts: a fixed component $x'\boldsymbol{\beta}$, which is a linear function of the independent coefficients, and a random noise, or error, component ϵ :

$$y = x'\boldsymbol{\beta} + \epsilon$$

The independent coefficients x are constructed from the model effects as described in the section “[Parameterization of PROC GLM Models](#)” on page 3673. Further, the errors for different observations are assumed to be uncorrelated with identical variances. Thus, this model can be written

$$E(Y) = \mathbf{X}\boldsymbol{\beta}, \quad \text{Var}(Y) = \sigma^2 I$$

where Y is the vector of dependent variable values, $\mathbf{X}$ is the matrix of independent coefficients, I is the identity matrix, and σ^2 is the common variance for the errors. For multiple dependent variables, the model is similar except that the errors for different dependent variables within the same observation are not assumed to be uncorrelated. This yields a multivariate linear model of the form

$$E(Y) = \mathbf{X}\mathbf{B}, \quad \text{Var}(\text{vec}(Y)) = \boldsymbol{\Sigma} \otimes I$$

where Y and B are now matrices, with one column for each dependent variable, $\text{vec}(Y)$ strings Y out by rows, and $\otimes$ indicates the Kronecker matrix product.

Under the assumptions thus far discussed, the least squares approach provides estimates of the linear parameters that are unbiased and have minimum variance among linear estimators. Under the further assumption that the errors have a normal (or Gaussian) distribution, the least squares estimates are the maximum likelihood estimates and their distribution is known. All of the significance levels (“ p values”) and confidence limits calculated by the GLM procedure require this assumption of normality in order to be exactly valid, although they are good approximations in many other cases.

Specification of Effects

Each term in a model, called an *effect*, is a variable or combination of variables. Effects are specified with a special notation that uses variable names and operators. There are two kinds of variables: *classification* (or

CLASS) variables and *continuous variables*. There are two primary operators: *crossing* and *nesting*. A third operator, the *bar operator*, is used to simplify effect specification.

In an analysis-of-variance model, independent variables must be variables that identify classification levels. In the SAS System, these are called *classification* (or *class*) variables and are declared in the **CLASS** statement. (They can also be called *categorical*, *qualitative*, *discrete*, or *nominal variables*.) Classification variables can be either *numeric* or *character*. The values of a classification variable are called *levels*. For example, the classification variable Sex has the levels “male” and “female.”

In a model, an independent variable that is not declared in the **CLASS** statement is assumed to be continuous. Continuous variables, which must be numeric, are used for response variables and covariates. For example, the heights and weights of subjects are continuous variables.

Types of Effects

There are seven different types of effects used in the GLM procedure. In the following list, assume that A, B, C, D, and E are **CLASS** variables and that X1, X2, and Y are continuous variables:

- Regressor effects are specified by writing continuous variables by themselves: X1 X2.
- Polynomial effects are specified by joining two or more continuous variables with asterisks: X1*X1 X1*X2.
- Main effects are specified by writing CLASS variables by themselves: A B C.
- Crossed effects (interactions) are specified by joining classification variables with asterisks: A*B B*C A*B*C.
- Nested effects are specified by following a main effect or crossed effect with a classification variable or list of classification variables enclosed in parentheses. The main effect or crossed effect is nested within the effects listed in parentheses: B(A) C(B*A) D*E(C*B*A). In this example, B(A) is read “B nested within A.”
- Continuous-by-class effects are written by joining continuous variables and classification variables with asterisks: X1*A.
- Continuous-nesting-class effects consist of continuous variables followed by a classification variable interaction enclosed in parentheses: X1(A) X1*X2(A*B).

One example of the general form of an effect involving several variables is

$$X1*X2*A*B*C(D*E)$$

This example contains crossed continuous terms by crossed classification terms nested within more than one classification variable. The continuous list comes first, followed by the crossed list, followed by the nesting list in parentheses. Note that asterisks can appear within the nested list but not immediately before the left parenthesis. For details on how the design matrix and parameters are defined with respect to the effects specified in this section, see the section “[Parameterization of PROC GLM Models](#)” on page 3673.

The **MODEL** statement and several other statements use these effects. Some examples of **MODEL** statements that use various kinds of effects are shown in the following table; a, b, and c represent classification variables, and y, y1, y2, x, and z represent continuous variables.

Specification	Type of Model
<code>model y=x;</code>	Simple regression
<code>model y=x z;</code>	Multiple regression
<code>model y=x x*x;</code>	Polynomial regression
<code>model y1 y2=x z;</code>	Multivariate regression
<code>model y=a;</code>	One-way ANOVA
<code>model y=a b c;</code>	Main-effects ANOVA
<code>model y=a b a*b;</code>	Factorial ANOVA with interaction
<code>model y=a b(a) c(b a);</code>	Nested ANOVA
<code>model y1 y2=a b;</code>	Multivariate analysis of variance (MANOVA)
<code>model y=a x;</code>	Analysis of covariance
<code>model y=a x(a);</code>	Separate-slopes regression
<code>model y=a x x*a;</code>	Homogeneity-of-slopes regression

The Bar Operator

You can shorten the specification of a large factorial model by using the bar operator. For example, two ways of writing the model for a full three-way factorial model follow:

```
model Y = A B C A*B A*C B*C A*B*C;
```

```
model Y = A|B|C;
```

When the bar (|) is used, the right and left sides become effects, and the cross of them becomes an effect. Multiple bars are permitted. The expressions are expanded from left to right, using rules 2–4 given in Searle (1971, p. 390).

- Multiple bars are evaluated from left to right. For instance, A|B|C is evaluated as follows:

$$\begin{aligned}
 A|B|C &\rightarrow \{A|B\}|C \\
 &\rightarrow \{A B A*B\}|C \\
 &\rightarrow A B A*B C A*C B*C A*B*C
 \end{aligned}$$

- Crossed and nested groups of variables are combined. For example, A(B) | C(D) generates A*C(B D), among other terms.
- Duplicate variables are removed. For example, A(C) | B(C) generates A*B(C C), among other terms, and the extra C is removed.
- Effects are discarded if a variable occurs on both the crossed and nested parts of an effect. For instance, A(B) | B(D E) generates A*B(B D E), but this effect is eliminated immediately.

You can also specify the maximum number of variables involved in any effect that results from bar evaluation by specifying that maximum number, preceded by an @ sign, at the end of the bar effect. For example, the

specification $A \mid B \mid C@2$ would result in only those effects that contain 2 or fewer variables: in this case, $A \ B \ A*B \ C \ A*C$ and $B*C$.

More examples of using the bar and at operators follow:

$A \mid C(B)$	is equivalent to	$A \ C(B) \ A*C(B)$
$A(B) \mid C(B)$	is equivalent to	$A(B) \ C(B) \ A*C(B)$
$A(B) \mid B(D \ E)$	is equivalent to	$A(B) \ B(D \ E)$
$A \mid B(A) \mid C$	is equivalent to	$A \ B(A) \ C \ A*C \ B*C(A)$
$A \mid B(A) \mid C@2$	is equivalent to	$A \ B(A) \ C \ A*C$
$A \mid B \mid C \mid D@2$	is equivalent to	$A \ B \ A*B \ C \ A*C \ B*C \ D \ A*D \ B*D \ C*D$
$A*B(C*D)$	is equivalent to	$A*B(C \ D)$

Using PROC GLM Interactively

You can use the GLM procedure interactively. After you specify a model with a [MODEL](#) statement and run PROC GLM with a RUN statement, you can execute a variety of statements without reinvoking PROC GLM.

The section “[Syntax: GLM Procedure](#)” on page 3622 describes which statements can be used interactively. These interactive statements can be executed singly or in groups by following the single statement or group of statements with a RUN statement. Note that the [MODEL](#) statement cannot be repeated; PROC GLM allows only one [MODEL](#) statement.

If you use PROC GLM interactively, you can end the GLM procedure with a DATA step, another PROC step, an ENDSAS statement, or a QUIT statement.

When you are using PROC GLM interactively, additional RUN statements do not end the procedure but tell PROC GLM to execute additional statements.

When you specify a WHERE statement with PROC GLM, it should appear before the first RUN statement. The WHERE statement enables you to select only certain observations for analysis without using a subsetting DATA step. For example, **where group ne 5** omits observations with GROUP=5 from the analysis. See *SAS Statements: Reference* for details on this statement.

When you specify a BY statement with PROC GLM, interactive processing is not possible; that is, once the first RUN statement is encountered, processing proceeds for each BY group in the data set, and no further statements are accepted by the procedure.

Interactivity is also disabled when there are different patterns of missing values among the dependent variables. For details, see the section “[Missing Values](#)” on page 3723.

Parameterization of PROC GLM Models

The GLM procedure constructs a linear model according to the specifications in the [MODEL](#) statement. Each effect generates one or more columns in a design matrix **X**. This section shows precisely how **X** is built.

Intercept

All models include a column of 1s by default to estimate an intercept parameter μ . You can use the [NOINT](#) option to suppress the intercept.

Regression Effects

Regression effects (covariates) have the values of the variables copied into the design matrix directly. Polynomial terms are multiplied out and then installed in **X**.

Main Effects

If a classification variable has m levels, PROC GLM generates m columns in the design matrix for its main effect. Each column is an indicator variable for one of the levels of the classification variable. The default order of the columns is the sort order of the values of their levels; this order can be controlled with the [ORDER=](#) option in the [PROC GLM](#) statement, as shown in the following table.

Data		Design Matrix					
A	B	μ	A		B		
			A1	A2	B1	B2	B3
1	1	1	1	0	1	0	0
1	2	1	1	0	0	1	0
1	3	1	1	0	0	0	1
2	1	1	0	1	1	0	0
2	2	1	0	1	0	1	0
2	3	1	0	1	0	0	1

There are more columns for these effects than there are degrees of freedom for them; in other words, PROC GLM is using an over-parameterized model.

Crossed Effects

First, PROC GLM reorders the terms to correspond to the order of the variables in the [CLASS](#) statement; thus, $B*A$ becomes $A*B$ if A precedes B in the [CLASS](#) statement. Then, PROC GLM generates columns for all combinations of levels that occur in the data. The order of the columns is such that the rightmost variables in the cross index faster than the leftmost variables. No columns are generated corresponding to combinations of levels that do not occur in the data.

Data		Design Matrix											
A	B	μ	A		B			A*B					
			A1	A2	B1	B2	B3	A1B1	A1B2	A1B3	A2B1	A2B2	A2B3
1	1	1	1	0	1	0	0	1	0	0	0	0	0
1	2	1	1	0	0	1	0	0	1	0	0	0	0
1	3	1	1	0	0	0	1	0	0	1	0	0	0
2	1	1	0	1	1	0	0	0	0	0	1	0	0
2	2	1	0	1	0	1	0	0	0	0	0	1	0
2	3	1	0	1	0	0	1	0	0	0	0	0	1

In this matrix, main-effects columns are not linearly independent of crossed-effect columns; in fact, the column space for the crossed effects contains the space of the main effect.

Nested Effects

Nested effects are generated in the same manner as crossed effects. Hence, the design columns generated by the following statements are the same (but the ordering of the columns is different):

```
model y=a b(a);      (B nested within A)
model y=a a*b;       (omitted main effect for B)
```

The nesting operator in PROC GLM is more a notational convenience than an operation distinct from crossing. Nested effects are characterized by the property that the nested variables never appear as main effects. The order of the variables within nesting parentheses is made to correspond to the order of these variables in the **CLASS** statement. The order of the columns is such that variables outside the parentheses index faster than those inside the parentheses, and the rightmost nested variables index faster than the leftmost variables.

Data			Design Matrix							
A	B	μ	A		B(A)					
			A1	A2	B1A1	B2A1	B3A1	B1A2	B2A2	B3A2
1	1	1	1	0	1	0	0	0	0	0
1	2	1	1	0	0	1	0	0	0	0
1	3	1	1	0	0	0	1	0	0	0
2	1	1	0	1	0	0	0	1	0	0
2	2	1	0	1	0	0	0	0	1	0
2	3	1	0	1	0	0	0	0	0	1

Continuous-Nesting-Class Effects

When a continuous variable nests with a classification variable, the design columns are constructed by multiplying the continuous values into the design columns for the class effect.

Data		Design Matrix				
X	A	μ	A		X(A)	
			A1	A2	X(A1)	X(A2)
21	1	1	1	0	21	0
24	1	1	1	0	24	0
22	1	1	1	0	22	0
28	2	1	0	1	0	28
19	2	1	0	1	0	19
23	2	1	0	1	0	23

This model estimates a separate slope for X within each level of A.

Continuous-by-Class Effects

Continuous-by-class effects generate the same design columns as continuous-nesting-class effects. The two models differ by the presence of the continuous variable as a regressor by itself, in addition to being a contributor to $X*A$.

Data		Design Matrix					
X	A	μ	X	A		$X*A$	
				A1	A2	$X*A1$	$X*A2$
21	1	1	21	1	0	21	0
24	1	1	24	1	0	24	0
22	1	1	22	1	0	22	0
28	2	1	28	0	1	0	28
19	2	1	19	0	1	0	19
23	2	1	23	0	1	0	23

Continuous-by-class effects are used to test the homogeneity of slopes. If the continuous-by-class effect is nonsignificant, the effect can be removed so that the response with respect to X is the same for all levels of the classification variables.

General Effects

An example that combines all the effects is

$$X1*X2*A*B*C(D\ E)$$

The continuous list comes first, followed by the crossed list, followed by the nested list in parentheses.

The sequencing of parameters is important to learn if you use the **CONTRAST** or **ESTIMATE** statement to compute or test some linear function of the parameter estimates.

Effects might be retitled by PROC GLM to correspond to ordering rules. For example, $B*A(E\ D)$ might be retitled $A*B(D\ E)$ to satisfy the following:

- Classification variables that occur outside parentheses (crossed effects) are sorted in the order in which they appear in the **CLASS** statement.
- Variables within parentheses (nested effects) are sorted in the order in which they appear in a **CLASS** statement.

The sequencing of the parameters generated by an effect can be described by which variables have their levels indexed faster:

- Variables in the crossed part index faster than variables in the nested list.
- Within a crossed or nested list, variables to the right index faster than variables to the left.

For example, suppose a model includes four effects— A , B , C , and D —each having two levels, 1 and 2. If the **CLASS** statement is

```
class A B C D;
```

then the order of the parameters for the effect B*A(C D), which is retitled A*B(C D), is as follows.

```
A1 B1 C1 D1
A1 B2 C1 D1
A2 B1 C1 D1
A2 B2 C1 D1
A1 B1 C1 D2
A1 B2 C1 D2
A2 B1 C1 D2
A2 B2 C1 D2
A1 B1 C2 D1
A1 B2 C2 D1
A2 B1 C2 D1
A2 B2 C2 D1
A1 B1 C2 D2
A1 B2 C2 D2
A2 B1 C2 D2
A2 B2 C2 D2
```

Note that first the crossed effects B and A are sorted in the order in which they appear in the CLASS statement so that A precedes B in the parameter list. Then, for each combination of the nested effects in turn, combinations of A and B appear. The B effect changes fastest because it is rightmost in the (renamed) cross list. Then A changes next fastest. The D effect changes next fastest, and C is the slowest since it is leftmost in the nested list.

When numeric classification variables are used, their levels are sorted by their character format, which might not correspond to their numeric sort sequence. Therefore, it is advisable to include a format for numeric classification variables or to use the **ORDER=INTERNAL** option in the **PROC GLM** statement to ensure that levels are sorted by their internal values.

Degrees of Freedom

For models with classification (categorical) effects, there are more design columns constructed than there are degrees of freedom for the effect. Thus, there are linear dependencies among the columns. In this event, the parameters are not jointly estimable; there is an infinite number of least squares solutions. The GLM procedure uses a generalized g_2 -inverse to obtain values for the estimates; see the section “**Computational Method**” on page 3726 for more details. The solution values are not produced unless the **SOLUTION** option is specified in the **MODEL** statement. The solution has the characteristic that estimates are zero whenever the design column for that parameter is a linear combination of previous columns. (Strictly termed, the solution values should not be called estimates, since the parameters might not be formally estimable.) With this full parameterization, hypothesis tests are constructed to test linear functions of the parameters that are estimable.

Other procedures (such as the CATMOD procedure) reparameterize models to full rank by using certain restrictions on the parameters. PROC GLM does not reparameterize, making the hypotheses that are commonly tested more understandable. See Goodnight (1978a) for additional reasons for not reparameterizing.

PROC GLM does not actually construct the entire design matrix **X**; rather, a row x_i of **X** is constructed for each observation in the data set and used to accumulate the crossproduct matrix $\mathbf{X}'\mathbf{X} = \sum_i x_i'x_i$.

Hypothesis Testing in PROC GLM

See Chapter 15, “The Four Types of Estimable Functions,” for a complete discussion of the four standard types of hypothesis tests.

Example

To illustrate the four types of tests and the principles upon which they are based, consider a two-way design with interaction based on the following data:

		B	
		1	2
A	1	23.5 23.7	28.7
	2	8.9	5.6 8.9
	3	10.3 12.5	13.6 14.6

Invoke PROC GLM and specify all the estimable functions options to examine what the GLM procedure can test. The following statements produce the summary ANOVA table displayed in [Figure 47.10](#).

```
data example;
  input a b y @@;
  datalines;
1 1 23.5  1 1 23.7  1 2 28.7  2 1  8.9  2 2  5.6
2 2  8.9  3 1 10.3  3 1 12.5  3 2 13.6  3 2 14.6
;

proc glm;
  class a b;
  model y=a b a*b / e e1 e2 e3 e4;
run;
```

Figure 47.10 Summary ANOVA Table from PROC GLM

The GLM Procedure					
Dependent Variable: y					
Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	5	520.4760000	104.0952000	49.66	0.0011
Error	4	8.3850000	2.0962500		
Corrected Total	9	528.8610000			

R-Square	Coeff Var	Root MSE	y Mean
0.984145	9.633022	1.447843	15.03000

The following sections show the general form of estimable functions and discuss the four standard tests, their properties, and abbreviated output for the two-way crossed example.

Estimability

Figure 47.11 is the general form of estimable functions for the example. In order to be testable, a hypothesis must be able to fit within the framework displayed here.

Figure 47.11 General Form of Estimable Functions

The GLM Procedure		
General Form of Estimable Functions		
Effect	Coefficients	
Intercept	L1	
a	1	L2
a	2	L3
a	3	L1-L2-L3
b	1	L5
b	2	L1-L5
a*b	1 1	L7
a*b	1 2	L2-L7
a*b	2 1	L9
a*b	2 2	L3-L9
a*b	3 1	L5-L7-L9
a*b	3 2	L1-L2-L3-L5+L7+L9

If a hypothesis is estimable, the **Ls** in the preceding scheme can be set to values that match the hypothesis. All the standard tests in PROC GLM can be shown in the preceding format, with some of the **Ls** zeroed and some set to functions of other **Ls**.

The following sections show how many of the hypotheses can be tested by comparing the model sum-of-squares regression from one model to a submodel. The notation used is

$$SS(B \text{ effects} | A \text{ effects}) = SS(B \text{ effects}, A \text{ effects}) - SS(A \text{ effects})$$

where $SS(A \text{ effects})$ denotes the regression model sum of squares for the model consisting of $A \text{ effects}$. This notation is equivalent to the reduction notation defined by Searle (1971) and summarized in Chapter 15, “The Four Types of Estimable Functions.”

Type I Tests

Type I sums of squares (SS), also called *sequential sums of squares*, are the incremental improvement in error sums of squares as each effect is added to the model. They can be computed by fitting the model in steps and recording the difference in error sum of squares at each step.

Source	Type I SS
A	$SS(A \mu)$
B	$SS(B \mu, A)$
$A * B$	$SS(A * B \mu, A, B)$

Type I sums of squares are displayed by default because they are easy to obtain and can be used in various hand calculations to produce sum of squares values for a series of different models. Nelder (1994) and others have argued that Type I and II sums are essentially the only appropriate ones for testing ANOVA effects; however, see also the discussion of Nelder's article, especially Rodriguez, Tobias, and Wolfinger (1995) and Searle (1995).

The Type I hypotheses have these properties:

- Type I sum of squares for all effects add up to the model sum of squares. None of the other sum of squares types have this property, except in special cases.
- Type I hypotheses can be derived from rows of the Forward-Dolittle transformation of $\mathbf{X}'\mathbf{X}$ (a transformation that reduces $\mathbf{X}'\mathbf{X}$ to an upper triangular matrix by row operations).
- Type I sum of squares are statistically independent of each other under the usual [assumption](#) that the true residual errors are independent and identically normally distributed (see page [3670](#)).
- Type I hypotheses depend on the order in which effects are specified in the **MODEL** statement.
- Type I hypotheses are uncontaminated by parameters corresponding to effects that precede the effect being tested; however, the hypotheses usually involve parameters for effects following the tested effect in the model. For example, in the model

$$\mathbf{Y} = \mathbf{A} \mathbf{B};$$

the Type I hypothesis for B does not involve A parameters, but the Type I hypothesis for A does involve B parameters.

- Type I hypotheses are functions of the cell counts for unbalanced data; the hypotheses are not usually the same hypotheses that are tested if the data are balanced.
- Type I sums of squares are useful for polynomial models where you want to know the contribution of a term as though it had been made orthogonal to preceding effects. Thus, in polynomial models, Type I sums of squares correspond to tests of the orthogonal polynomial effects.

The Type I estimable functions and associated tests for the example are shown in [Figure 47.12](#).

Figure 47.12 Type I Estimable Functions and Tests

Type I Estimable Functions					
Coefficients					
Effect		a	b	a*b	
Intercept		0	0	0	
a	1	L2	0	0	
a	2	L3	0	0	
a	3	-L2-L3	0	0	
b	1	0.1667*L2-0.1667*L3	L5	0	
b	2	-0.1667*L2+0.1667*L3	-L5	0	
a*b	1 1	0.6667*L2	0.2857*L5	L7	
a*b	1 2	0.3333*L2	-0.2857*L5	-L7	
a*b	2 1	0.3333*L3	0.2857*L5	L9	
a*b	2 2	0.6667*L3	-0.2857*L5	-L9	
a*b	3 1	-0.5*L2-0.5*L3	0.4286*L5	-L7-L9	
a*b	3 2	-0.5*L2-0.5*L3	-0.4286*L5	L7+L9	

Source	DF	Type I SS	Mean Square	F Value	Pr > F
a	2	494.0310000	247.0155000	117.84	0.0003
b	1	10.7142857	10.7142857	5.11	0.0866
a*b	2	15.7307143	7.8653571	3.75	0.1209

Type II Tests

The Type II tests can also be calculated by comparing the error sums of squares (SS) for subset models. The Type II SS are the reduction in error SS due to adding the term after all other terms have been added to the model except terms that contain the effect being tested. An effect is contained in another effect if it can be derived by deleting variables from the latter effect. For example, A and B are both contained in A*B. For this model, the Type II SS are given by the reduced sums of squares as shown in the following table.

Source	Type II SS
A	$SS(A \mid \mu, B)$
B	$SS(B \mid \mu, A)$
A * B	$SS(A * B \mid \mu, A, B)$

Type II SS have these properties:

- Type II SS do not necessarily sum to the model SS.
- The hypothesis for an effect does not involve parameters of other effects except for containing effects (which it must involve to be estimable).
- Type II SS are invariant to the ordering of effects in the model.
- For unbalanced designs, Type II hypotheses for effects that are contained in other effects are not usually the same hypotheses that are tested if the data are balanced. The hypotheses are generally functions of the cell counts.

The Type II estimable functions and associated tests for the example are shown in Figure 47.13.

Figure 47.13 Type II Estimable Functions and Tests

Type II Estimable Functions					
Coefficients					
Effect		a	b	a*b	
Intercept		0	0	0	
a	1	L2	0	0	
a	2	L3	0	0	
a	3	-L2-L3	0	0	
b	1	0	L5	0	
b	2	0	-L5	0	
a*b	1 1	0.619*L2+0.0476*L3	0.2857*L5	L7	
a*b	1 2	0.381*L2-0.0476*L3	-0.2857*L5	-L7	
a*b	2 1	-0.0476*L2+0.381*L3	0.2857*L5	L9	
a*b	2 2	0.0476*L2+0.619*L3	-0.2857*L5	-L9	
a*b	3 1	-0.5714*L2-0.4286*L3	0.4286*L5	-L7-L9	
a*b	3 2	-0.4286*L2-0.5714*L3	-0.4286*L5	L7+L9	

Source	DF	Type II SS	Mean Square	F Value	Pr > F
a	2	499.1202857	249.5601429	119.05	0.0003
b	1	10.7142857	10.7142857	5.11	0.0866
a*b	2	15.7307143	7.8653571	3.75	0.1209

Type III and Type IV Tests

Type III and Type IV sums of squares (SS), sometimes referred to as *partial sums of squares*, are considered by many to be the most desirable; see Searle (1987, Section 4.6). Using PROC GLM's singular parameterization, these SS cannot, in general, be computed by comparing model SS from different models. However, they can sometimes be computed by reduction for methods that reparameterize to full rank, when such a reparameterization effectively imposes Type III linear constraints on the parameters. In PROC GLM, they are computed by constructing a hypothesis matrix $\mathbf{L}$ and then computing the SS associated with the hypothesis $\mathbf{L}\boldsymbol{\beta} = 0$. As long as there are no missing cells in the design, Type III and Type IV SS are the same.

These are properties of Type III and Type IV SS:

- The hypothesis for an effect does not involve parameters of other effects except for containing effects (which it must involve to be estimable).
- The hypotheses to be tested are invariant to the ordering of effects in the model.
- The hypotheses are the same hypotheses that are tested if there are no missing cells. They are not functions of cell counts.
- The SS do not generally add up to the model SS and, in some cases, can exceed the model SS.

The SS are constructed from the general form of estimable functions. Type III and Type IV tests are different only if the design has missing cells. In this case, the Type III tests have an orthogonality property, while the

Type IV tests have a balancing property. These properties are discussed in Chapter 15, “[The Four Types of Estimable Functions](#).” For this example, since the data contain observations for all pairs of levels of A and B, Type IV tests are identical to the Type III tests that are shown in [Figure 47.14](#). (This combines tables from several pages of output.)

Figure 47.14 Type III Estimable Functions and Tests

Type III Estimable Functions					
Coefficients					
Effect		a	b	a*b	
Intercept		0	0	0	
a	1	L2	0	0	
a	2	L3	0	0	
a	3	-L2-L3	0	0	
b	1	0	L5	0	
b	2	0	-L5	0	
a*b	1 1	0.5*L2	0.3333*L5	L7	
a*b	1 2	0.5*L2	-0.3333*L5	-L7	
a*b	2 1	0.5*L3	0.3333*L5	L9	
a*b	2 2	0.5*L3	-0.3333*L5	-L9	
a*b	3 1	-0.5*L2-0.5*L3	0.3333*L5	-L7-L9	
a*b	3 2	-0.5*L2-0.5*L3	-0.3333*L5	L7+L9	

Source	DF	Type III SS	Mean Square	F Value	Pr > F
a	2	479.1078571	239.5539286	114.28	0.0003
b	1	9.4556250	9.4556250	4.51	0.1009
a*b	2	15.7307143	7.8653571	3.75	0.1209

Effect Size Measures for *F* Tests in GLM

A significant *F* test in a linear model indicates that the effect of the term or contrast being tested might be real. The next thing you want to know is, How big is the effect? Various measures have been devised to give answers to this question that are comparable over different experimental designs. If you specify the **EFFECTSIZE** option in the **MODEL** statement, then GLM adds to each ANOVA table estimates and confidence intervals for three different measures of effect size:

- the noncentrality parameter for the *F* test
- the proportion of total variation accounted for (also known as the semipartial correlation ratio or the squared semipartial correlation)
- the proportion of partial variation accounted for (also known as the full partial correlation ratio or the squared full partial correlation)

The adjectives “semipartial” and “full partial” might seem strange. They refer to how other effects are “partialed out” of the dependent variable and the effect being tested. For “semipartial” statistics, all other effects are partialed out of the effect in question, but not the dependent variable. This measures the (adjusted)

effect as a proportion of the total variation in the dependent variable. On the other hand, for “full partial” statistics, all other effects are partialled out of *both* the dependent variable and the effect in question. This measures the (adjusted) effect as a proportion of only the dependent variation remaining after partialing, or in other words the partial variation. Details about the computation and interpretation of these estimates and confidence intervals are discussed in the remainder of this section.

The noncentrality parameter is directly related to the true distribution of the F statistic when the effect being tested has a non-null effect. The uniformly minimum variance unbiased estimate for the noncentrality is

$$NC_{UMVUE} = \frac{DF(DFE - 2)FValue}{DFE} - DF$$

where $FValue$ is the observed value of the F statistic for the test and DF and DFE are the numerator and denominator degrees of freedom for the test, respectively. An alternative estimate that can be slightly biased but has a somewhat lower expected mean square error is

$$NC_{minMSE} = \frac{DF(DFE - 4)FValue}{DFE} - \frac{DF(DFE - 4)}{DFE - 2}$$

(See Perlman and Rasmussen (1975), cited in Johnson, Kotz, and Balakrishnan (1994).) A $p \times 100\%$ lower confidence bound for the noncentrality is given by the value of NC for which $\text{probf}(FValue, DF, DFE, NC) = p$, where $\text{probf}()$ is the cumulative probability function for the non-central F distribution. This result can be used to form a $(1 - \alpha) \times 100\%$ confidence interval for the noncentrality.

The partial proportion of variation accounted for by the effect being tested is easiest to define by its natural sample estimate,

$$\hat{\eta}_{partial}^2 = \frac{SS}{SS + SSE}$$

where SSE is the sample error sum of squares. Note that $\hat{\eta}_{partial}^2$ is actually sometimes denoted $R_{partial}^2$ or just R square, but in this context the R square notation is reserved for the $\hat{\eta}^2$ corresponding to the overall model, which is just the familiar R square for the model. $\hat{\eta}_{partial}^2$ is actually a biased estimate of the true $\eta_{partial}^2$; an alternative that is approximately unbiased is given by

$$\omega_{partial}^2 = \frac{SS - DF \times MSE}{SS + (N - DF)MSE}$$

where $MSE = SSE/DFE$ is the sample mean square for error and N is the number of observations. The true $\eta_{partial}^2$ is related to the true noncentrality parameter NC by the formula

$$\eta_{partial}^2 = \frac{NC}{NC + N}$$

This fact can be employed to transform a confidence interval for NC into one for $\eta_{partial}^2$. Note that some authors (Steiger and Fouladi 1997; Fidler and Thompson 2001; Smithson 2003) have published slightly different confidence intervals for $\eta_{partial}^2$, based on a slightly different formula for the relationship between $\eta_{partial}^2$ and NC , apparently due to Cohen (1988). Cohen’s formula appears to be approximately correct for

random predictor values (Maxwell 2000), but the one given previously is correct if the predictor values are assumed fixed, as is standard for the GLM procedure.

Finally, the proportion of total variation accounted for by the effect being tested is again easiest to define by its natural sample estimate, which is known as the (*semipartial*) $\hat{\eta}^2$ statistic,

$$\hat{\eta}^2 = \frac{SS}{SS_{total}}$$

where SS_{total} is the total sample (corrected) sum of squares, and SS is the observed sum of squares due to the effect being tested. As with $\hat{\eta}_{partial}^2$, $\hat{\eta}^2$ is actually a biased estimate of the true η^2 ; an alternative that is approximately unbiased is the (*semipartial*) ω^2 statistic

$$\omega^2 = \frac{SS - DF \times MSE}{SS_{total} + MSE}$$

where $MSE = SSE/DFE$ is the sample mean square for error. Whereas $\eta_{partial}^2$ depends only on the noncentrality for its associated F test, the presence of the total sum of squares in the previous formulas indicates that η^2 depends on the noncentralities for all effects in the model. An exact confidence interval is not available, but if you write the formula for $\hat{\eta}^2$ as

$$\hat{\eta}^2 = \frac{SS}{SS + (SS_{total} - SS)}$$

then a conservative confidence interval can be constructed as for $\eta_{partial}^2$, treating $SS_{total} - SS$ as the SSE and $N - DF - 1$ as the DFE (Smithson 2004). This confidence interval is conservative in the sense that it implies values of the true η^2 that are smaller than they should be.

Estimates and confidence intervals for effect sizes require some care in interpretation. For example, while the true proportions of total and partial variation accounted for are nonnegative quantities, their estimates might be less than zero. Also, confidence intervals for effect sizes are not directly related to the corresponding estimates. In particular, it is possible for the estimate to lie outside the confidence interval.

As for interpreting the actual values of effect size measures, the approximately unbiased ω^2 estimates are usually preferred for point estimates. Some authors have proposed certain ranges as indicating “small,” “medium,” and “large” effects (Cohen 1988), but general benchmarks like this depend on the nature of the data and the typical signal-to-noise ratio; they should not be expected to apply across various disciplines. For example, while an ω^2 value of 10% might be viewed as “large” for psychometric data, it can be a relatively small effect for industrial experimentation. Whatever the standard, confidence intervals for true effect sizes typically span more than one category, indicating that in small experiments, it can be difficult to make firm statements about the size of effects.

Example

The data for this example are similar to data analyzed in Steiger and Fouladi (1997); Fidler and Thompson (2001); Smithson (2003). Consider the following hypothetical design, testing 28 men and 28 women on seven different tasks.

```

data Test;
  do Task = 1 to 7;
    do Gender = 'M', 'F';
      do i = 1 to 4;
        input Response @@;
        output;
      end;
    end;
  end;
  datalines;
7.1 2.8 3.9 3.7      6.5 6.5 6.5 6.6
7.1 5.5 4.8 2.6      3.6 5.4 5.6 4.5
7.2 4.6 4.9 4.6      3.3 5.4 2.8 1.5
5.6 6.2 5.4 6.5      5.6 2.7 3.8 2.3
2.2 5.4 5.6 8.4      1.2 2.0 4.3 4.6
9.1 4.5 7.6 4.9      4.3 7.7 6.5 7.7
4.5 3.8 5.9 6.1      1.7 2.5 4.3 2.7
;

```

This is a balanced two-way design with four replicates per cell. The following statements analyze this data. Since this is a balanced design, you can use the **SS1** option in the **MODEL** statement to display only the Type I sums of squares.

```

proc glm data=Test;
  class Gender Task;
  model Response = Gender|Task / ss1;
run;

```

The analysis of variance results are shown in Figure 47.15.

Figure 47.15 Two-Way Analysis of Variance

The GLM Procedure

Dependent Variable: Response

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Gender	1	14.40285714	14.40285714	6.00	0.0185
Task	6	38.15964286	6.35994048	2.65	0.0285
Gender*Task	6	35.99964286	5.99994048	2.50	0.0369

You can see that the two main effects as well as their interaction are all significant. Suppose you want to compare the main effect of Gender with the interaction between Gender and Task. The sums of squares for the interaction are more than twice as large, but it's not clear how experimental variability might affect this. The following statements perform the same analysis as before, but add the **EFFECTSIZE** option to the **MODEL** statement; also, with **ALPHA=0.1** option displays 90% confidence intervals, ensuring that inferences based on the *p*-values at the 0.05 levels will agree with the lower confidence limit.

```

proc glm data=Test;
  class Gender Task;
  model Response = Gender|Task / ss1 effectsize alpha=0.1;
run;

```

The Type I analysis of variance results with added effect size information are shown in Figure 47.16.

Figure 47.16 Two-Way Analysis of Variance with Effect Sizes

The GLM Procedure

Dependent Variable: Response

Source	DF	Type I SS	Mean Square	F Value	Pr > F	Noncentrality Parameter			
						Min Var Unbiased Estimate	Low MSE Estimate	90% Confidence Limits	
Gender	1	14.40285714	14.40285714	6.00	0.0185	4.72	4.48	0.521	17.1
Task	6	38.15964286	6.35994048	2.65	0.0285	9.14	8.69	0.870	27.3
Gender*Task	6	35.99964286	5.99994048	2.50	0.0369	8.29	7.87	0.463	25.9

Source	Total Variation Accounted For				Partial Variation Accounted For			
	Semipartial Eta-Square	Semipartial Omega-Square	Conservative 90% Confidence Limits		Partial Eta-Square	Partial Omega-Square	90% Confidence Limits	
Gender	0.0761	0.0626	0.0019	0.2030	0.1250	0.0820	0.0092	0.2342
Task	0.2015	0.1239	0.0000	0.2772	0.2746	0.1502	0.0153	0.3277
Gender*Task	0.1901	0.1126	0.0000	0.2639	0.2632	0.1385	0.0082	0.3160

The estimated effect sizes for Gender and the interaction all tell pretty much the same story: the effect of the interaction is appreciably greater than the effect of Gender. However, the confidence intervals suggest that this inference should be treated with some caution, since the lower confidence bound for the Gender effect is greater than the lower confidence bound for the interaction in all three cases. Follow-up testing is probably in order, using the estimated effect sizes in this preliminary study to design a large enough sample to distinguish the sizes of the effects.

Absorption

Absorption is a computational technique used to reduce computing resource needs in certain cases. The classic use of absorption occurs when a blocking factor with a large number of levels is a term in the model.

For example, the statements

```
proc glm;
  absorb herd;
  class a b;
  model y=a b a*b;
run;
```

are equivalent to

```
proc glm;
  class herd a b;
  model y=herd a b a*b;
run;
```

The exception to the previous statements is that the Type II, Type III, or Type IV SS for HERD are not computed when HERD is absorbed.

The algorithm for absorbing variables is similar to the one used by the NESTED procedure for computing a nested analysis of variance. As each new row of $[X|Y]$ (corresponding to the nonabsorbed independent effects and the dependent variables) is constructed, it is adjusted for the absorbed effects in a Type I fashion. The efficiency of the absorption technique is due to the fact that this adjustment can be done in one pass of the data and without solving any linear equations, assuming that the data have been sorted by the absorbed variables.

Several effects can be absorbed at one time. For example, these statements

```
proc glm;
  absorb herd cow;
  class a b;
  model y=a b a*b;
run;
```

are equivalent to

```
proc glm;
  class herd cow a b;
  model y=herd cow(herd) a b a*b;
run;
```

When you use absorption, the size of the $X'X$ matrix is a function only of the effects in the **MODEL** statement. The effects being absorbed do not contribute to the size of the $X'X$ matrix.

For the preceding example, a and b can be absorbed:

```
proc glm;
  absorb a b;
  class herd cow;
  model y=herd cow(herd);
run;
```

Although the sources of variation in the results are listed as

```
a b(a) herd cow(herd)
```

all types of estimable functions for herd and cow(herd) are free of a, b, and a*b parameters.

To illustrate the savings in computing by using the **ABSORB** statement, PROC GLM is run on generated data with 1147 degrees of freedom in the model with the following statements.

```
data a;
  do herd=1 to 40;
    do cow=1 to 30;
      do treatment=1 to 3;
        do rep=1 to 2;
          y = herd/5 + cow/10 + treatment + rannor(1);
          output;
        end;
      end;
    end;
  end;
run;
```

```
proc glm data=a;
  class herd cow treatment;
  model y=herd cow(herd) treatment;
run;
```

This analysis would have required over 6 megabytes of memory for the $X'X$ matrix had PROC GLM solved it directly. However, in the following statements, the GLM procedure needs only a 4×4 matrix for the intercept and treatment because the other effects are absorbed.

```
proc glm data=a;
  absorb herd cow;
  class treatment;
  model y = treatment;
run;
```

These statements produce the results shown in Figure 47.17.

Figure 47.17 Absorption of Effects

The GLM Procedure

Class Level Information			
Class	Levels	Values	
treatment	3	1 2 3	

Number of Observations Read 7200

Number of Observations Used 7200

The GLM Procedure

Dependent Variable: y

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	1201	49465.40242	41.18685	41.57	<.0001
Error	5998	5942.23647	0.99070		
Corrected Total	7199	55407.63889			

R-Square	Coeff Var	Root MSE	y Mean
0.892754	13.04236	0.995341	7.631598

Source	DF	Type I SS	Mean Square	F Value	Pr > F
herd	39	38549.18655	988.44068	997.72	<.0001
cow(herd)	1160	6320.18141	5.44843	5.50	<.0001
treatment	2	4596.03446	2298.01723	2319.58	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
treatment	2	4596.034455	2298.017228	2319.58	<.0001

Specification of ESTIMATE Expressions

Consider the model

$$E(Y) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3$$

The corresponding **MODEL** statement for PROC GLM is

```
model y=x1 x2 x3;
```

To estimate the difference between the parameters for x_1 and x_2 ,

$$\beta_1 - \beta_2 = (0 \ 1 \ -1 \ 0) \boldsymbol{\beta}, \text{ where } \boldsymbol{\beta} = (\beta_0 \ \beta_1 \ \beta_2 \ \beta_3)'$$

you can use the following **ESTIMATE** statement:

```
estimate 'B1-B2' x1 1 x2 -1;
```

To predict y at $x_1 = 1$, $x_2 = 0$, and $x_3 = -2$, you can estimate

$$\beta_0 + \beta_1 - 2\beta_3 = (1 \ 1 \ 0 \ -2) \boldsymbol{\beta}$$

with the following **ESTIMATE** statement:

```
estimate 'B0+B1-2B3' intercept 1 x1 1 x3 -2;
```

Now consider models involving classification variables such as

```
model y=A B A*B;
```

with the associated parameters:

$$(\mu \ \alpha_1 \ \alpha_2 \ \alpha_3 \ \beta_1 \ \beta_2 \ \gamma_{11} \ \gamma_{12} \ \gamma_{21} \ \gamma_{22} \ \gamma_{31} \ \gamma_{32})$$

The LS-mean for the first level of **A** is $\mathbf{L}\boldsymbol{\beta}$, where

$$\mathbf{L} = (1 \mid 1 \ 0 \ 0 \mid 0.5 \ 0.5 \mid 0.5 \ 0.5 \ 0 \ 0 \ 0 \ 0)$$

You can estimate this with the following **ESTIMATE** statement:

```
estimate 'LS-mean(A1)' intercept 1 A 1 B 0.5 0.5 A*B 0.5 0.5;
```

Note in this statement that only one element of **L** is specified following the **A** effect, even though **A** has three levels. Whenever the list of constants following an effect name is shorter than the effect's number of levels, zeros are used as the remaining constants. (If the list of constants is longer than the number of levels for the effect, the extra constants are ignored, and a warning message is displayed.)

To estimate the **A** linear effect in the preceding model, assuming equally spaced levels for **A**, you can use the following **L**:

$$\mathbf{L} = (0 \mid -1 \ 0 \ 1 \mid 0 \ 0 \mid -0.5 \ -0.5 \ 0 \ 0 \ 0.5 \ 0.5)$$

The **ESTIMATE** statement for this **L** is written as

```
estimate 'A Linear' A -1 0 1;
```

If you do not specify the elements of **L** for an effect that contains a specified effect, then the elements of the specified effect are equally distributed over the corresponding levels of the higher-order effect. In addition, if you specify the intercept in an **ESTIMATE** or **CONTRAST** statement, it is distributed over all classification effects that are not contained by any other specified effect.

The distribution of lower-order coefficients to higher-order effect coefficients follows the same general rules as in the **LSMEANS** statement, and it is similar to that used to construct Type IV tests. In the previous example, the -1 associated with α_1 is divided by the number n_{1j} of γ_{1j} parameters; then each γ_{1j} coefficient is set to $-1/n_{1j}$. The 1 associated with α_3 is distributed among the γ_{3j} parameters in a similar fashion. In the event that an unspecified effect contains several specified effects, only that specified effect with the most factors in common with the unspecified effect is used for distribution of coefficients to the higher-order effect.

Numerous syntactical expressions for the **ESTIMATE** statement were considered, including many that involved specifying the effect and level information associated with each coefficient. For models involving higher-level effects, the requirement of specifying level information can lead to very bulky specifications. Consequently, the simpler form of the **ESTIMATE** statement described earlier was implemented.

The syntax of this **ESTIMATE** statement puts a burden on you to know a priori the order of the parameter list associated with each effect. You can use the **ORDER=** option in the **PROC GLM** statement to ensure that the levels of the classification effects are sorted appropriately.

NOTE: If you use the **ESTIMATE** statement with unspecified effects, use the **E** option to make sure that the actual **L** constructed by the preceding rules is the one you intended.

A Check for Estimability

Each **L** is checked for estimability using the relationship $\mathbf{L} = \mathbf{LH}$, where $\mathbf{H} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}$. The **L** vector is declared nonestimable, if for any i

$$\text{ABS}(\mathbf{L}_i - (\mathbf{LH})_i) > \begin{cases} \epsilon & \text{if } \mathbf{L}_i = 0 \text{ or} \\ \epsilon \times \text{ABS}(\mathbf{L}_i) & \text{otherwise} \end{cases}$$

where $\epsilon = 10^{-4}$ by default; you can change this with the **SINGULAR=** option. Continued fractions (like $1/3$) should be specified to at least six decimal places, or the **DIVISOR** parameter should be used.

Comparing Groups

An important task in analyzing data with classification effects is to estimate the typical response for each level of a given effect; often, you also want to compare these estimates to determine which levels are equivalent in terms of the response. You can perform this task in two ways with the **GLM** procedure: with direct, arithmetic group means; and with so-called *least squares means* (LS-means).

Means versus LS-Means

Computing and comparing arithmetic means—either simple or weighted within-group averages of the input data—is a familiar and well-studied statistical process. This is the right approach to summarizing and comparing groups for one-way and balanced designs. However, in unbalanced designs with more than one

effect, the arithmetic mean for a group might not accurately reflect the “typical” response for that group, since it does not take other effects into account.

For example, the following analysis of an unbalanced two-way design produces the ANOVA, means, and LS-means shown in Figure 47.18, Figure 47.19, and Figure 47.20.

```
data twoway;
  input Treatment Block y @@;
  datalines;
1 1 17   1 1 28   1 1 19   1 1 21   1 1 19
1 2 43   1 2 30   1 2 39   1 2 44   1 2 44
1 3 16
2 1 21   2 1 21   2 1 24   2 1 25
2 2 39   2 2 45   2 2 42   2 2 47
2 3 19   2 3 22   2 3 16
3 1 22   3 1 30   3 1 33   3 1 31
3 2 46
3 3 26   3 3 31   3 3 26   3 3 33   3 3 29   3 3 25
;

title "Unbalanced Two-way Design";
ods select ModelANOVA Means LSMeans;

proc glm data=twoway;
  class Treatment Block;
  model y = Treatment|Block;
  means Treatment;
  lsmeans Treatment;
run;

ods select all;
```

Figure 47.18 ANOVA Results for Unbalanced Two-Way Design

Unbalanced Two-way Design

The GLM Procedure

Dependent Variable: y

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Treatment	2	8.060606	4.030303	0.24	0.7888
Block	2	2621.864124	1310.932062	77.95	<.0001
Treatment*Block	4	32.684361	8.171090	0.49	0.7460

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Treatment	2	266.130682	133.065341	7.91	0.0023
Block	2	1883.729465	941.864732	56.00	<.0001
Treatment*Block	4	32.684361	8.171090	0.49	0.7460

Figure 47.19 Treatment Means for Unbalanced Two-Way Design**Unbalanced Two-way Design****The GLM Procedure**

		y	
Level of Treatment	N	Mean	Std Dev
1	11	29.0909091	11.5104695
2	11	29.1818182	11.5569735
3	11	30.1818182	6.3058414

Figure 47.20 Treatment LS-means for Unbalanced Two-Way Design**Unbalanced Two-way Design****The GLM Procedure****Least Squares Means**

Treatment	y LSMEAN
1	25.6000000
2	28.3333333
3	34.4444444

No matter how you look at them, these data exhibit a strong effect due to the blocks (F test $p < 0.0001$) and no significant interaction between treatments and blocks (F test $p > 0.7$). But the lack of balance affects how the treatment effect is interpreted: in a main-effects-only model, there are no significant differences between the treatment means themselves (Type I F test $p > 0.7$), but there are highly significant differences between the treatment means corrected for the block effects (Type III F test $p < 0.01$).

LS-means are, in effect, within-group means appropriately adjusted for the other effects in the model. More precisely, they estimate the marginal means for a balanced population (as opposed to the unbalanced design). For this reason, they are also called *estimated population marginal means* by Searle, Speed, and Milliken (1980). In the same way that the Type I F test assesses differences between the arithmetic treatment means (when the treatment effect comes first in the model), the Type III F test assesses differences between the LS-means. Accordingly, for the unbalanced two-way design, the discrepancy between the Type I and Type III tests is reflected in the arithmetic treatment means and treatment LS-means, as shown in Figure 47.19 and Figure 47.20. See the section “Construction of Least Squares Means” on page 3707 for more on LS-means.

Note that, while the arithmetic means are always uncorrelated (under the usual assumptions for analysis of variance; see page 3670), the LS-means might not be. This fact complicates the problem of multiple comparisons for LS-means; see the following section.

Multiple Comparisons

When comparing more than two means, an ANOVA F test tells you whether the means are significantly different from each other, but it does not tell you which means differ from which other means. Multiple-comparison procedures (MCPs), also called *mean separation tests*, give you more detailed information about the differences among the means. The goal in multiple comparisons is to compare the average effects of three or more “treatments” (for example, drugs, groups of subjects) to decide which treatments are better,

which ones are worse, and by how much, while controlling the probability of making an incorrect decision. A variety of multiple-comparison methods are available with the **MEANS** and **LSMEANS** statement in the GLM procedure.

The following classification is due to Hsu (1996). Multiple-comparison procedures can be categorized in two ways: by the comparisons they make and by the strength of inference they provide. With respect to which comparisons are made, the GLM procedure offers two types:

- comparisons between all pairs of means
- comparisons between a control and all other means

The strength of inference says what can be inferred about the structure of the means when a test is significant; it is related to what type of error rate the MCP controls. MCPs available in the GLM procedure provide one of the following types of inference, in order from weakest to strongest:

- Individual: differences between means, unadjusted for multiplicity
- Inhomogeneity: means are different
- Inequalities: which means are different
- Intervals: simultaneous confidence intervals for mean differences

Methods that control only individual error rates are not true MCPs at all. Methods that yield the strongest level of inference, simultaneous confidence intervals, are usually preferred, since they enable you not only to say which means are different but also to put confidence bounds on *how much* they differ, making it easier to assess the practical significance of a difference. They are also less likely to lead nonstatisticians to the invalid conclusion that nonsignificantly different sample means imply equal population means. Interval MCPs are available for both arithmetic means and LS-means via the **MEANS** and **LSMEANS** statements, respectively.¹

Table 47.12 and Table 47.13 display MCPs available in PROC GLM for all pairwise comparisons and comparisons with a control, respectively, along with associated strength of inference and the syntax (when applicable) for both the **MEANS** and the **LSMEANS** statements.

¹The Duncan-Waller method does not fit into the preceding scheme, since it is based on the Bayes risk rather than any particular error rate.

Table 47.12 Multiple-Comparison Procedures for All Pairwise Comparisons

Method	Strength of Inference	Syntax	
		MEANS	LSMEANS
Student's <i>t</i>	Individual	T	PDIFF ADJUST=T
Duncan	Individual	DUNCAN	
Student-Newman-Keuls	Inhomogeneity	SNK	
REGWQ	Inequalities	REGWQ	
Tukey-Kramer	Intervals	TUKEY	PDIFF ADJUST=TUKEY
Bonferroni	Intervals	BON	PDIFF ADJUST=BON
Sidak	Intervals	SIDAK	PDIFF ADJUST=SIDAK
Scheffé	Intervals	SCHEFFE	PDIFF ADJUST=SCHEFFE
SMM	Intervals	SMM	PDIFF ADJUST=SMM
Gabriel	Intervals	GABRIEL	
Simulation	Intervals		PDIFF ADJUST=SIMULATE

Table 47.13 Multiple-Comparison Procedures for Comparisons with a Control

Method	Strength of Inference	Syntax	
		MEANS	LSMEANS
Student's <i>t</i>	Individual		PDIFF=CONTROL ADJUST=T
Dunnett	Intervals	DUNNETT	PDIFF=CONTROL ADJUST=DUNNETT
Bonferroni	Intervals		PDIFF=CONTROL ADJUST=BON
Sidak	Intervals		PDIFF=CONTROL ADJUST=SIDAK
Scheffé	Intervals		PDIFF=CONTROL ADJUST=SCHEFFE
SMM	Intervals		PDIFF=CONTROL ADJUST=SMM
Simulation	Intervals		PDIFF=CONTROL ADJUST=SIMULATE

NOTE: One-sided Dunnett's tests are also available from the **MEANS** statement with the **DUNNETTL** and **DUNNETTU** options and from the **LSMEANS** statement with **PDIFF=CONTROL** and **PDIFF=CONTROLL**.

A note concerning the ODS tables for the results of the **PDIFF** or **TDIFF** options in the **LSMEANS** statement: The *p/t*-values for differences are displayed in columns of the LSMeans table for **PDIFF/TDIFF=CONTROL** or **PDIFF/TDIFF=ANOM**, and for **PDIFF/TDIFF=ALL** when there are only two LS-means. Otherwise (for **PDIFF/TDIFF=ALL** when there are more than two LS-means), the *p/t*-values for differences are displayed in a separate table called Diff.

Details of these multiple comparison methods are given in the following sections.

Pairwise Comparisons

All the methods discussed in this section depend on the standardized pairwise differences $t_{ij} = (\bar{y}_i - \bar{y}_j)/\hat{\sigma}_{ij}$, where the parts of this expression are defined as follows:

- *i* and *j* are the indices of two groups
- $\bar{y}_i$ and $\bar{y}_j$ are the means or LS-means for groups *i* and *j*

- $\hat{\sigma}_{ij}$ is the square root of the estimated variance of $\bar{y}_i - \bar{y}_j$. For simple arithmetic means, $\hat{\sigma}_{ij}^2 = s^2(1/n_i + 1/n_j)$, where n_i and n_j are the sizes of groups i and j , respectively, and s^2 is the mean square for error, with ν degrees of freedom. For weighted arithmetic means, $\hat{\sigma}_{ij}^2 = s^2(1/w_i + 1/w_j)$, where w_i and w_j are the sums of the weights in groups i and j , respectively. Finally, for LS-means defined by the linear combinations $\mathbf{l}_i' \mathbf{b}$ and $\mathbf{l}_j' \mathbf{b}$ of the parameter estimates, $\hat{\sigma}_{ij}^2 = s^2 \mathbf{l}_i' (\mathbf{X}'\mathbf{X})^{-1} \mathbf{l}_j$.

Furthermore, all of the methods are discussed in terms of significance tests of the form

$$|t_{ij}| \geq c(\alpha)$$

where $c(\alpha)$ is some constant depending on the significance level. Such tests can be inverted to form confidence intervals of the form

$$(\bar{y}_i - \bar{y}_j) - \hat{\sigma}_{ij}c(\alpha) \leq \mu_i - \mu_j \leq (\bar{y}_i - \bar{y}_j) + \hat{\sigma}_{ij}c(\alpha)$$

The simplest approach to multiple comparisons is to do a t test on every pair of means (the **T** option in the **MEANS** statement, **ADJUST=T** in the **LSMEANS** statement). For the i th and j th means, you can reject the null hypothesis that the population means are equal if

$$|t_{ij}| \geq t(\alpha; \nu)$$

where α is the significance level, ν is the number of error degrees of freedom, and $t(\alpha; \nu)$ is the two-tailed critical value from a Student's t distribution. If the cell sizes are all equal to, say, n , the preceding formula can be rearranged to give

$$|\bar{y}_i - \bar{y}_j| \geq t(\alpha; \nu)s\sqrt{\frac{2}{n}}$$

the value of the right-hand side being Fisher's least significant difference (LSD).

There is a problem with repeated t tests, however. Suppose there are 10 means and each t test is performed at the 0.05 level. There are $10(10 - 1)/2 = 45$ pairs of means to compare, each with a 0.05 probability of a type 1 error (a false rejection of the null hypothesis). The chance of making at least one type 1 error is much higher than 0.05. It is difficult to calculate the exact probability, but you can derive a pessimistic approximation by assuming that the comparisons are independent, giving an upper bound to the probability of making at least one type 1 error (the experimentwise error rate) of

$$1 - (1 - 0.05)^{45} = 0.90$$

The actual probability is somewhat less than 0.90, but as the number of means increases, the chance of making at least one type 1 error approaches 1.

If you decide to control the individual type 1 error rates for each comparison, you are controlling the individual or comparisonwise error rate. On the other hand, if you want to control the overall type 1 error rate for all the comparisons, you are controlling the experimentwise error rate. It is up to you to decide whether to control the comparisonwise error rate or the experimentwise error rate, but there are many situations in which the experimentwise error rate should be held to a small value. Statistical methods for comparing three or more

means while controlling the probability of making at least one type 1 error are called *multiple-comparison procedures*.

It has been suggested that the experimentwise error rate can be held to the α level by performing the overall ANOVA F test at the α level and making further comparisons only if the F test is significant, as in Fisher's protected LSD. This assertion is false if there are more than three means (Einot and Gabriel 1975). Consider again the situation with 10 means. Suppose that one population mean differs from the others by such a sufficiently large amount that the power (probability of correctly rejecting the null hypothesis) of the F test is near 1 but that all the other population means are equal to each other. There will be $9(9 - 1)/2 = 36$ t tests of true null hypotheses, with an upper limit of 0.84 on the probability of at least one type 1 error. Thus, you must distinguish between the experimentwise error rate under the complete null hypothesis, in which all population means are equal, and the experimentwise error rate under a partial null hypothesis, in which some means are equal but others differ. The following abbreviations are used in the discussion:

CER comparisonwise error rate

EERC experimentwise error rate under the complete null hypothesis

MEER maximum experimentwise error rate under any complete or partial null hypothesis

These error rates are associated with the different [strengths of inference](#) discussed on page 3694: individual tests control the CER; tests for inhomogeneity of means control the EERC; tests that yield confidence inequalities or confidence intervals control the MEER. A preliminary F test controls the EERC but not the MEER.

You can control the MEER at the α level by setting the CER to a sufficiently small value. The Bonferroni inequality (Miller 1981) has been widely used for this purpose. If

$$\text{CER} = \frac{\alpha}{c}$$

where c is the total number of comparisons, then the MEER is less than α . Bonferroni t tests (the [BON](#) option in the [MEANS](#) statement, [ADJUST=BON](#) in the [LSMEANS](#) statement) with $\text{MEER} < \alpha$ declare two means to be significantly different if

$$|t_{ij}| \geq t(\epsilon; \nu)$$

where

$$\epsilon = \frac{2\alpha}{k(k-1)}$$

for comparison of k means.

Šidák (1967) has provided a tighter bound, showing that

$$\text{CER} = 1 - (1 - \alpha)^{1/c}$$

also ensures that $\text{MEER} \leq \alpha$ for any set of c comparisons. A Sidak t test (Games 1977), provided by the [SIDAK](#) option, is thus given by

$$|t_{ij}| \geq t(\epsilon; \nu)$$

where

$$\epsilon = 1 - (1 - \alpha)^{\frac{2}{k(k-1)}}$$

for comparison of k means.

You can use the Bonferroni additive inequality and the Sidak multiplicative inequality to control the MEER for any set of contrasts or other hypothesis tests, not just pairwise comparisons. The Bonferroni inequality can provide simultaneous inferences in any statistical application requiring tests of more than one hypothesis. Other methods discussed in this section for pairwise comparisons can also be adapted for general contrasts (Miller 1981).

Scheffé (1953, 1959) proposes another method to control the MEER for any set of contrasts or other linear hypotheses in the analysis of linear models, including pairwise comparisons, obtained with the SCHEFFE option. Two means are declared significantly different if

$$|t_{ij}| \geq \sqrt{DF \cdot F(\alpha; DF, \nu)}$$

where $F(\alpha; DF, \nu)$ is the α -level critical value of an F distribution with DF numerator degrees of freedom and ν denominator degrees of freedom. The value of DF is $k - 1$ for the **MEANS** statement, but in other statements the precise definition depends on context. For the **LSMEANS** statement, DF is the rank of the contrast matrix **L** for LS-means differences. In more general contexts—for example, the **ESTIMATE** or **LSMESTIMATE** statements in PROC GLIMMIX— DF is the rank of the contrast covariance matrix $LCov(b)L'$.

Scheffé's test is compatible with the overall ANOVA F test in that Scheffé's method never declares a contrast significant if the overall F test is nonsignificant. Most other multiple-comparison methods can find significant contrasts when the overall F test is nonsignificant and, therefore, suffer a loss of power when used with a preliminary F test.

Scheffé's method might be more powerful than the Bonferroni or Sidak method if the number of comparisons is large relative to the number of means. For pairwise comparisons, Sidak t tests are generally more powerful.

Tukey (1952, 1953) proposes a test designed specifically for pairwise comparisons based on the studentized range, sometimes called the “honestly significant difference test,” that controls the MEER when the sample sizes are equal. Tukey (1953) and Kramer (1956) independently propose a modification for unequal cell sizes. The Tukey or Tukey-Kramer method is provided by the **TUKEY** option in the **MEANS** statement and the **ADJUST=TUKEY** option in the **LSMEANS** statement. This method has fared extremely well in Monte Carlo studies (Dunnett 1980). In addition, Hayter (1984) gives a proof that the Tukey-Kramer procedure controls the MEER for means comparisons, and Hayter (1989) describes the extent to which the Tukey-Kramer procedure has been proven to control the MEER for LS-means comparisons. The Tukey-Kramer method is more powerful than the Bonferroni, Sidak, or Scheffé method for pairwise comparisons. Two means are considered significantly different by the Tukey-Kramer criterion if

$$|t_{ij}| \geq q(\alpha; k, \nu) / \sqrt{2}$$

where $q(\alpha; k, \nu)$ is the α -level critical value of a studentized range distribution of k independent normal random variables with ν degrees of freedom.

Hochberg (1974) devised a method (the GT2 or SMM option) similar to Tukey's, but it uses the studentized maximum modulus instead of the studentized range and employs the uncorrelated t inequality of Šidák (1967). It is proven to hold the MEER at a level not exceeding α with unequal sample sizes. It is generally less powerful than the Tukey-Kramer method and always less powerful than Tukey's test for equal cell sizes. Two means are declared significantly different if

$$|t_{ij}| \geq m(\alpha; c, \nu)$$

where $m(\alpha; c, \nu)$ is the α -level critical value of the studentized maximum modulus distribution of c independent normal random variables with ν degrees of freedom and $c = k(k - 1)/2$.

Gabriel (1978) proposes another method (the **GABRIEL** option) based on the studentized maximum modulus. This method is applicable only to arithmetic means. It rejects if

$$\frac{|\bar{y}_i - \bar{y}_j|}{s \left(\frac{1}{\sqrt{2n_i}} + \frac{1}{\sqrt{2n_j}} \right)} \geq m(\alpha; k, \nu)$$

For equal cell sizes, Gabriel's test is equivalent to Hochberg's GT2 method. For unequal cell sizes, Gabriel's method is more powerful than GT2 but might become liberal with highly disparate cell sizes (see also Dunnett (1980)). Gabriel's test is the only method for unequal sample sizes that lends itself to a graphical representation as intervals around the means. Assuming $\bar{y}_i > \bar{y}_j$, you can rewrite the preceding inequality as

$$\bar{y}_i - m(\alpha; k, \nu) \frac{s}{\sqrt{2n_i}} \geq \bar{y}_j + m(\alpha; k, \nu) \frac{s}{\sqrt{2n_j}}$$

The expression on the left does not depend on j , nor does the expression on the right depend on i . Hence, you can form what Gabriel calls an (l, u) -interval around each sample mean and declare two means to be significantly different if their (l, u) -intervals do not overlap. See Hsu (1996, section 5.2.1.1) for a discussion of other methods of graphically representing all pairwise comparisons.

Comparing All Treatments to a Control

One special case of means comparison is that in which the only comparisons that need to be tested are between a set of new treatments and a single control. In this case, you can achieve better power by using a method that is restricted to test only comparisons to the single control mean. Dunnett (1955) proposes a test for this situation that declares a mean significantly different from the control if

$$|t_{i0}| \geq d(\alpha; k, \nu, \rho_1, \dots, \rho_{k-1})$$

where $\bar{y}_0$ is the control mean and $d(\alpha; k, \nu, \rho_1, \dots, \rho_{k-1})$ is the critical value of the "many-to-one t statistic" (Miller 1981; Krishnaiah and Armitage 1966) for k means to be compared to a control, with ν error degrees of freedom and correlations $\rho_1, \dots, \rho_{k-1}$, $\rho_i = n_i/(n_0 + n_i)$. The correlation terms arise because each of the treatment means is being compared to the same control. Dunnett's test holds the MEER to a level not exceeding the stated α .

Analysis of Means: Comparing Each Treatments to the Average

Analysis of means (ANOM) refers to a technique for comparing group means and displaying the comparisons graphically so that you can easily see which ones are different. Means are judged as different if they are significantly different from the overall average, with significance adjusted for multiplicity. The overall average is computed as a weighted mean of the LS-means, the weights being inversely proportional to the variances. If you use the **PDIF=ANOM** option in the **LSMEANS** statement, the procedure will display the p -values (adjusted for multiplicity, by default) for tests of the differences between each LS-mean and the average LS-mean. The ANOM procedure in SAS/QC software displays both tables and graphics for the analysis of means with a variety of response types. For one-way designs, confidence intervals for **PDIF=ANOM** comparisons are equivalent to the results of PROC ANOM. The difference is that PROC GLM directly displays the confidence intervals for the differences, while the graphical output of PROC ANOM displays them as decision limits around the overall mean.

If the LS-means being compared are uncorrelated, exact adjusted p -values and critical values for confidence limits can be computed; see Nelson (1982, 1991, 1993) and Guirguis and Tobias (2004). For correlated LS-means, an approach similar to that of Hsu (1992) is employed, using a factor-analytic approximation of the correlation between the LS-means to derive approximate “effective sample sizes” for which exact critical values are computed. Note that computing the exact adjusted p -values and critical values for unbalanced designs can be computationally intensive. A simulation-based approach, as specified by the **ADJUST=SIM** option, while nondeterministic, might provide inferences that are accurate enough in much less time. See the section “[Approximate and Simulation-Based Methods](#)” on page 3700 for more details.

Approximate and Simulation-Based Methods

Tukey’s, Dunnett’s, and Nelson’s tests are all based on the same general quantile calculation:

$$q^t(\alpha, \nu, R) = \{q \ni P(\max(|t_1|, \dots, |t_n|) > q) = \alpha\}$$

where the t_i have a joint multivariate t distribution with ν degrees of freedom and correlation matrix R . In general, evaluating $q^t(\alpha, \nu, R)$ requires repeated numerical calculation of an $(n + 1)$ -fold integral. This is usually intractable, but the problem reduces to a feasible 2-fold integral when R has a certain symmetry in the case of Tukey’s test, and a *factor analytic structure* (Hsu 1992) in the case of Dunnett’s and Nelson’s tests. The R matrix has the required symmetry for exact computation of Tukey’s test in the following two cases:

- The t_i s are studentized differences between $k(k - 1)/2$ pairs of k uncorrelated means with equal variances—that is, equal sample sizes.
- The t_i s are studentized differences between $k(k - 1)/2$ pairs of k LS-means from a *variance-balanced* design (for example, a balanced incomplete block design).

See Hsu (1992, 1996) for more information. The R matrix has the factor analytic structure for exact computation of Dunnett’s and Nelson’s tests in the following two cases:

- if the t_i s are studentized differences between $k - 1$ means and a control mean, all uncorrelated. (Dunnett’s one-sided methods depend on a similar probability calculation, without the absolute values.) Note that it is not required that the variances of the means (that is, the sample sizes) be equal.
- if the t_i s are studentized differences between $k - 1$ LS-means and a control LS-mean from either a *variance-balanced* design, or a design in which the other factors are *orthogonal* to the treatment factor (for example, a randomized block design with proportional cell frequencies)

However, other important situations that do **not** result in a correlation matrix R that has the structure for exact computation are the following:

- all pairwise differences with unequal sample sizes
- differences between LS-means in many unbalanced designs

In these situations, exact calculation of $q^t(\alpha, \nu, R)$ is intractable in general. Most of the preceding methods can be viewed as using various approximations for $q^t(\alpha, \nu, R)$. When the sample sizes are unequal, the Tukey-Kramer test is equivalent to another approximation. For comparisons with a control when the correlation R does not have a factor analytic structure, Hsu (1992) suggests approximating R with a matrix R^* that does have such a structure and correspondingly approximating $q^t(\alpha, \nu, R)$ with $q^t(\alpha, \nu, R^*)$. When you request Dunnett's or Nelson's test for LS-means (the `PDIFF=CONTROL` and `ADJUST=DUNNETT` options or the `PDIFF=ANOM` and `ADJUST=NELSON` options, respectively), the GLM procedure automatically uses Hsu's approximation when appropriate.

Finally, Edwards and Berry (1987) suggest calculating $q^t(\alpha, \nu, R)$ by simulation. Multivariate t vectors are sampled from a distribution with the appropriate ν and R parameters, and Edwards and Berry (1987) suggest estimating $q^t(\alpha, \nu, R)$ by $\hat{q}$, the α th percentile of the observed values of $\max(|t_1|, \dots, |t_n|)$. Sufficient samples are generated for the true $P(\max(|t_1|, \dots, |t_n|) > \hat{q})$ to be within a certain accuracy radius γ of α with accuracy confidence $100(1 - \epsilon)$. You can approximate $q^t(\alpha, \nu, R)$ by simulation for comparisons between LS-means by specifying `ADJUST=SIM` (with any `PDIFF=` type). By default, $\gamma = 0.005$ and $\epsilon = 0.01$, so that the tail area of $\hat{q}$ is within 0.005 of α with 99% confidence. You can use the `ACC=` and `EPS=` options with `ADJUST=SIM` to reset γ and ϵ , or you can use the `NSAMP=` option to set the sample size directly. You can also control the random number sequence with the `SEED=` option.

Hsu and Nelson (1998) suggest a more accurate simulation method for estimating $q^t(\alpha, \nu, R)$, using a control variate adjustment technique. The same independent, standardized normal variates that are used to generate multivariate t vectors from a distribution with the appropriate ν and R parameters are also used to generate multivariate t vectors from a distribution for which the exact value of $q^t(\alpha, \nu, R)$ is known. $\max(|t_1|, \dots, |t_n|)$ for the second sample is used as a control variate for adjusting the quantile estimate based on the first sample; see Hsu and Nelson (1998) for more details. The control variate adjustment has the drawback that it takes somewhat longer than the crude technique of Edwards and Berry (1987), but it typically yields an estimate that is many times more accurate. In most cases, if you are using `ADJUST=SIM`, then you should specify `ADJUST=SIM(CVADJUST)`. You can also specify `ADJUST=SIM(CVADJUST REPORT)` to display a summary of the simulation that includes, among other things, the actual accuracy radius γ , which should be substantially smaller than the target accuracy radius (0.005 by default).

Multiple-Stage Tests

You can use all of the methods discussed so far to obtain simultaneous confidence intervals (Miller 1981). By sacrificing the facility for simultaneous estimation, you can obtain simultaneous tests with greater power by using multiple-stage tests (MSTs). MSTs come in both step-up and step-down varieties (Welsch 1977). The step-down methods, which have been more widely used, are available in SAS/STAT software.

Step-down MSTs first test the homogeneity of all the means at a level γ_k . If the test results in a rejection, then each subset of $k - 1$ means is tested at level γ_{k-1} ; otherwise, the procedure stops. In general, if the hypothesis of homogeneity of a set of p means is rejected at the γ_p level, then each subset of $p - 1$ means is tested at the γ_{p-1} level; otherwise, the set of p means is considered not to differ significantly and none of its subsets are tested. The many varieties of MSTs that have been proposed differ in the levels γ_p and the

statistics on which the subset tests are based. Clearly, the EERC of a step-down MST is not greater than γ_k , and the CER is not greater than γ_2 , but the MEER is a complicated function of γ_p , $p = 2, \dots, k$.

With unequal cell sizes, PROC GLM uses the harmonic mean of the cell sizes as the common sample size. However, since the resulting operating characteristics can be undesirable, MSTs are recommended only for the balanced case. When the sample sizes are equal, using the range statistic enables you to arrange the means in ascending or descending order and test only contiguous subsets. But if you specify the F statistic, this shortcut cannot be taken. For this reason, only range-based MSTs are implemented. It is common practice to report the results of an MST by writing the means in such an order and drawing lines parallel to the list of means spanning the homogeneous subsets. This form of presentation is also convenient for pairwise comparisons with equal cell sizes.

The best-known MSTs are the Duncan (the DUNCAN option) and Student-Newman-Keuls (the SNK option) methods (Miller 1981). Both use the studentized range statistic and, hence, are called *multiple range tests*. Duncan's method is often called the "new" multiple range test despite the fact that it is one of the oldest MSTs in current use.

The Duncan and SNK methods differ in the γ_p values used. For Duncan's method, they are

$$\gamma_p = 1 - (1 - \alpha)^{p-1}$$

whereas the SNK method uses

$$\gamma_p = \alpha$$

Duncan's method controls the CER at the α level. Its operating characteristics appear similar to those of Fisher's unprotected LSD or repeated t tests at level α (Petrinovich and Hardyck 1969). Since repeated t tests are easier to compute, easier to explain, and applicable to unequal sample sizes, Duncan's method is not recommended. Several published studies (for example, Carmer and Swanson (1973)) have claimed that Duncan's method is superior to Tukey's because of greater power without considering that the greater power of Duncan's method is due to its higher type 1 error rate (Einot and Gabriel 1975).

The SNK method holds the EERC to the α level but does not control the MEER (Einot and Gabriel 1975). Consider ten population means that occur in five pairs such that means within a pair are equal, but there are large differences between pairs. If you make the usual sampling assumptions and also assume that the sample sizes are very large, all subset homogeneity hypotheses for three or more means are rejected. The SNK method then comes down to five independent tests, one for each pair, each at the α level. Letting α be 0.05, the probability of at least one false rejection is

$$1 - (1 - 0.05)^5 = 0.23$$

As the number of means increases, the MEER approaches 1. Therefore, the SNK method cannot be recommended.

A variety of MSTs that control the MEER have been proposed, but these methods are not as well known as those of Duncan and SNK. One approach (Ryan 1959, 1960; Einot and Gabriel 1975; Welsch 1977) sets

$$\gamma_p = \begin{cases} 1 - (1 - \alpha)^{p/k} & \text{for } p < k - 1 \\ \alpha & \text{for } p \geq k - 1 \end{cases}$$

You can use range statistics, leading to what is called the REGWQ method, after the authors' initials. If you assume that the sample means have been arranged in descending order from $\bar{y}_1$ through $\bar{y}_k$, the homogeneity of means $\bar{y}_i, \dots, \bar{y}_j, i < j$, is rejected by REGWQ if

$$\bar{y}_i - \bar{y}_j \geq q(\gamma_p; p, \nu) \frac{s}{\sqrt{n}}$$

where $p = j - i + 1$ and the summations are over $u = i, \dots, j$ (Einot and Gabriel 1975). To ensure that the MEER is controlled, the current implementation checks whether $q(\gamma_p; p, \nu)$ is monotonically increasing in p . If not, then a set of critical values that are increasing in p is substituted instead.

REGWQ appears to be the most powerful step-down MST in the current literature (for example, Ramsey 1978). Use of a preliminary F test decreases the power of all the other multiple-comparison methods discussed previously except for Scheffé's test.

Bayesian Approach

Waller and Duncan (1969) and Duncan (1975) take an approach to multiple comparisons that differs from all the methods previously discussed in minimizing the Bayes risk under additive loss rather than controlling type 1 error rates. For each pair of population means μ_i and μ_j , null (H_0^{ij}) and alternative (H_a^{ij}) hypotheses are defined:

$$H_0^{ij}: \mu_i - \mu_j \leq 0$$

$$H_a^{ij}: \mu_i - \mu_j > 0$$

For any i, j pair, let d_0 indicate a decision in favor of H_0^{ij} and d_a indicate a decision in favor of H_a^{ij} , and let $\delta = \mu_i - \mu_j$. The loss function for the decision on the i, j pair is

$$L(d_0 | \delta) = \begin{cases} 0 & \text{if } \delta \leq 0 \\ \delta & \text{if } \delta > 0 \end{cases}$$

$$L(d_a | \delta) = \begin{cases} -k\delta & \text{if } \delta \leq 0 \\ 0 & \text{if } \delta > 0 \end{cases}$$

where k represents a constant that you specify rather than the number of means. The loss for the joint decision involving all pairs of means is the sum of the losses for each individual decision. The population means are assumed to have a normal prior distribution with unknown variance, the logarithm of the variance of the means having a uniform prior distribution. For the i, j pair, the null hypothesis is rejected if

$$\bar{y}_i - \bar{y}_j \geq t_B s \sqrt{\frac{2}{n}}$$

where t_B is the Bayesian t value (Waller and Kemp 1976) depending on k , the F statistic for the one-way ANOVA, and the degrees of freedom for F . The value of t_B is a decreasing function of F , so the Waller-Duncan test (specified by the **WALLER** option) becomes more liberal as F increases.

Recommendations

In summary, if you are interested in several individual comparisons and are not concerned about the effects of multiple inferences, you can use repeated t tests or Fisher's unprotected LSD. If you are interested in all pairwise comparisons or all comparisons with a control, you should use Tukey's or Dunnett's test, respectively, in order to make the strongest possible inferences. If you have weaker inferential requirements and, in particular, if you do not want confidence intervals for the mean differences, you should use the REGWQ method. Finally, if you agree with the Bayesian approach and Waller and Duncan's assumptions, you should use the Waller-Duncan test.

Interpretation of Multiple Comparisons

When you interpret multiple comparisons, remember that failure to reject the hypothesis that two or more means are equal should not lead you to conclude that the population means are, in fact, equal. Failure to reject the null hypothesis implies only that the difference between population means, if any, is not large enough to be detected with the given sample size. A related point is that nonsignificance is nontransitive: that is, given three sample means, the largest and smallest might be significantly different from each other, while neither is significantly different from the middle one. Nontransitive results of this type occur frequently in multiple comparisons.

Multiple comparisons can also lead to counterintuitive results when the cell sizes are unequal. Consider four cells labeled A, B, C, and D, with sample means in the order $A > B > C > D$. If A and D each have two observations, and B and C each have 10,000 observations, then the difference between B and C might be significant, while the difference between A and D is not.

Simple Effects

Suppose you use the following statements to fit a full factorial model to a two-way design:

```
data twoway;
  input A B Y @@;
  datalines;
1 1 10.6   1 1 11.0   1 1 10.6   1 1 11.3
1 2 -0.2   1 2  1.3   1 2 -0.2   1 2  0.2
1 3  0.1   1 3  0.4   1 3 -0.4   1 3  1.0
2 1 19.7   2 1 19.3   2 1 18.5   2 1 20.4
2 2 -0.2   2 2  0.5   2 2  0.8   2 2 -0.4
2 3 -0.9   2 3 -0.1   2 3 -0.2   2 3 -1.7
3 1 29.7   3 1 29.6   3 1 29.0   3 1 30.2
3 2  1.5   3 2  0.2   3 2 -1.5   3 2  1.3
3 3  0.2   3 3  0.4   3 3 -0.4   3 3 -2.2
;

proc glm data=twoway;
  class A B;
  model Y = A B A*B;
run;
```

Partial results for the analysis of variance are shown in [Figure 47.21](#). The Type I and Type III results are the same because this is a balanced design.

Figure 47.21 Two-Way Design with Significant Interaction**The GLM Procedure****Dependent Variable: Y**

Source	DF	Type I SS	Mean Square	F Value	Pr > F
A	2	219.905000	109.952500	165.11	<.0001
B	2	3206.101667	1603.050833	2407.25	<.0001
A*B	4	487.103333	121.775833	182.87	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
A	2	219.905000	109.952500	165.11	<.0001
B	2	3206.101667	1603.050833	2407.25	<.0001
A*B	4	487.103333	121.775833	182.87	<.0001

The interaction A*B is significant, indicating that the effect of A depends on the level of B. In some cases, you might be interested in looking at the differences between predicted values across A for different levels of B. Winer (1971) calls this the *simple effects* of A. You can compute simple effects with the **LSMEANS** statement by specifying the **SLICE=** option. In this case, since the GLM procedure is interactive, you can compute the simple effects of A by submitting the following statements after the preceding statements.

```
lsmeans A*B / slice=B;
run;
```

The results are shown [Figure 47.22](#). Note that A has a significant effect for B=1 but not for B=2 and B=3.

Figure 47.22 Interaction LS-Means and Simple Effects**The GLM Procedure**
Least Squares Means

A	B	Y LSMEAN
1	1	10.8750000
1	2	0.2750000
1	3	0.2750000
2	1	19.4750000
2	2	0.1750000
2	3	-0.7250000
3	1	29.6250000
3	2	0.3750000
3	3	-0.5000000

The GLM Procedure
Least Squares Means

A*B Effect Sliced by B for Y					
B	DF	Sum of Squares	Mean Square	F Value	Pr > F
1	2	704.726667	352.363333	529.13	<.0001
2	2	0.080000	0.040000	0.06	0.9418
3	2	2.201667	1.100833	1.65	0.2103

Homogeneity of Variance in One-Way Models

One of the usual [assumptions](#) in using the GLM procedure is that the underlying errors are all uncorrelated with homogeneous variances (see page 3670). You can test this assumption in PROC GLM by using the [HOVTEST](#) option in the [MEANS](#) statement, requesting a *homogeneity of variance* test. This section discusses the computational details behind these tests. Note that the GLM procedure allows homogeneity of variance testing for simple one-way models only. Homogeneity of variance testing for more complex models is a subject of current research.

Bartlett (1937) proposes a test for equal variances that is a modification of the normal-theory likelihood ratio test (the [HOVTEST=BARTLETT](#) option). While Bartlett's test has accurate Type I error rates and optimal power when the underlying distribution of the data is normal, it can be very inaccurate if that distribution is even slightly nonnormal (Box 1953). Therefore, Bartlett's test is not recommended for routine use.

An approach that leads to tests that are much more robust to the underlying distribution is to transform the original values of the dependent variable to derive a *dispersion variable* and then to perform analysis of variance on this variable. The significance level for the test of homogeneity of variance is the *p*-value for the ANOVA *F* test on the dispersion variable. All of the homogeneity of variance tests available in PROC GLM except Bartlett's use this approach.

Levene's test (Levene 1960) is widely considered to be the standard homogeneity of variance test (the [HOVTEST=LEVENE](#) option). Levene's test is of the dispersion-variable-ANOVA form discussed previously, where the dispersion variable is either of the following:

$$\begin{aligned} z_{ij}^2 &= (y_{ij} - \bar{y}_i)^2 \quad (\text{TYPE=SQUARE, the default}) \\ z_{ij} &= |y_{ij} - \bar{y}_i| \quad (\text{TYPE=ABS}) \end{aligned}$$

O'Brien (1979) proposes a test ([HOVTEST=OBRIEN](#)) that is basically a modification of Levene's z_{ij}^2 , using the dispersion variable

$$z_{ij}^W = \frac{(W + n_i - 2)n_i(y_{ij} - \bar{y}_i)^2 - W(n_i - 1)\sigma_i^2}{(n_i - 1)(n_i - 2)}$$

where n_i is the size of the *i*th group and σ_i^2 is its sample variance. You can use the *W=* option in parentheses to tune O'Brien's z_{ij}^W dispersion variable to match the suspected kurtosis of the underlying distribution. The choice of the value of the *W=* option is rarely critical. By default, *W=0.5*, as suggested by O'Brien (1979, 1981).

Finally, Brown and Forsythe (1974) suggest using the absolute deviations from the group *medians*:

$$z_{ij}^{\text{BF}} = |y_{ij} - m_i|$$

where m_i is the median of the *i*th group. You can use the [HOVTEST=BF](#) option to specify this test.

Simulation results (Conover, Johnson, and Johnson 1981; Olejnik and Algina 1987) show that, while all of these ANOVA-based tests are reasonably robust to the underlying distribution, the Brown-Forsythe test seems best at providing power to detect variance differences while protecting the Type I error probability. However, since the within-group medians are required for the Brown-Forsythe test, it can be resource intensive if there are very many groups or if some groups are very large.

If one of these tests rejects the assumption of homogeneity of variance, you should use Welch's ANOVA instead of the usual ANOVA to test for differences between group means. However, this conclusion holds only if you use one of the robust homogeneity of variance tests (that is, not for `HOVTEST=BARTLETT`); even then, any homogeneity of variance test has too little power to be relied upon to always detect when Welch's ANOVA is appropriate. Unless the group variances are extremely different or the number of groups is large, the usual ANOVA test is relatively robust when the groups are all about the same size. As Box (1953) notes, "To make the preliminary test on variances is rather like putting to sea in a rowing boat to find out whether conditions are sufficiently calm for an ocean liner to leave port!"

Example 47.10 illustrates the use of the `HOVTEST` and `WELCH` options in the `MEANS` statement in testing for equal group variances and adjusting for unequal group variances in a one-way ANOVA.

Weighted Means

If you specify a `WEIGHT` statement and one or more of the multiple comparisons options, the variance of the difference between weighted group means for group i and j is computed as

$$\text{MSE} \times \left(\frac{1}{w_i} + \frac{1}{w_j} \right)$$

where w_i is the sum of the weights for the observations in group i .

Construction of Least Squares Means

To construct a least squares mean (LS-mean) for a particular level of a particular effect, construct a row vector $\mathbf{L}$ according to the following rules and use it in an `ESTIMATE` statement to compute the value of the LS-mean:

1. Set all L_i that correspond to covariates (continuous variables) to their mean value.
2. Consider effects that are contained by the particular effect. (For a definition of *containing*, see Chapter 15, "The Four Types of Estimable Functions.") Set the L_i that correspond to levels associated with the particular level equal to 1. Set all other L_i in these effects equal to 0.
3. Consider the particular effect. Set the L_i that correspond to the particular level equal to 1. Set the L_i that correspond to other levels equal to 0.
4. Consider the effects that contain the particular effect. If these effects are not nested within the particular effect, then set the L_i that correspond to the particular level to $1/k$, where k is the number of such columns. If these effects are nested within the particular effect, then set the L_i that correspond to the particular level to $1/(k_1 k_2)$, where k_1 is the number of nested levels within this combination of nested effects and k_2 is the number of such combinations. For L_i that correspond to other levels, use 0.
5. Consider other effects that are not yet considered. For each effect that has no nested factors, set all L_i that correspond to this effect to $1/j$, where j is the number of levels in the effect. For each effect that has nested factors, set all L_i that correspond to this effect to $1/(j_1 j_2)$, where j_1 is the number of nested levels within a particular combination of nested effects and j_2 is the number of such combinations.

The consequence of these rules is that the sum of the X s within any classification effect is 1. This set of X s forms a linear combination of the parameters that is checked for estimability before it is evaluated.

For example, consider the following model:

```
proc glm;
  class A B C;
  model Y=A B A*B C Z;
  lsmeans A B A*B C;
run;
```

Assume A has 3 levels, B has 2 levels, and C has 2 levels, and assume that every combination of levels of A and B exists in the data. Assume also that Z is a continuous variable with an average of 12.5. Then the least squares means are computed by the following linear combinations of the parameter estimates:

		A			B		A*B						C		Z
	μ	1	2	3	1	2	11	12	21	22	31	32	1	2	
LSM()	1	1/3	1/3	1/3	1/2	1/2	1/6	1/6	1/6	1/6	1/6	1/6	1/2	1/2	12.5
LSM(A1)	1	1	0	0	1/2	1/2	1/2	1/2	0	0	0	0	1/2	1/2	12.5
LSM(A2)	1	0	1	0	1/2	1/2	0	0	1/2	1/2	0	0	1/2	1/2	12.5
LSM(A3)	1	0	0	1	1/2	1/2	0	0	0	0	1/2	1/2	1/2	1/2	12.5
LSM(B1)	1	1/3	1/3	1/3	1	0	1/3	0	1/3	0	1/3	0	1/2	1/2	12.5
LSM(B2)	1	1/3	1/3	1/3	0	1	0	1/3	0	1/3	0	1/3	1/2	1/2	12.5
LSM(AB11)	1	1	0	0	1	0	1	0	0	0	0	0	1/2	1/2	12.5
LSM(AB12)	1	1	0	0	0	1	0	1	0	0	0	0	1/2	1/2	12.5
LSM(AB21)	1	0	1	0	1	0	0	0	1	0	0	0	1/2	1/2	12.5
LSM(AB22)	1	0	1	0	0	1	0	0	0	1	0	0	1/2	1/2	12.5
LSM(AB31)	1	0	0	1	1	0	0	0	0	0	1	0	1/2	1/2	12.5
LSM(AB32)	1	0	0	1	0	1	0	0	0	0	0	1	1/2	1/2	12.5
LSM(C1)	1	1/3	1/3	1/3	1/2	1/2	1/6	1/6	1/6	1/6	1/6	1/6	1	0	12.5
LSM(C2)	1	1/3	1/3	1/3	1/2	1/2	1/6	1/6	1/6	1/6	1/6	1/6	0	1	12.5

Setting Covariate Values

By default, all covariate effects are set equal to their mean values for computation of standard LS-means. The **AT** option in the **LSMEANS** statement enables you to set the covariates to whatever values you consider interesting.

If any effect contains two or more covariates, the **AT** option sets the effect equal to the product of the individual means rather than the mean of the product (as with standard LS-means calculations). The **AT MEANS** option leaves covariates equal to their mean values (as with standard LS-means) and incorporates this adjustment to crossproducts of covariates.

For example, the following statements specify a model with a classification variable A and two continuous variables, x1 and x2:

```
class A;
model y = A x1 x2 x1*x2;
```

The coefficients for the continuous effects with various **AT** specifications are shown in the following table.

Syntax	x1	x2	x1*x2
lsmeans A;	$\bar{x}_1$	$\bar{x}_2$	$\bar{x}_1\bar{x}_2$
lsmeans A / at means ;	$\bar{x}_1$	$\bar{x}_2$	$\bar{x}_1 \cdot \bar{x}_2$
lsmeans A / at x1=1.2;	1.2	$\bar{x}_2$	$1.2 \cdot \bar{x}_2$
lsmeans A / at (x1 x2)=(1.2 0.3);	1.2	0.3	$1.2 \cdot 0.3$

For the first two **LSMEANS** statements, the A LS-mean coefficient for x1 is $\bar{x}_1$ (the mean of x1) and for x2 is $\bar{x}_2$ (the mean of x2). However, the coefficient for x1*x2 is $\bar{x}_1\bar{x}_2$ for the first **LSMEANS** statement, but it is $\bar{x}_1 \cdot \bar{x}_2$ for the second **LSMEANS** statement. The third **LSMEANS** statement sets the coefficient for x1 equal to 1.2 and leaves the coefficient for x2 at $\bar{x}_2$, and the final **LSMEANS** statement sets these values to 1.2 and 0.3, respectively.

The covariate means used in computing the LS-means are affected by the **AT** option in the **LSMEANS** statement as follows:

- If you use a **WEIGHT** statement:
 - If you do not specify the **AT** option, unweighted covariate means are used for the covariate coefficients.
 - If you specify the **AT** option, weighted covariate means are used for the covariate coefficients for which no explicit **AT** values are specified (thus for all the covariate means if you specify **AT MEANS**).
- If any observations have missing dependent variable values:
 - If you do not specify the **AT** option, these observations are not included in computing the covariate means.
 - If you specify the **AT** option, these observations are included unless they form a missing cell (a combination of CLASS variables all of whose responses are missing).

These conditions are summarized in Table 47.14. You can use the **E** option in conjunction with the **AT** option to verify that the modified LS-means coefficients are the ones you want.

Table 47.14 Treatment of Covariate Means with and without the **AT** Option

	Without AT	With AT
WEIGHT statement	Weights not used	Weights used
Observations with missing dependent variable values	Observations not included	Observations included

The **AT** option is disabled if you specify the **BYLEVEL** option, in which case the coefficients for the covariates are set equal to their means within each level of the LS-mean effect in question.

Changing the Weighting Scheme

The standard LS-means have equal coefficients across classification effects; however, the **OM** option in the **LSMEANS** statement changes these coefficients to be proportional to those found in the input data set. This adjustment is reasonable when you want your inferences to apply to a population that is not necessarily balanced but has the margins observed in the original data set.

In computing the observed margins, PROC GLM uses all observations for which there are no missing independent variables, including those for which there are missing dependent variables. Also, if there is a **WEIGHT** variable, PROC GLM uses weighted margins to construct the LS-means coefficients. If the analysis data set is balanced or if you specify a simple one-way model, the LS-means will be unchanged by the OM option.

The **BYLEVEL** option modifies the observed-margins LS-means. Instead of computing the margins across the entire data set, PROC GLM computes separate margins for each level of the LS-mean effect in question. The resulting LS-means are actually equal to raw means in this case. The **BYLEVEL** option disables the **AT** option if it is specified.

Note that the MIXED procedure implements a more versatile form of the **OM** option, enabling you to specifying an alternative data set over which to compute observed margins. If you use the **BYLEVEL** option, too, then this data set is effectively the “population” over which the population marginal means are computed. For more information, see Chapter 78, “The MIXED Procedure.”

You might want to use the **E** option in conjunction with either the **OM** or **BYLEVEL** option to verify that the modified LS-means coefficients are the ones you want. It is possible that the modified LS-means are not estimable when the standard ones are, or vice versa.

Estimability of LS-Means

LS-means are defined as certain linear combinations of the parameters. As such, it is possible for them to be inestimable. In fact, it is possible for a pair of LS-means to be both inestimable but their difference estimable. When this happens, only the entries that correspond to the estimable difference are computed and displayed in the Diff's table. If **ADJUST=SIMULATE** is specified when there are inestimable LS-means differences, adjusted results for all differences are displayed as missing.

Multivariate Analysis of Variance

If you fit several dependent variables to the same effects, you might want to make joint tests involving parameters of several dependent variables. Suppose you have p dependent variables, k parameters for each dependent variable, and n observations. The models can be collected into one equation:

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

where $\mathbf{Y}$ is $n \times p$, $\mathbf{X}$ is $n \times k$, $\boldsymbol{\beta}$ is $k \times p$, and $\boldsymbol{\epsilon}$ is $n \times p$. Each of the p models can be estimated and tested separately. However, you might also want to consider the joint distribution and test the p models simultaneously.

For multivariate tests, you need to make some assumptions about the errors. With p dependent variables, there are $n \times p$ errors that are independent across observations but not across dependent variables. Assume

$$\text{vec}(\boldsymbol{\epsilon}) \sim N(\mathbf{0}, \mathbf{I}_n \otimes \boldsymbol{\Sigma})$$

where $\text{vec}(\boldsymbol{\epsilon})$ strings $\boldsymbol{\epsilon}$ out by rows, $\otimes$ denotes Kronecker product multiplication, and $\boldsymbol{\Sigma}$ is $p \times p$. $\boldsymbol{\Sigma}$ can be estimated by

$$\mathbf{S} = \frac{\mathbf{e}'\mathbf{e}}{n - r} = \frac{(\mathbf{Y} - \mathbf{X}\mathbf{b})'(\mathbf{Y} - \mathbf{X}\mathbf{b})}{n - r}$$

where $\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}$, r is the rank of the $\mathbf{X}$ matrix, and $\mathbf{e}$ is the matrix of residuals.

If $\mathbf{S}$ is scaled to unit diagonals, the values in $\mathbf{S}$ are called *partial correlations of the Y s adjusting for the X s*. This matrix can be displayed by PROC GLM if **PRINTE** is specified as a **MANOVA** option.

The multivariate general linear hypothesis is written

$$\mathbf{L}\boldsymbol{\beta}\mathbf{M} = 0$$

You can form hypotheses for linear combinations across columns, as well as across rows of $\boldsymbol{\beta}$.

The **MANOVA** statement of the GLM procedure tests special cases where $\mathbf{L}$ corresponds to Type I, Type II, Type III, or Type IV tests, and $\mathbf{M}$ is the $p \times p$ identity matrix. These tests are joint tests that the given type of hypothesis holds for all dependent variables in the model, and they are often sufficient to test all hypotheses of interest.

Finally, when these special cases are not appropriate, you can specify your own $\mathbf{L}$ and $\mathbf{M}$ matrices by using the **CONTRAST** statement before the **MANOVA** statement and the **M=** specification in the **MANOVA** statement, respectively. Another alternative is to use a **REPEATED** statement, which automatically generates a variety of $\mathbf{M}$ matrices useful in repeated measures analysis of variance. See the section “**REPEATED Statement**” on page 3664 and the section “**Repeated Measures Analysis of Variance**” on page 3712 for more information.

One useful way to think of a MANOVA analysis with an $\mathbf{M}$ matrix other than the identity is as an analysis of a set of transformed variables defined by the columns of the $\mathbf{M}$ matrix. You should note, however, that PROC GLM always displays the $\mathbf{M}$ matrix in such a way that the transformed variables are defined by the rows, not the columns, of the displayed $\mathbf{M}$ matrix.

All multivariate tests carried out by the GLM procedure first construct the matrices $\mathbf{H}$ and $\mathbf{E}$ corresponding to the numerator and denominator, respectively, of a univariate F test:

$$\begin{aligned}\mathbf{H} &= \mathbf{M}'(\mathbf{Lb})'(\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}')^{-1}(\mathbf{Lb})\mathbf{M} \\ \mathbf{E} &= \mathbf{M}'(\mathbf{Y}'\mathbf{Y} - \mathbf{b}'(\mathbf{X}'\mathbf{X})\mathbf{b})\mathbf{M}\end{aligned}$$

The diagonal elements of $\mathbf{H}$ and $\mathbf{E}$ correspond to the hypothesis and error SS for univariate tests. When the $\mathbf{M}$ matrix is the identity matrix (the default), these tests are for the original dependent variables on the left side of the **MODEL** statement. When an $\mathbf{M}$ matrix other than the identity is specified, the tests are for transformed variables defined by the columns of the $\mathbf{M}$ matrix. These tests can be studied by requesting the **SUMMARY** option, which produces univariate analyses for each original or transformed variable.

Four multivariate test statistics, all functions of the eigenvalues of $\mathbf{E}^{-1}\mathbf{H}$ (or $(\mathbf{E} + \mathbf{H})^{-1}\mathbf{H}$), are constructed:

- Wilks' lambda = $\det(\mathbf{E})/\det(\mathbf{H} + \mathbf{E})$
- Pillai's trace = $\text{trace}(\mathbf{H}(\mathbf{H} + \mathbf{E})^{-1})$
- Hotelling-Lawley trace = $\text{trace}(\mathbf{E}^{-1}\mathbf{H})$
- Roy's greatest root = λ , largest eigenvalue of $\mathbf{E}^{-1}\mathbf{H}$

By default, all four are reported with p -values based on F approximations, as discussed in the “Multivariate Tests” section in Chapter 4, “**Introduction to Regression Procedures**.” Alternatively, if you specify **MSTAT=EXACT** in the associated **MANOVA** or **REPEATED** statement, p -values for three of the four tests

are computed exactly (Wilks' lambda, the Hotelling-Lawley trace, and Roy's greatest root), and the p -values for the fourth (Pillai's trace) are based on an F approximation that is more accurate than the default. See the "Multivariate Tests" section in Chapter 4, "[Introduction to Regression Procedures](#)," for more details on the exact calculations.

Repeated Measures Analysis of Variance

When several measurements are taken on the same experimental unit (person, plant, machine, and so on), the measurements tend to be correlated with each other. When the measurements represent qualitatively different things, such as weight, length, and width, this correlation is best taken into account by use of multivariate methods, such as multivariate analysis of variance. When the measurements can be thought of as responses to levels of an experimental factor of interest, such as time, treatment, or dose, the correlation can be taken into account by performing a repeated measures analysis of variance.

PROC GLM provides both univariate and multivariate tests for repeated measures for one response. For an overall reference on univariate repeated measures, see Winer (1971). The multivariate approach is covered in Cole and Grizzle (1966). For a discussion of the relative merits of the two approaches, see LaTour and Miniard (1983).

Another approach to analysis of repeated measures is via general mixed models. This approach can handle balanced as well as unbalanced or missing within-subject data, and it offers more options for modeling the within-subject covariance. The main drawback of the mixed models approach is that it generally requires iteration and, thus, might be less computationally efficient. For further details on this approach, see Chapter 78, "[The MIXED Procedure](#)," and Wolfinger and Chang (1995).

Organization of Data for Repeated Measure Analysis

In order to deal efficiently with the correlation of repeated measures, the GLM procedure uses the multivariate method of specifying the model, even if only a univariate analysis is desired. In some cases, data might already be entered in the univariate mode, with each repeated measure listed as a separate observation along with a variable that represents the experimental unit (subject) on which measurement is taken. Consider the following data set Old:

```
data Old;
  input Subject Group Time y;
  datalines;
1 1 1 15
1 1 2 19
1 1 3 25
2 1 1 21
2 1 2 18
2 1 3 17
1 2 1 14
1 2 2 12
1 2 3 16
2 2 1 11
2 2 2 20
2 2 3 21

... more lines ...
```

```

10 3 1 14
10 3 2 18
10 3 3 16
;

```

There are three observations for each subject, corresponding to measurements taken at times 1, 2, and 3. These data could be analyzed using the following statements:

```

proc glm data=Old;
  class Group Subject Time;
  model y=Group Subject (Group) Time Group*Time;
  test h=Group e=Subject (Group);
run;

```

However, this analysis assumes subjects' measurements are uncorrelated across time. A repeated measures analysis does not make this assumption. It uses the following data set New:

```

data New;
  input Group y1 y2 y3;
  datalines;
1 15 19 25
1 21 18 17
2 14 12 16
2 11 20 21
2 24 15 12

... more lines ...

3 14 18 16
;

```

In the data set New, the three measurements for a subject are all in one observation. For example, the measurements for subject 1 for times 1, 2, and 3 are 15, 19, and 25, respectively. For these data, the statements for a repeated measures analysis (assuming default options) are

```

proc glm data=New;
  class Group;
  model y1-y3 = Group / nouni;
  repeated Time;
run;

```

To convert the univariate form of repeated measures data to the multivariate form, you can use a program like the following:

```

proc sort data=Old;
  by Group Subject;
run;

data New(keep=y1-y3 Group);
  array yy(3) y1-y3;
  do Time = 1 to 3;
    set Old;
    by Group Subject;
    yy(Time) = y;
  end;
run;

```

```

        if last.Subject then return;
    end;
run;

```

Alternatively, you could use PROC TRANSPOSE to achieve the same results with a program like this one:

```

proc sort data=Old;
    by Group Subject;
run;

proc transpose out=New(rename=( _1=y1 _2=y2 _3=y3 ));
    by Group Subject;
    id Time;
run;

```

See the discussions in *SAS Language Reference: Concepts* for more information about rearrangement of data sets.

Hypothesis Testing in Repeated Measures Analysis

In repeated measures analysis of variance, the effects of interest are as follows:

- between-subject effects (such as GROUP in the previous example)
- within-subject effects (such as TIME in the previous example)
- interactions between the two types of effects (such as GROUP*TIME in the previous example)

Repeated measures analyses are distinguished from MANOVA because of interest in testing hypotheses about the within-subject effects and the within-subject-by-between-subject interactions.

For tests that involve only between-subjects effects, both the multivariate and univariate approaches give rise to the same tests. These tests are provided for all effects in the **MODEL** statement, as well as for any **CONTRAST**s specified. The ANOVA table for these tests is labeled “Tests of Hypotheses for Between Subjects Effects” in the PROC GLM results. These tests are constructed by first adding together the dependent variables in the model. Then an analysis of variance is performed on the sum divided by the square root of the number of dependent variables. For example, the statements

```

model y1-y3=group;
repeated time;

```

give a one-way analysis of variance that uses $(Y1 + Y2 + Y3)/\sqrt{3}$ as the dependent variable for performing tests of hypothesis on the between-subject effect GROUP. Tests for between-subject effects are equivalent to tests of the hypothesis $L\beta M = 0$, where **M** is simply a vector of 1s.

For within-subject effects and for within-subject-by-between-subject interaction effects, the univariate and multivariate approaches yield different tests. These tests are provided for the within-subject effects and for the interactions between these effects and the other effects in the **MODEL** statement, as well as for any **CONTRAST**s specified. The univariate tests are displayed in a table labeled “Univariate Tests of Hypotheses for Within Subject Effects.” Results for multivariate tests are displayed in a table labeled “Repeated Measures Analysis of Variance.”

The multivariate tests provided for within-subjects effects and interactions involving these effects are Wilks’ lambda, Pillai’s trace, Hotelling-Lawley trace, and Roy’s greatest root. For further details on these four

statistics, see the “Multivariate Tests” section in Chapter 4, “[Introduction to Regression Procedures](#).” As an example, the statements

```
model y1-y3=group;
repeated time;
```

produce multivariate tests for the within-subject effect TIME and the interaction TIME*GROUP.

The multivariate tests for within-subject effects are produced by testing the hypothesis $\mathbf{L}\boldsymbol{\beta}\mathbf{M} = 0$, where the $\mathbf{L}$ matrix is the usual matrix corresponding to the Type I, Type II, Type III, or Type IV hypotheses test, and the $\mathbf{M}$ matrix is one of several matrices depending on the transformation that you specify in the [REPEATED](#) statement. These multivariate tests require that the column rank of $\mathbf{M}$ be less than or equal to the number of error degrees of freedom. Besides that, the only assumption required for valid tests is that the dependent variables in the model have a multivariate normal distribution with a common covariance matrix across the between-subject effects.

The univariate tests for within-subject effects and interactions involving these effects require some assumptions for the probabilities provided by the ordinary F tests to be correct. Specifically, these tests require certain patterns of covariance matrices, known as Type H covariances (Huynh and Feldt 1970). Data with these patterns in the covariance matrices are said to satisfy the Huynh-Feldt condition. You can test this assumption (and the Huynh-Feldt condition) by applying a sphericity test (Anderson 1958) to any set of variables defined by an orthogonal contrast transformation. Such a set of variables is known as a set of orthogonal components. When you use the [PRINTE](#) option in the [REPEATED](#) statement, this sphericity test is applied both to the transformed variables defined by the [REPEATED](#) statement and to a set of orthogonal components if the specified transformation is not orthogonal. It is the test applied to the orthogonal components that is important in determining whether your data have a Type H covariance structure. When there are only two levels of the within-subject effect, there is only one transformed variable, and a sphericity test is not needed. The sphericity test is labeled “Test for Sphericity” in the output.

If your data satisfy the preceding assumptions, use the usual F tests to test univariate hypotheses for the within-subject effects and associated interactions.

If your data do not satisfy the assumption of Type H covariance, an adjustment to numerator and denominator degrees of freedom can be used. Several such adjustments, based on a degrees-of-freedom adjustment factor known as ϵ (epsilon) (Box 1954), are provided in PROC GLM. All these adjustments estimate ϵ and then multiply the numerator and denominator degrees of freedom by this estimate before determining significance levels for the F tests. Significance levels associated with the adjusted tests are labeled “Adj Pr > F” in the output. Two such adjustments are displayed. One is the maximum likelihood estimate of Box’s ϵ factor, which is known to be conservative, possibly very much so. The other adjustment is intended to be unbiased although possibly at the cost of being liberal. The first adjustment is labeled as the “Greenhouse-Geisser Epsilon.” It has the form

$$\hat{\epsilon}_{GG} = \frac{\text{trace}^2(\mathbf{E})/b}{\text{trace}(\mathbf{E}^2)}$$

where $\mathbf{E}$ is the error matrix for the corresponding multivariate test and b is the degrees of freedom for the hypothesis being tested. $\hat{\epsilon}_{GG}$ was initially proposed for use in data analysis by Greenhouse and Geisser (1959). Significance levels associated with F tests thus adjusted are labeled “G-G” in the output.

Huynh and Feldt (1976) showed that $\hat{\epsilon}_{GG}$ tends to be biased downward (that is, conservative), especially for small samples. Alternative estimates have been proposed to overcome this conservative bias, and there are several options for which estimate to display along with $\hat{\epsilon}_{GG}$.

- Huynh and Feldt (1976) proposed an estimate of Box's epsilon, constructed using estimators of its numerator and denominator that are intended to be unbiased. The Huynh-Feldt epsilon has the form of a modification of the Greenhouse-Geisser epsilon,

$$\hat{\epsilon}_{\text{HF}} = \frac{nb\hat{\epsilon}_{\text{GG}} - 2}{b(\text{DFE} - b\hat{\epsilon}_{\text{GG}})}$$

where n is the number of subjects and DFE is the degrees of freedom for error. The numerator of this estimate is precisely unbiased only when there are no between-subject effects, but $\hat{\epsilon}_{\text{HF}}$ is still often employed even with nontrivial between-subject models; it was the only unbiased epsilon alternative in SAS/STAT releases before SAS/STAT 9.22. The Huynh-Feldt epsilon is no longer the default, but you can request it and its corresponding F test by using the **UEPSDEF=HF** option in the **REPEATED** statement. The estimate is labeled “Huynh-Feldt Epsilon” in the PROC GLM output, and the significance levels associated with adjusted F tests are labeled “H-F.”

- Lecoutre (1991) gave the unbiased form of the numerator of Box's epsilon when there is one between-subject effect. The correct form of Huynh and Feldt's idea in this case is

$$\hat{\epsilon}_{\text{HFL}} = \frac{(\text{DFE} + 1)b\hat{\epsilon}_{\text{GG}} - 2}{b(\text{DFE} - b\hat{\epsilon}_{\text{GG}})}$$

More recently, Gribbin (2007) showed that $\hat{\epsilon}_{\text{HFL}}$ applies to general between-subject models, and Chi et al. (2012) showed that it extends even to situations where the number of error degrees of freedom is less than the column rank of the within-subject contrast matrix. Thus, the Lecoutre correction of the Huynh-Feldt epsilon is displayed by default along with the Greenhouse-Geisser epsilon; you can also explicitly request it by using the **UEPSDEF=HFL** option in the **REPEATED** statement. The estimate is labeled “Huynh-Feldt-Lecoutre Epsilon” in the PROC GLM output, and the significance levels associated with adjusted F tests are labeled “H-F-L.”

- Finally, Chi et al. (2012) suggest that Box's epsilon might be better estimated by replacing the reciprocal of an unbiased form of the denominator with an approximately unbiased form of the reciprocal itself. The resulting estimator can be written as a multiple of the corrected Huynh-Feldt epsilon $\hat{\epsilon}_{\text{HFL}}$,

$$\hat{\epsilon}_{\text{CM}} = \hat{\epsilon}_{\text{HFL}}(\nu_a - 2)(\nu_a - 4)/\nu_a^2$$

where $\nu_a = (\text{DFE} - 1) + \text{DFE}(\text{DFE} - 1)/2$. Simulations indicate that $\hat{\epsilon}_{\text{CM}}$ does a good job of providing accurate p -values without being either too conservative or too liberal. Over a wide range of cases, it is never much worse than any other alternative epsilon and often much better. You can request that the Chi-Muller epsilon estimate and its corresponding F test be displayed by using the **UEPSDEF=CM** option in the **REPEATED** statement. The estimate is labeled “Chi-Muller Epsilon” in the PROC GLM output, and the significance levels associated with adjusted F tests are labeled “C-M.”

Although ϵ must be in the range of 0 to 1, the three approximately unbiased estimators can be outside this range. When any of these estimators is greater than 1, a value of 1 is used in all calculations for probabilities—in other words, the probabilities are not adjusted. Additionally, if $\hat{\epsilon}_{\text{CM}} < 1/b$, then the degrees of freedom are adjusted by $1/b$ instead of $\hat{\epsilon}_{\text{CM}}$.

In summary, if your data do not meet the assumptions, use adjusted F tests. However, when you strongly suspect that your data might not have Type H covariance, all these univariate tests should be interpreted cautiously. In such cases, you should consider using the multivariate tests instead.

The univariate sums of squares for hypotheses involving within-subject effects can be easily calculated from the **H** and **E** matrices corresponding to the multivariate tests described in the section “[Multivariate Analysis of Variance](#)” on page 3710. If the **M** matrix is orthogonal, the univariate sums of squares is calculated as the trace (sum of diagonal elements) of the appropriate **H** matrix; if it is not orthogonal, PROC GLM calculates the trace of the **H** matrix that results from an orthogonal **M** matrix transformation. The appropriate error term for the univariate *F* tests is constructed in a similar way from the error SSCP matrix and is labeled `Error(factorname)`, where *factorname* indicates the **M** matrix that is used in the transformation.

When the design specifies more than one repeated measures factor, PROC GLM computes the **M** matrix for a given effect as the direct (Kronecker) product of the **M** matrices defined by the [REPEATED](#) statement if the factor is involved in the effect or as a vector of 1s if the factor is not involved. The test for the main effect of a repeated measures factor is constructed using an **L** matrix that corresponds to a test that the mean of the observation is zero. Thus, the main effect test for repeated measures is a test that the means of the variables defined by the **M** matrix are all equal to zero, while interactions involving repeated measures effects are tests that the between-subjects factors involved in the interaction have no effect on the means of the transformed variables defined by the **M** matrix. In addition, you can specify other **L** matrices to test hypotheses of interest by using the [CONTRAST](#) statement, since hypotheses defined by [CONTRAST](#) statements are also tested in the [REPEATED](#) analysis. To see which combinations of the original variables the transformed variables represent, you can specify the [PRINTM](#) option in the [REPEATED](#) statement. This option displays the transpose of **M**, which is labeled as **M** in the PROC GLM results. The tests produced are the same for any choice of transformation (**M**) matrix specified in the [REPEATED](#) statement; however, depending on the nature of the repeated measurements being studied, a particular choice of transformation matrix, coupled with the [CANONICAL](#) or [SUMMARY](#) option, can provide additional insight into the data being studied.

Transformations Used in Repeated Measures Analysis of Variance

As mentioned in the specifications of the [REPEATED](#) statement, several different **M** matrices can be generated automatically, based on the transformation that you specify in the [REPEATED](#) statement. Remember that both the univariate and multivariate tests that PROC GLM performs are unaffected by the choice of transformation; the choice of transformation is important only when you are trying to study the nature of a repeated measures effect, particularly with the [CANONICAL](#) and [SUMMARY](#) options. If one of these matrices does not meet your needs for a particular analysis, you might want to use the **M=** option in the [MANOVA](#) statement to perform the tests of interest.

The following sections describe the transformations available in the [REPEATED](#) statement, provide an example of the **M** matrix that is produced, and give guidelines for the use of the transformation. As in the PROC GLM output, the displayed matrix is labeled **M**. This is the **M'** matrix.

[CONTRAST Transformation](#)

This is the default transformation used by the [REPEATED](#) statement. It is useful when one level of the repeated measures effect can be thought of as a control level against which the others are compared. For example, if five drugs are administered to each of several animals and the first drug is a control or placebo, the statements

```
proc glm;
  model d1-d5= / nouni;
  repeated drug 5 contrast(1) / summary printm;
run;
```

produce the following M matrix:

$$M = \begin{bmatrix} -1 & 1 & 0 & 0 & 0 \\ -1 & 0 & 1 & 0 & 0 \\ -1 & 0 & 0 & 1 & 0 \\ -1 & 0 & 0 & 0 & 1 \end{bmatrix}$$

When you examine the analysis of variance tables produced by the **SUMMARY** option, you can tell which of the drugs differed significantly from the placebo.

POLYNOMIAL Transformation

This transformation is useful when the levels of the repeated measure represent quantitative values of a treatment, such as dose or time. If the levels are unequally spaced, *level values* can be specified in parentheses after the number of levels in the **REPEATED** statement. For example, if five levels of a drug corresponding to 1, 2, 5, 10, and 20 milligrams are administered to different treatment groups, represented by the variable group, the statements

```
proc glm;
  class group;
  model r1-r5=group / nouni;
  repeated dose 5 (1 2 5 10 20) polynomial / summary printm;
run;
```

produce the following M matrix:

$$M = \begin{bmatrix} -0.4250 & -0.3606 & -0.1674 & 0.1545 & 0.7984 \\ 0.4349 & 0.2073 & -0.3252 & -0.7116 & 0.3946 \\ -0.4331 & 0.1366 & 0.7253 & -0.5108 & 0.0821 \\ 0.4926 & -0.7800 & 0.3743 & -0.0936 & 0.0066 \end{bmatrix}$$

The **SUMMARY** option in this example provides univariate ANOVAs for the variables defined by the rows of this M matrix. In this case, they represent the linear, quadratic, cubic, and quartic trends for dose and are labeled dose_1, dose_2, dose_3, and dose_4, respectively.

HELMERT Transformation

Since the Helmert transformation compares a level of a repeated measure to the mean of subsequent levels, it is useful when interest lies in the point at which responses cease to change. For example, if four levels of a repeated measures factor represent responses to treatments administered over time to males and females, the statements

```
proc glm;
  class sex;
  model resp1-resp4=sex / nouni;
  repeated trtmnt 4 helmert / canon printm;
run;
```

produce the following M matrix:

$$M = \begin{bmatrix} 1 & -0.33333 & -0.33333 & -0.33333 \\ 0 & 1 & -0.50000 & -0.50000 \\ 0 & 0 & 1 & -1 \end{bmatrix}$$

MEAN Transformation

This transformation can be useful in the same types of situations in which the **CONTRAST** transformation is useful. If you substitute the following statement for the **REPEATED** statement shown in the **CONTRAST Transformation** section,

```
repeated drug 5 mean / printm;
```

the following **M** matrix is produced:

$$M = \begin{bmatrix} 1 & -0.25 & -0.25 & -0.25 & -0.25 \\ -0.25 & 1 & -0.25 & -0.25 & -0.25 \\ -0.25 & -0.25 & 1 & -0.25 & -0.25 \\ -0.25 & -0.25 & -0.25 & 1 & -0.25 \end{bmatrix}$$

As with the **CONTRAST** transformation, if you want to omit a level other than the last, you can specify it in parentheses after the keyword **MEAN** in the **REPEATED** statement.

PROFILE Transformation

When a repeated measure represents a series of factors administered over time, but a polynomial response is unreasonable, a profile transformation might prove useful. As an example, consider a training program in which four different methods are employed to teach students at several different schools. The repeated measure is the score on tests administered after each of the methods is completed. The statements

```
proc glm;
  class school;
  model t1-t4=school / nouni;
  repeated method 4 profile / summary nom printm;
run;
```

produce the following **M** matrix:

$$M = \begin{bmatrix} 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & -1 \end{bmatrix}$$

To determine the point at which an improvement in test scores takes place, you can examine the analyses of variance for the transformed variables representing the differences between adjacent tests. These analyses are requested by the **SUMMARY** option in the **REPEATED** statement, and the variables are labeled **METHOD.1**, **METHOD.2**, and **METHOD.3**.

Random-Effects Analysis

When some model effects are random (that is, assumed to be sampled from a normal population of effects), you can specify these effects in the **RANDOM** statement in order to compute the expected values of mean squares for various model effects and contrasts and, optionally, to perform random-effects analysis of variance tests.

PROC GLM versus PROC MIXED for Random-Effects Analysis

Other SAS procedures that can be used to analyze models with random effects include the MIXED and VARCOMP procedures. Note that, for these procedures, the random-effects specification is an integral part of the model, affecting how both random and fixed effects are fit; for PROC GLM, the random effects are treated in a *post hoc* fashion after the complete fixed-effect model is fit. This distinction affects other features in the GLM procedure, such as the results of the LSMEANS and ESTIMATE statements. These features assume that all effects are fixed, so that all tests and estimability checks for these statements are based on a fixed-effects model, even when you use a RANDOM statement. Standard errors for estimates and LS-means based on the fixed-effects model might be significantly smaller than those based on a true random-effects model; in fact, some functions that are estimable under a true random-effects model might not even be estimable under the fixed-effects model. Therefore, you should use the MIXED procedure to compute tests involving these features that take the random effects into account; see Chapter 78, “The MIXED Procedure,” for more information.

Note that, for balanced data, the test statistics computed when you specify the TEST option in the RANDOM statement have an exact F distribution only when the design is balanced; for unbalanced designs, the p values for the F tests are approximate. For balanced data, the values obtained by PROC GLM and PROC MIXED agree; for unbalanced data, they usually do not.

Computation of Expected Mean Squares for Random Effects

The RANDOM statement in PROC GLM declares one or more effects in the model to be random rather than fixed. By default, PROC GLM displays the coefficients of the expected mean squares for all terms in the model. In addition, when you specify the TEST option in the RANDOM statement, the procedure determines what tests are appropriate and provides F ratios and probabilities for these tests.

The expected mean squares are computed as follows. Consider the model

$$Y = X_0\beta_0 + X_1\beta_1 + \cdots + X_k\beta_k + \epsilon$$

where β_0 represents the fixed effects and $\beta_1, \beta_2, \dots, \epsilon$ represent the random effects. Random effects are assumed to be normally and independently distributed. For any $\mathbf{L}$ in the row space of $\mathbf{X} = (X_0 \mid X_1 \mid X_2 \mid \cdots \mid X_k)$, the expected value of the sum of squares for $\mathbf{L}\beta$ is

$$E(SS_L) = \beta_0' C_0' C_0 \beta_0 + SSQ(C_1)\sigma_1^2 + SSQ(C_2)\sigma_2^2 + \cdots + SSQ(C_k)\sigma_k^2 + \text{rank}(\mathbf{L})\sigma_\epsilon^2$$

where $\mathbf{C}$ is of the same dimensions as $\mathbf{L}$ and is partitioned as the $\mathbf{X}$ matrix. In other words,

$$\mathbf{C} = (\mathbf{C}_0 \mid \mathbf{C}_1 \mid \cdots \mid \mathbf{C}_k)$$

Furthermore, $\mathbf{C} = \mathbf{M}\mathbf{L}$, where $\mathbf{M}$ is the inverse of the lower triangular Cholesky decomposition matrix of $\mathbf{L}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{L}'$. $SSQ(\mathbf{A})$ is defined as $\text{tr}(\mathbf{A}'\mathbf{A})$.

For the model in the following MODEL statement

```
model Y=A B(A) C A*C;
random B(A);
```

with B(A) declared as random, the expected mean square of each effect is displayed as

$$\text{Var}(\text{Error}) + \text{constant} \times \text{Var}(\text{B(A)}) + Q(\mathbf{A}, \mathbf{C}, \mathbf{A} * \mathbf{C})$$

If any fixed effects appear in the expected mean square of an effect, the letter Q followed by the list of fixed effects in the expected value is displayed. The actual numeric values of the quadratic form (Q matrix) can be displayed using the Q option.

To determine appropriate means squares for testing the effects in the model, the TEST option in the RANDOM statement performs the following steps:

1. First, it forms a matrix of coefficients of the expected mean squares of those effects that were declared to be random.
2. Next, for each effect in the model, it determines the combination of these expected mean squares that produce an expectation that includes all the terms in the expected mean square of the effect of interest except the one corresponding to the effect of interest. For example, if the expected mean square of an effect A*B is

$$\text{Var}(\text{Error}) + 3 \times \text{Var}(\text{A}) + \text{Var}(\text{A} * \text{B})$$

PROC GLM determines the combination of other expected mean squares in the model that has expectation

$$\text{Var}(\text{Error}) + 3 \times \text{Var}(\text{A})$$

3. If the preceding criterion is met by the expected mean square of a single effect in the model (as is often the case in balanced designs), the *F* test is formed directly. In this case, the mean square of the effect of interest is used as the numerator, the mean square of the single effect with an expected mean square that satisfies the criterion is used as the denominator, and the degrees of freedom for the test are simply the usual model degrees of freedom.
4. When more than one mean square must be combined to achieve the appropriate expectation, an approximation is employed to determine the appropriate degrees of freedom (Satterthwaite 1946). When effects other than the effect of interest are listed after the Q in the output, tests of hypotheses involving the effect of interest are not valid unless all other fixed effects involved in it are assumed to be zero. When tests such as these are performed by using the TEST option in the RANDOM statement, a note is displayed reminding you that further assumptions are necessary for the validity of these tests. Remember that although the tests are not valid unless these assumptions are made, this does not provide a basis for these assumptions to be true. The particulars of a given experiment must be examined to determine whether the assumption is reasonable.

For further theoretical discussion, see Goodnight and Speed (1978), Milliken and Johnson (1984, Chapters 22 and 23), and Hocking (1985).

Sum-to-Zero Assumptions

The formulation and parameterization of the expected mean squares for random effects in mixed models are ongoing items of controversy in the statistical literature. Confusion arises over whether or not to assume that terms involving fixed effects sum to zero. Cornfield and Tukey (1956); Winer (1971), and others assume that they do sum to zero; Searle (1971); Hocking (1973), and others (including PROC GLM) do not.

Different assumptions about these sum-to-zero constraints can lead to different expected mean squares for certain terms, and hence to different *F* and *p* values.

For arguments in favor of not assuming that terms involving fixed effects sum to zero, see Section 9.7 of Searle (1971) and Sections 1 and 4 of McLean, Sanders, and Stroup (1991). Other references are Hartley and Searle (1969) and Searle, Casella, and McCulloch (1992).

Computing Type I, II, and IV Expected Mean Squares

When you use the **RANDOM** statement, by default the GLM procedure produces the Type III expected mean squares for model effects and for contrasts specified before the **RANDOM** statement. In order to obtain expected values for other types of mean squares, you need to specify which types of mean squares are of interest in the **MODEL** statement. For example, in order to obtain the Type IV expected mean squares for effects in the **RANDOM** and **CONTRAST** statements, specify the **SS4** option in the **MODEL** statement. If you want both Type III and Type IV expected mean squares, specify both the **SS3** and **SS4** options in the **MODEL** statement. Since the estimable function basis is not automatically calculated for Type I and Type II SS, the **E1** (for Type I) or **E2** (for Type II) option must be specified in the **MODEL** statement in order for the **RANDOM** statement to produce the expected mean squares for the Type I or Type II sums of squares. Note that it is important to list the fixed effects first in the **MODEL** statement when requesting the Type I expected mean squares.

For example, suppose you have a two-way design with factors A and B in which the main effect for B and the interaction are random. In order to compute the Type III expected mean squares (in addition to the fixed-effect analysis), you can use the following statements:

```
proc glm;
  class A B;
  model Y = A B A*B;
  random B A*B;
run;
```

Suppose you use the **SS4** option in the **MODEL** statement, as follows:

```
proc glm;
  class A B;
  model Y = A B A*B / ss4;
  random B A*B;
run;
```

Then only the Type IV expected mean squares are computed (as well as the Type IV fixed-effect tests). For the Type I expected mean squares, you can use the following statements:

```
proc glm;
  class A B;
  model Y = A B A*B / e1;
  random B A*B;
run;
```

For each of these cases, in order to perform random-effect analysis of variance tests for each effect specified in the model, you need to specify the **TEST** option in the **RANDOM** statement, as follows:

```
proc glm;
  class A B;
  model Y = A B A*B;
  random B A*B / test;
run;
```

The GLM procedure automatically determines the appropriate error term for each test, based on the expected mean squares.

Missing Values

For an analysis involving one dependent variable, PROC GLM uses an observation if values are nonmissing for that dependent variable and all the classification variables.

For an analysis involving multiple dependent variables without the [MANOVA](#) or [REPEATED](#) statement, or without the [MANOVA](#) option in the [PROC GLM](#) statement, a missing value in one dependent variable does not eliminate the observation from the analysis of other nonmissing dependent variables. On the other hand, for an analysis with the [MANOVA](#) or [REPEATED](#) statement, or with the [MANOVA](#) option in the PROC GLM statement, PROC GLM uses an observation if values are nonmissing for all dependent variables and all the variables used in independent effects.

During processing, the GLM procedure groups the dependent variables by their pattern of missing values across observations so that sums and crossproducts can be collected in the most efficient manner.

If your data have different patterns of missing values among the dependent variables, interactivity is disabled. This can occur when some of the variables in your data set have missing values and either of the following conditions obtain:

- You do not use the [MANOVA](#) option in the PROC GLM statement.
- You do not use a [MANOVA](#) or [REPEATED](#) statement before the first RUN statement.

Note that the REG procedure handles missing values differently in this case; see Chapter 99, “[The REG Procedure](#),” for more information.

Computational Resources

Memory

For large problems, most of the memory resources are required for holding the $X'X$ matrix of the sums and crossproducts. The section “[Parameterization of PROC GLM Models](#)” on page 3673 describes how columns of the X matrix are allocated for various types of effects. For each level that occurs in the data for a combination of classification variables in a given effect, a row and a column for $X'X$ are needed.

The following example illustrates the calculation. Suppose A has 20 levels, B has 4 levels, and C has 3 levels. Then consider the model

```
proc glm;
  class A B C;
  model Y1 Y2 Y3=A B A*B C A*C B*C A*B*C X1 X2;
run;
```

The $X'X$ matrix (bordered by $X'Y$ and $Y'Y$) can have as many as 425 rows and columns:

1	for the intercept term
20	for A
4	for B
80	for A * B
3	for C
60	for A * C
12	for B * C
240	for A * B * C
2	for X1 and X2 (continuous variables)
3	for Y1, Y2, and Y3 (dependent variables)

The matrix has 425 rows and columns only if all combinations of levels occur for each effect in the model. For m rows and columns, $8m^2$ bytes are needed for crossproducts. In this case, $8 \cdot 425^2 = 1,445,000$ bytes, or about $1,445,000/1024 = 1411K$.

The required memory grows as the square of the number of columns of X ; most of the memory is for the A*B*C interaction. Without A*B*C, you have 185 columns and need 268K for $X'X$. Without either A*B*C or A*B, you need 86K. If A is recoded to have 10 levels, then the full model has only 220 columns and requires 378K.

The second time that a large amount of memory is needed is when Type III, Type IV, or contrast sums of squares are being calculated. This memory requirement is a function of the number of degrees of freedom of the model being analyzed and the maximum degrees of freedom for any single source. Let Rank equal the sum of the model degrees of freedom, MaxDF be the maximum number of degrees of freedom for any single source, and N_y be the number of dependent variables in the model. Then the memory requirement in bytes is 8 times

$$\begin{aligned}
 N_y \times \text{Rank} &+ (\text{Rank} \times (\text{Rank} + 1)) / 2 \\
 &+ \text{MaxDF} \times \text{Rank} \\
 &+ (\text{MaxDF} \times (\text{MaxDF} + 1)) / 2 \\
 &+ \text{MaxDF} \times N_y
 \end{aligned}$$

The first two components of this formula are for the estimable model coefficients and their variance; the rest correspond to L , $L(X'X)^{-1}L'$, and Lb in the computation of $SS(L\beta = 0) = (Lb)'(L(X'X)^{-1}L')^{-1}(Lb)$. If the operating system enables SAS to run parallel computational threads on multiple CPUs, then GLM will attempt to allocate another $8 \times \text{Rank} \times \text{Rank}$ bytes in order to perform these calculations in parallel. If this much memory is not available, then the estimability calculations are performed in a single thread.

Unfortunately, these quantities are not available when the $X'X$ matrix is being constructed, so PROC GLM might occasionally request additional memory even after you have increased the memory allocation available to the program.

If you have a large model that exceeds the memory capacity of your computer, these are your options:

- Eliminate terms, especially high-level interactions.
- Reduce the number of levels for variables with many levels.
- Use the [ABSORB](#) statement for parts of the model that are large.
- Use the [REPEATED](#) statement for repeated measures variables.
- Use PROC ANOVA or PROC REG rather than PROC GLM, if your design allows.

A related limitation is that for any model effect involving classification variables (interactions as well as main effects), the number of levels cannot exceed 32,767. This is because GLM internally indexes effect levels with signed short (16-bit) integers, for which the maximum value is $2^{15} - 1 = 32,767$.

CPU Time

Typically, if the GLM procedure requires a lot of CPU time, it will be for one of several reasons. Suppose that the input data has n rows (observations) and the model has E effects that together produce a design matrix X with m columns. Then if m or n is relatively large, the procedure might spend a lot of time in any of the following areas:

- collecting the sums of squares and crossproducts
- solving the normal equations
- computing the Type III tests

The time required for collecting sums and crossproducts is difficult to calculate because it is a complicated function of the model. The worst case occurs if all columns are continuous variables, involving $nm^2/2$ multiplications and additions. If the columns are levels of a classification, then only m sums might be needed, but a significant amount of time might be spent in look-up operations. Solving the normal equations requires time for approximately $m^3/2$ multiplications and additions, and the number of operations required to compute the Type III tests is also proportional to both E and m^3 .

Suppose that you know that Type IV sums of squares are appropriate for the model you are analyzing (for example, if your design has no missing cells). You can specify the [SS4](#) option in your [MODEL](#) statement, which saves CPU time by requesting the Type IV sums of squares instead of the more computationally burdensome Type III sums of squares. This proves especially useful if you have a factor in your model that has many levels and is involved in several interactions.

If the operating system enables SAS to run parallel computational threads on multiple CPUs, then both the solution of the normal equations and the computation of Type III tests can take advantage of this to reduce the computational time for large models. In solving the normal equations, the fundamental row sweep operations (Goodnight 1979) are performed in parallel. In computing the Type III tests, both the orthogonalization for the estimable functions and the sums of squares calculation have been parallelized (if there is sufficient memory).

The reduction in computational time due to parallel processing depends on the size of the model, the number of processors, and the parallel architecture of the operating system. If the model is large enough that the overwhelming proportion of CPU time for the procedure is accounted for in solving the normal equations and/or computing the Type III tests, then you can expect a reduction in computational time approximately

inversely proportional to the number of CPUs. However, as you increase the number of processors, the efficiency of this scaling can be reduced by several effects. One mitigating factor is a purely mathematical one known as “Amdahl’s law,” which is related to the fact that only part of the processing time for the procedure can be parallelized. Even taking Amdahl’s law into account, the parallelization efficiency can be reduced by cache effects related to how fast the multiple processors can access memory. See Cohen (2002) for a discussion of these issues. For additional information about parallel processing in SAS, see the chapter on “Support for Parallel Processing” in *SAS Language Reference: Concepts*.

Computational Method

Let $\mathbf{X}$ represent the $n \times p$ design matrix and $\mathbf{Y}$ the $n \times 1$ vector of dependent variables. (See the section “Parameterization of PROC GLM Models” on page 3673 for information about how $\mathbf{X}$ is formed from your model specification.)

The normal equations $\mathbf{X}'\mathbf{X}\boldsymbol{\beta} = \mathbf{X}'\mathbf{Y}$ are solved using a modified sweep routine that produces a generalized inverse $(\mathbf{X}'\mathbf{X})^-$ and a solution $\mathbf{b} = (\mathbf{X}'\mathbf{X})^- \mathbf{X}'\mathbf{y}$. The modification is that rows and columns corresponding to diagonal elements that are found during sweeping to be zero (or within the expected level of numerical error of zero) are zeroed out. The $(\mathbf{X}'\mathbf{X})^-$ produced by this procedure satisfies the following two equations:

$$\begin{aligned}(\mathbf{X}'\mathbf{X}) (\mathbf{X}'\mathbf{X})^- (\mathbf{X}'\mathbf{X}) &= (\mathbf{X}'\mathbf{X}) \\ (\mathbf{X}'\mathbf{X})^- (\mathbf{X}'\mathbf{X}) (\mathbf{X}'\mathbf{X})^- &= (\mathbf{X}'\mathbf{X})^-\end{aligned}$$

Pringle and Rayner (1971) call a generalized inverse with these characteristics a g_2 -inverse, and this is the term usually used in SAS documentation and output. Urquhart (1968) uses the term *reflexive g-inverse* to emphasize that $(\mathbf{X}'\mathbf{X})^-$ is a generalized inverse of $\mathbf{X}'\mathbf{X}$ in the same way that $\mathbf{X}'\mathbf{X}$ is a generalized inverse of $(\mathbf{X}'\mathbf{X})^-$. Note that a g_2 -inverse is not necessarily unique: if $\mathbf{X}'\mathbf{X}$ is singular, then sweeping the matrix in a different order will result in a different g_2 -inverse that also satisfies the two preceding equations.

For each effect in the model, a matrix $\mathbf{L}$ is computed such that the rows of $\mathbf{L}$ are estimable. Tests of the hypothesis $\mathbf{L}\boldsymbol{\beta} = \mathbf{0}$ are then made by first computing

$$SS(\mathbf{L}\boldsymbol{\beta} = \mathbf{0}) = (\mathbf{Lb})'(\mathbf{L}(\mathbf{X}'\mathbf{X})^- \mathbf{L}')^{-1}(\mathbf{Lb})$$

and then computing the associated F value by using the mean squared error.

Output Data Sets

OUT= Data Set Created by the OUTPUT Statement

The **OUTPUT** statement produces an output data set that contains the following:

- all original data from the SAS data set input to PROC GLM
- the new variables corresponding to the diagnostic measures specified with statistics keywords in the **OUTPUT** statement (PREDICTED=, RESIDUAL=, and so on)

With multiple dependent variables, a name can be specified for any of the diagnostic measures for each of the dependent variables in the order in which they occur in the **MODEL** statement.

For example, suppose that the input data set **A** contains the variables **y1**, **y2**, **y3**, **x1**, and **x2**. Then you can use the following statements:

```
proc glm data=A;
  model y1 y2 y3=x1;
  output out=out p=y1hat y2hat y3hat
             r=y1resid lclm=y1lcl uclm=y1ucl;
run;
```

The output data set **out** contains **y1**, **y2**, **y3**, **x1**, **x2**, **y1hat**, **y2hat**, **y3hat**, **y1resid**, **y1lcl**, and **y1ucl**. The variable **x2** is output even though it is not used by PROC GLM. Although predicted values are generated for all three dependent variables, residuals are output for only the first dependent variable.

When any independent variable in the analysis (including all class variables) is missing for an observation, then all new variables that correspond to diagnostic measures are missing for the observation in the output data set.

When a dependent variable in the analysis is missing for an observation, then some new variables that correspond to diagnostic measures are missing for the observation in the output data set, and some are still available. Specifically, in this case, the new variables that correspond to **COOKD**, **COVRATIO**, **DFFITS**, **PRESS**, **R**, **RSTUDENT**, **STDR**, and **STUDENT** are missing in the output data set. The variables corresponding to **H**, **LCL**, **LCLM**, **P**, **STDI**, **STDP**, **UCL**, and **UCLM** are not missing.

OUT= Data Set Created by the LSMEANS Statement

The **OUT=** option in the **LSMEANS** statement produces an output data set that contains the following:

- the unformatted values of each classification variable specified in any effect in the **LSMEANS** statement
- a new variable, **LSMEAN**, which contains the LS-mean for the specified levels of the classification variables
- a new variable, **STDERR**, which contains the standard error of the LS-mean

The variances and covariances among the LS-means are also output when the **COV** option is specified along with the **OUT=** option. In this case, only one effect can be specified in the **LSMEANS** statement, and the following variables are included in the output data set:

- new variables, **COV1**, **COV2**, ..., **COV_n**, where *n* is the number of levels of the effect specified in the **LSMEANS** statement. These variables contain the covariances of each LS-mean with every other LS-mean.
- a new variable, **NUMBER**, which provides an index for each observation to identify the covariances that correspond to that observation. The covariances for the observation with **NUMBER** equal to *n* can be found in the variable **COV_n**.

OUTSTAT= Data Set

The **OUTSTAT=** option in the PROC GLM statement produces an output data set that contains the following:

- the BY variables, if any
- **_TYPE_**, a new character variable. **_TYPE_** can take the values 'SS1', 'SS2', 'SS3', 'SS4', or 'CONTRAST', corresponding to the various types of sums of squares generated, or the values 'CANCORR', 'STRUCTUR', or 'SCORE', if a canonical analysis is performed through the **MANOVA** statement and no **M=** matrix is specified.
- **_SOURCE_**, a new character variable. For each observation in the data set, **_SOURCE_** contains the name of the model effect or contrast label from which the corresponding statistics are generated.
- **_NAME_**, a new character variable. For each observation in the data set, **_NAME_** contains the name of one of the dependent variables in the model or, in the case of canonical statistics, the name of one of the canonical variables (CAN1, CAN2, and so forth).
- four new numeric variables: SS, DF, F, and PROB, containing sums of squares, degrees of freedom, *F* values, and probabilities, respectively, for each model or contrast sum of squares generated in the analysis. For observations resulting from canonical analyses, these variables have missing values.
- if there is more than one dependent variable, then variables with the same names as the dependent variables represent the following:
 - for **_TYPE_=SS1, SS2, SS3, SS4, or CONTRAST**, the crossproducts of the hypothesis matrices
 - for **_TYPE_=CANCORR**, canonical correlations for each variable
 - for **_TYPE_=STRUCTUR**, coefficients of the total structure matrix
 - for **_TYPE_=SCORE**, raw canonical score coefficients

The output data set can be used to perform special hypothesis tests (for example, with the IML procedure in SAS/IML software), to reformat output, to produce canonical variates (through the SCORE procedure), or to rotate structure matrices (through the FACTOR procedure).

Displayed Output

The GLM procedure produces the following output by default:

- The overall analysis-of-variance table breaks down the Total Sum of Squares for the dependent variable into the portion attributed to the Model and the portion attributed to Error.
- The Mean Square term is the Sum of Squares divided by the degrees of freedom (DF).
- The Mean Square for Error is an estimate of σ^2 , the variance of the true errors.
- The *F* Value is the ratio produced by dividing the Mean Square for the Model by the Mean Square for Error. It tests how well the model as a whole (adjusted for the mean) accounts for the dependent variable's behavior. An *F* test is a joint test to determine that all parameters except the intercept are zero.

- A small significance probability, $\Pr > F$, indicates that some linear function of the parameters is significantly different from zero.
- R-Square, R^2 , measures how much variation in the dependent variable can be accounted for by the model. R square, which can range from 0 to 1, is the ratio of the sum of squares for the model to the corrected total sum of squares. In general, the larger the value of R square, the better the model's fit.
- Coeff Var, the coefficient of variation, which describes the amount of variation in the population, is 100 times the standard deviation estimate of the dependent variable, Root MSE (Mean Square for Error), divided by the Mean. The coefficient of variation is often a preferred measure because it is unitless.
- Root MSE estimates the standard deviation of the dependent variable (or equivalently, the error term) and equals the square root of the Mean Square for Error.
- Mean is the sample mean of the dependent variable.

These tests are used primarily in analysis-of-variance applications:

- The Type I SS (sum of squares) measures incremental sums of squares for the model as each variable is added.
- The Type III SS is the sum of squares for a balanced test of each effect, adjusted for every other effect.

These items are used primarily in regression applications:

- The Estimates for the model Parameters (the intercept and the coefficients)
- t Value is the Student's t value for testing the null hypothesis that the parameter (if it is estimable) equals zero.
- The significance level, $\Pr > |t|$, is the probability of getting a larger value of t if the parameter is truly equal to zero. A very small value for this probability leads to the conclusion that the independent variable contributes significantly to the model.
- The Standard Error is the square root of the estimated variance of the estimate of the true value of the parameter.

Other portions of output are discussed in the following examples.

ODS Table Names

PROC GLM assigns a name to each table 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 47.15](#). For more information about ODS, see Chapter 20, “Using the Output Delivery System.”

Table 47.15 ODS Tables Produced by PROC GLM

ODS Table Name	Description	Statement / Option
Aliasing	Type 1,2,3,4 aliasing structure	MODEL / (E1 E2 E3 or E4) and ALIASING
AltErrContrasts	ANOVA table for contrasts with alternative error	CONTRAST / E=
AltErrTests	ANOVA table for tests with alternative error	TEST / E=
Bartlett	Bartlett's homogeneity of variance test	MEANS / HOVTEST=BARTLETT
CLDiffs	Multiple comparisons of pairwise differences	MEANS / CLDIFF or DUNNETT or (Unequal cells and not LINES)
CLDiffsInfo	Information for multiple comparisons of pairwise differences	MEANS / CLDIFF or DUNNETT or (Unequal cells and not LINES)
CLMeans	Multiple comparisons of means with confidence/comparison interval	MEANS / CLM
CLMeansInfo	Information for multiple comparison of means with confidence/comparison interval	MEANS / CLM
CanAnalysis	Canonical analysis	(MANOVA or REPEATED) / CANONICAL
CanCoef	Canonical coefficients	(MANOVA or REPEATED) / CANONICAL
CanStructure	Canonical structure	(MANOVA or REPEATED) / CANONICAL
CharStruct	Characteristic roots and vectors	(MANOVA / not CANONICAL) or (REPEATED / PRINTRV)
ClassLevels	Classification variable levels	CLASS statement
ContrastCoef	L matrix for contrast or estimate	CONTRAST / E or ESTIMATE / E
Contrasts	ANOVA table for contrasts	CONTRAST statement
DependentInfo	Simultaneously analyzed dependent variables	default when there are multiple dependent variables with different patterns of missing values
Diff	PDiff matrix of least squares means	LSMEANS / PDIFF=ALL and more than two LS-means
Epsilons	Greenhouse-Geisser and Huynh-Feldt epsilons	REPEATED statement
ErrorSSCP	Error SSCP matrix	(MANOVA or REPEATED) / PRINTE
EstFunc	Type 1,2,3,4 estimable functions	MODEL / (E1 E2 E3 or E4)
Estimates	Estimate statement results	ESTIMATE statement
ExpectedMeanSquares	Expected mean squares	RANDOM statement
FitStatistics	R-Square, Coeff Var, Root MSE, and dependent mean	default

Table 47.15 *continued*

ODS Table Name	Description	Statement / Option
GAliaising	General form of aliasing structure	MODEL / E and ALIASING
GEstFunc	General form of estimable functions	MODEL / E
HOVFTest	Homogeneity of variance ANOVA	MEANS / HOVTEST
HypothesisSSCP	Hypothesis SSCP matrix	(MANOVA or REPEATED) / PRINTH
InvXPX	$\text{inv}(\mathbf{X}'\mathbf{X})$ matrix	MODEL / INVERSE
LSMeanCL	Confidence interval for LS-means	LSMEANS / CL
LSMeanCoef	Coefficients of least squares means	LSMEANS / E
LSMeanDiffCL	Confidence interval for LS-mean differences	LSMEANS / PDIFF and CL
LSMeans	Least squares means	LSMEANS statement
LSMLines	Least squares means comparison lines	LSMEANS / PDIFF=ALL LINES
MANOVATransform	Multivariate transformation matrix	MANOVA / M=
MCLines	Multiple comparisons LINES output	MEANS / LINES or ((DUNCAN or WALLER or SNK or REGWQ) and not (CLDIFF or CLM)) or (Equal cells and not CLDIFF)
MCLinesInfo	Information for multiple comparison LINES output	MEANS / LINES or ((DUNCAN or WALLER or SNK or REGWQ) and not (CLDIFF or CLM)) or (Equal cells and not CLDIFF)
MCLinesRange	Ranges for multiple range MC tests	MEANS / LINES or ((DUNCAN or WALLER or SNK or REGWQ) and not (CLDIFF or CLM)) or (Equal cells and not CLDIFF)
MatrixRepresentation	$\mathbf{X}$ matrix element representation	as needed for other options
Means	Group means	MEANS statement
ModelANOVA	ANOVA for model terms	default
MultStat	Multivariate tests	MANOVA statement
NObs	Number of observations	default
OverallANOVA	Overall ANOVA	default
OverallEffectSize	Effect size measures for overall ANOVA	MODEL / EFFECTSIZE
ParameterEstimates	Estimated linear model coefficients	MODEL / SOLUTION
PartialCorr	Partial correlation matrix	(MANOVA or REPEATED) / PRINTE
PredictedInfo	Predicted values info	MODEL / P or CLM or CLI

Table 47.15 *continued*

ODS Table Name	Description	Statement / Option
PredictedValues	Predicted values	MODEL / P or CLM or CLI
QForm	Quadratic form for expected mean squares	RANDOM / Q
RandomModelANOVA	Random-effect tests	RANDOM / TEST
RepeatedLevelInfo	Correspondence between dependents and repeated measures levels	REPEATED statement
RepeatedTransform	Repeated measures transformation matrix	REPEATED / PRINTM
SimDetails	Details of difference quantile simulation	LSMEANS / ADJUST=SIMULATE(REPORT)
SimResults	Evaluation of difference quantile simulation	LSMEANS / ADJUST=SIMULATE(REPORT)
SlicedANOVA	Sliced-effect ANOVA table	LSMEANS / SLICE
Sphericity	Sphericity tests	REPEATED / PRINTE
Tests	Summary ANOVA for specified MANOVA H= effects	MANOVA / H= SUMMARY
Tolerances	$X'X$ tolerances	MODEL / TOLERANCE
Welch	Welch's ANOVA	MEANS / WELCH
XPX	$X'X$ matrix	MODEL / XPX

With the **PDIF** or **TDIF** option in the **LSMEANS** statement, the *p/t*-values for differences are displayed in columns of the LSMeans table for **PDIF/TDIF=CONTROL** or **PDIF/TDIF=ANOM**, and for **PDIF/TDIF=ALL** when there are only two LS-means. Otherwise (for **PDIF/TDIF=ALL** when there are more than two LS-means), the *p/t*-values for differences are displayed in a separate table called Diff.

ODS Graphics

Statistical procedures use ODS Graphics to create graphs as part of their output. ODS Graphics is described in detail in Chapter 21, “Statistical Graphics Using ODS.”

Before you create graphs, ODS Graphics must be enabled (for example, by specifying the ODS GRAPHICS ON statement). For more information about enabling and disabling ODS Graphics, see the section “Enabling and Disabling ODS Graphics” on page 607 in Chapter 21, “Statistical Graphics Using ODS.”

The overall appearance of graphs is controlled by ODS styles. Styles and other aspects of using ODS Graphics are discussed in the section “A Primer on ODS Statistical Graphics” on page 606 in Chapter 21, “Statistical Graphics Using ODS.”

When ODS Graphics is enabled, then for particular models the GLM procedure produces default graphics as follows:

- If you specify a one-way analysis of variance model, with just one **CLASS** variable, the GLM procedure produces a grouped box plot of the response values versus the **CLASS** levels. For an example of the box plot, see the section “One-Way Layout with Means Comparisons” on page 972 in Chapter 26,

“The ANOVA Procedure.” Outliers are labeled by the observation number within the entire data set, including observations that are not used in the analysis. Therefore, the labels of the outliers in this box plot might differ from those in the box plot that is produced by the MEANS statement.

- If you specify a two-way analysis of variance model, with just two **CLASS** variables, the GLM procedure produces an interaction plot of the response values, with horizontal position representing one **CLASS** variable and marker style representing the other, and with predicted response values connected by lines that represent the two-way analysis. For an example of the interaction plot, see the section “PROC GLM for Unbalanced ANOVA” on page 3613.
- If you specify a model that has a single continuous predictor, the GLM procedure produces a fit plot of the response values versus the covariate values, with a curve representing the fitted relationship. For an example of the fit plot, see the section “PROC GLM for Quadratic Least Squares Regression” on page 3616.
- If you specify a model that has a two continuous predictors and no **CLASS** variables, the GLM procedure produces a panel of fit plots as in the single predictor case, with a plot of the response values versus one of the covariates at each of several values of the other covariate.
- If you specify an analysis of covariance model that has one or two **CLASS** variables and one continuous variable, the GLM procedure produces an analysis-of-covariance plot of the response values versus the covariate values, with lines representing the fitted relationship within each classification level. For an example of the analysis of covariance plot, see [Example 47.4](#).
- If you specify an **LSMEANS** statement with the **PDIF** option, the GLM procedure produces a plot appropriate for the type of LS-means comparison. For **PDIF**=ALL (which is the default if you specify only **PDIF**), the procedure produces a diffogram, which displays all pairwise LS-means differences and their significance. The display is also known as a “mean-mean scatter plot” (Hsu 1996). For **PDIF**=CONTROL, the procedure produces a display of each noncontrol LS-mean compared to the control LS-mean, with two-sided confidence intervals for the comparison. For **PDIF**=CONTROL and **PDIF**=CONTROLU, a similar display is produced, but with one-sided confidence intervals. Finally, for the **PDIF**=ANOM option, the procedure produces an “analysis of means” plot, comparing each LS-mean to the average LS-mean.
- If you specify a **MEANS** statement, the GLM procedure produces a grouped box plot of the response values versus the effect for which means are being calculated. Outliers are labeled by the observation number within only the data that have no missing values, thus excluding observations that are not used in the analysis. Therefore, the labels of the outliers in this plot might differ from the labels of outliers in the box plot that is the default for one-way models.

In addition to the default graphics mentioned previously, you can request plots that help you diagnose the quality of the fitted model.

- The **PLOTS**=DIAGNOSTICS option in the **PROC GLM** statement requests that a panel of summary diagnostics for the fit be displayed. The panel displays scatter plots of residuals, absolute residuals, studentized residuals, and observed responses by predicted values; studentized residuals by leverage; Cook’s *D* by observation; a Q-Q plot of residuals; a residual histogram; and a residual-fit spread plot.
- The **PLOTS**=RESIDUALS option in the **PROC GLM** statement requests scatter plots of the residuals against each continuous covariate.

ODS Graph Names

PROC GLM assigns a name to each graph it creates using ODS. You can use these names to refer to the graphs when you use ODS. The names are listed in Table 47.16. ODS Graphics must be enabled before requesting plots. For more information about ODS Graphics, see Chapter 21, “Statistical Graphics Using ODS.”

Table 47.16 Graphs Produced by PROC GLM

ODS Graph Name	Plot Description	Option
ANCOVAPlot	Analysis-of-covariance plot	Analysis-of-covariance model
ANOMPlot	Plot of LS-mean differences against average LS-mean	LSMEANS / PDIFF=ANOM
BoxPlot	Box plot of group means	One-way ANOVA model or MEANS statement
ContourFit	Plot of predicted response surface	Two-predictor response surface model
ControlPlot	Plot of LS-mean differences against a control level	LSMEANS / PDIFF=CONTROL
DiagnosticsPanel	Panel of summary diagnostics for the fit	PLOTS=DIAGNOSTICS
CooksDPlot	Cook's <i>D</i> plot	PLOTS=DIAGNOSTICS(UNPACK)
ObservedByPredicted	Observed by predicted	PLOTS=DIAGNOSTICS(UNPACK)
QQPlot	Residual Q-Q plot	PLOTS=DIAGNOSTICS(UNPACK)
ResidualByPredicted	Residual by predicted values	PLOTS=DIAGNOSTICS(UNPACK)
ResidualHistogram	Residual histogram	PLOTS=DIAGNOSTICS(UNPACK)
RFPlot	RF plot	PLOTS=DIAGNOSTICS(UNPACK)
RStudentByPredicted	Studentized residuals by predicted	PLOTS=DIAGNOSTICS(UNPACK)
RStudentByLeverage	RStudent by hat diagonals	PLOTS=DIAGNOSTICS(UNPACK)
DiffPlot	Plot of LS-mean pairwise differences	LSMEANS / PDIFF
IntPlot	Interaction plot	Two-way ANOVA model
FitPlot	Plot of predicted response by predictor	Model with one continuous predictor
ResidualPlots	Plots of the residuals against each continuous covariate	PLOTS=RESIDUALS


```

contrast 'Glacial vs. drift'   Type   -1    0    1    1    0    0   -1;
contrast 'Clarion vs. Webster' Type   -1    0    0    0    0    0    1;
contrast "Knox vs. O'Neill"   Type    0    0    1   -1    0    0    0;

run;

means Type / waller regwq;

run;

```

Output 47.1.1 Analysis of Variance for Randomized Complete Blocks**Balanced Data from Randomized Complete Block****The GLM Procedure**

Class Level Information	
Class	Levels Values
Block	3 1 2 3
Type	7 Clarion Clinton Compost Knox O'Neill Wabash Webster

Number of Observations Read	21
Number of Observations Used	21

Balanced Data from Randomized Complete Block**The GLM Procedure****Dependent Variable: StemLength**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	8	142.1885714	17.7735714	10.80	0.0002
Error	12	19.7428571	1.6452381		
Corrected Total	20	161.9314286			

R-Square	Coeff Var	Root MSE	StemLength Mean
0.878079	3.939745	1.282668	32.55714

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Block	2	39.0371429	19.5185714	11.86	0.0014
Type	6	103.1514286	17.1919048	10.45	0.0004

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Block	2	39.0371429	19.5185714	11.86	0.0014
Type	6	103.1514286	17.1919048	10.45	0.0004

This analysis shows that the stem length is significantly different for the different soil types. In addition, there are significant differences in stem length among the three blocks in the experiment.

The GLM procedure is invoked again, this time with the **ORDER=DATA** option. This enables you to write accurate contrast statements more easily because you know the order SAS is using for the levels of the variable **Type**. The standard analysis is displayed again, this time including the tests for contrasts that you specified as well as the estimated parameters. These additional results are shown in [Output 47.1.2](#).

Output 47.1.2 Contrasts and Solutions**Balanced Data from Randomized Complete Block****The GLM Procedure****Dependent Variable: StemLength**

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Compost vs. others	1	29.24198413	29.24198413	17.77	0.0012
River soils vs. non	2	48.24694444	24.12347222	14.66	0.0006
Glacial vs. drift	1	22.14083333	22.14083333	13.46	0.0032
Clarion vs. Webster	1	1.70666667	1.70666667	1.04	0.3285
Knox vs. O'Neill	1	1.81500000	1.81500000	1.10	0.3143

Parameter	Estimate	Standard Error	t Value	Pr > t
Intercept	29.35714286	B 0.83970354	34.96	<.0001
Block 1	3.32857143	B 0.68561507	4.85	0.0004
Block 2	1.90000000	B 0.68561507	2.77	0.0169
Block 3	0.00000000	B .	.	.
Type Clarion	1.06666667	B 1.04729432	1.02	0.3285
Type Clinton	-0.80000000	B 1.04729432	-0.76	0.4597
Type Knox	3.80000000	B 1.04729432	3.63	0.0035
Type O'Neill	2.70000000	B 1.04729432	2.58	0.0242
Type Compost	-1.43333333	B 1.04729432	-1.37	0.1962
Type Wabash	4.86666667	B 1.04729432	4.65	0.0006
Type Webster	0.00000000	B .	.	.

Note: The X'X matrix has been found to be singular, and a generalized inverse was used to solve the normal equations. Terms whose estimates are followed by the letter 'B' are not uniquely estimable.

The contrast label, degrees of freedom, sum of squares, Mean Square, F Value, and Pr > F are shown for each contrast requested. In this example, the contrast results indicate the following inferences, at the 5% significance level:

- The stem length of plants grown in compost soil is significantly different from the average stem length of plants grown in other soils.
- The stem length of plants grown in river soils is significantly different from the average stem length of those grown in nonriver soils.
- The average stem length of plants grown in glacial soils (Clarion and Webster types) is significantly different from the average stem length of those grown in drift soils (Knox and O'Neill types).

- Stem lengths for Clarion and Webster types are not significantly different.
- Stem lengths for Knox and O'Neill types are not significantly different.

In addition to the estimates for the parameters of the model, the results of t tests about the parameters are also displayed. The 'B' following the parameter estimates indicates that the estimates are biased and do not represent a unique solution to the normal equations.

Output 47.1.3 Waller-Duncan tests

Balanced Data from Randomized Complete Block

The GLM Procedure

Waller-Duncan K-ratio t Test for StemLength

Note: This test minimizes the Bayes risk under additive loss and certain other assumptions.

Kratio	100
Error Degrees of Freedom	12
Error Mean Square	1.645238
F Value	10.45
Critical Value of t	2.12034
Minimum Significant Difference	2.2206

Means with the same letter are not significantly different.			
Waller Grouping	Mean	N	Type
A	35.967	3	Wabash
A			
A	34.900	3	Knox
A			
B	33.800	3	O'Neill
B			
B	32.167	3	Clarion
D	31.100	3	Webster
D			
D	30.300	3	Clinton
D			
D	29.667	3	Compost

Output 47.1.4 Ryan-Einot-Gabriel-Welsch Multiple Range Test**Balanced Data from Randomized Complete Block****The GLM Procedure****Ryan-Einot-Gabriel-Welsch Multiple Range Test for StemLength**

Note: This test controls the Type I experimentwise error rate.

Alpha	0.05
Error Degrees of Freedom	12
Error Mean Square	1.645238

Number of Means	2	3	4	5	6	7
Critical Range	2.9875528	3.2837322	3.4395625	3.5402383	3.5402383	3.6653133

Means with the same letter are not significantly different.					
REGWQ	Grouping	Mean	N	Type	
	A	35.967	3	Wabash	
	A				
B	A	34.900	3	Knox	
B	A				
B	A	C	33.800	3	O'Neill
B		C			
B	D	C	32.167	3	Clarion
	D	C			
	D	C	31.100	3	Webster
	D				
	D		30.300	3	Clinton
	D				
	D		29.667	3	Compost

The final two pages of output ([Output 47.1.3](#) and [Output 47.1.4](#)) present results of the Waller-Duncan and REGWQ multiple-comparison procedures. For each test, notes and information pertinent to the test are given in the output. The Type means are arranged from highest to lowest. Means with the same letter are not significantly different. For this example, while some pairs of means are significantly different, there are no clear equivalence classes among the different soils.

For an alternative method of analyzing and displaying mean differences, including high-resolution graphics, see [Example 47.3](#).

Example 47.2: Regression with Mileage Data

A car is tested for gas mileage at various speeds to determine at what speed the car achieves the highest gas mileage. A quadratic model is fit to the experimental data. The following statements produce [Output 47.2.1](#) through [Output 47.2.4](#).

```

title 'Gasoline Mileage Experiment';
data mileage;
    input mph mpg @@;
    datalines;
20 15.4
30 20.2
40 25.7
50 26.2  50 26.6  50 27.4
55 .
60 24.8
;

ods graphics on;
proc glm;
    model mpg=mph mph*mph / p clm;
run;
ods graphics off;

```

Output 47.2.1 Standard Regression Analysis**Gasoline Mileage Experiment****The GLM Procedure**

Number of Observations Read 8

Number of Observations Used 7

Gasoline Mileage Experiment**The GLM Procedure****Dependent Variable: mpg**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	2	111.8086183	55.9043091	77.96	0.0006
Error	4	2.8685246	0.7171311		
Corrected Total	6	114.6771429			

R-Square	Coeff Var	Root MSE	mpg Mean
0.974986	3.564553	0.846836	23.75714

Source	DF	Type I SS	Mean Square	F Value	Pr > F
mph	1	85.64464286	85.64464286	119.43	0.0004
mph*mph	1	26.16397541	26.16397541	36.48	0.0038

Source	DF	Type III SS	Mean Square	F Value	Pr > F
mph	1	41.01171219	41.01171219	57.19	0.0016
mph*mph	1	26.16397541	26.16397541	36.48	0.0038

Output 47.2.1 *continued*

Parameter	Estimate	Standard Error	t Value	Pr > t
Intercept	-5.985245902	3.18522249	-1.88	0.1334
mph	1.305245902	0.17259876	7.56	0.0016
mph*mph	-0.013098361	0.00216852	-6.04	0.0038

The overall F statistic is significant. The tests of mph and mph*mph in the Type I sums of squares show that both the linear and quadratic terms in the regression model are significant. The model fits well, with an R square of 0.97. The table of parameter estimates indicates that the estimated regression equation is

$$\text{mpg} = -5.9852 + 1.3052 \times \text{mph} - 0.0131 \times \text{mph}^2$$

Output 47.2.2 Results of Requesting the P and CLM Options

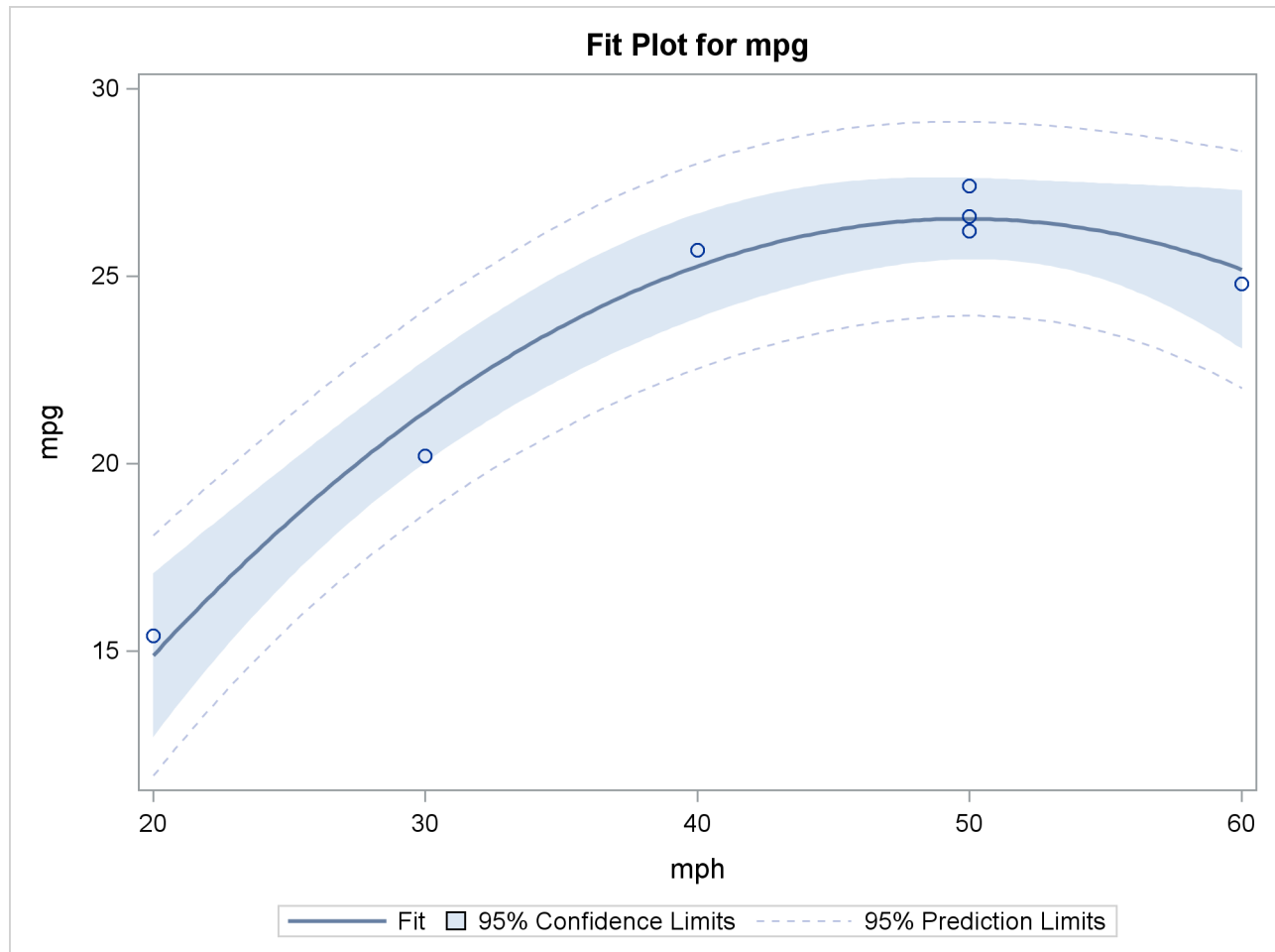
Observation	Observed	Predicted	Residual	95% Confidence Limits for Mean Predicted Value	
1	15.40000000	14.88032787	0.51967213	12.69701317	17.06364257
2	20.20000000	21.38360656	-1.18360656	20.01727192	22.74994119
3	25.70000000	25.26721311	0.43278689	23.87460041	26.65982582
4	26.20000000	26.53114754	-0.33114754	25.44573423	27.61656085
5	26.60000000	26.53114754	0.06885246	25.44573423	27.61656085
6	27.40000000	26.53114754	0.86885246	25.44573423	27.61656085
7 *	.	26.18073770	.	24.88679308	27.47468233
8	24.80000000	25.17540984	-0.37540984	23.05954977	27.29126990

The **P** and **CLM** options in the **MODEL** statement produce the table shown in [Output 47.2.2](#). For each observation, the observed, predicted, and residual values are shown. In addition, the 95% confidence limits for a mean predicted value are shown for each observation. Note that the observation with a missing value for mph is not used in the analysis, but predicted and confidence limit values are shown.

Output 47.2.3 Additional Results of Requesting the P and CLM Options

Sum of Residuals	-0.00000000
Sum of Squared Residuals	2.86852459
Sum of Squared Residuals - Error SS	-0.00000000
PRESS Statistic	23.18107335
First Order Autocorrelation	-0.54376613
Durbin-Watson D	2.94425592

The last portion of the output listing, shown in [Output 47.2.3](#), gives some additional information about the residuals. The Press statistic gives the sum of squares of predicted residual errors, as described in Chapter 4, “[Introduction to Regression Procedures](#).” The First Order Autocorrelation and the Durbin-Watson D statistic, which measures first-order autocorrelation, are also given.

Output 47.2.4 Plot of Mileage Data

Finally, the ODS GRAPHICS ON command in the previous statements enables ODS Graphics, which in this case produces the plot shown in [Output 47.2.4](#) of the actual and predicted values for the data, as well as a band representing the confidence limits for individual predictions. The quadratic relationship between mpg and mph is evident.

Example 47.3: Unbalanced ANOVA for Two-Way Design with Interaction

This example uses data from Kutner (1974, p. 98) to illustrate a two-way analysis of variance. The original data source is Afifi and Azen (1972, p. 166). These statements produce [Output 47.3.1](#) and [Output 47.3.2](#).

```

title 'Unbalanced Two-Way Analysis of Variance';
data a;
  input drug disease @;
  do i=1 to 6;
    input y @;
    output;
  end;
  datalines;

```

```

1 1 42 44 36 13 19 22
1 2 33 . 26 . 33 21
1 3 31 -3 . 25 25 24
2 1 28 . 23 34 42 13
2 2 . 34 33 31 . 36
2 3 3 26 28 32 4 16
3 1 . . 1 29 . 19
3 2 . 11 9 7 1 -6
3 3 21 1 . 9 3 .
4 1 24 . 9 22 -2 15
4 2 27 12 12 -5 16 15
4 3 22 7 25 5 12 .
;

proc glm;
  class drug disease;
  model y=drug disease drug*disease / ss1 ss2 ss3 ss4;
run;

```

Output 47.3.1 Classes and Levels for Unbalanced Two-Way Design**Unbalanced Two-Way Analysis of Variance****The GLM Procedure**

Class Level Information		
Class	Levels	Values
drug	4	1 2 3 4
disease	3	1 2 3

Number of Observations Read 72

Number of Observations Used 58

Output 47.3.2 Analysis of Variance for Unbalanced Two-Way Design**Unbalanced Two-Way Analysis of Variance****The GLM Procedure****Dependent Variable: y**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	11	4259.338506	387.212591	3.51	0.0013
Error	46	5080.816667	110.452536		
Corrected Total	57	9340.155172			

R-Square	Coeff Var	Root MSE	y Mean
0.456024	55.66750	10.50964	18.87931

Output 47.3.2 continued

Source	DF	Type I SS	Mean Square	F Value	Pr > F
drug	3	3133.238506	1044.412835	9.46	<.0001
disease	2	418.833741	209.416870	1.90	0.1617
drug*disease	6	707.266259	117.877710	1.07	0.3958

Source	DF	Type II SS	Mean Square	F Value	Pr > F
drug	3	3063.432863	1021.144288	9.25	<.0001
disease	2	418.833741	209.416870	1.90	0.1617
drug*disease	6	707.266259	117.877710	1.07	0.3958

Source	DF	Type III SS	Mean Square	F Value	Pr > F
drug	3	2997.471860	999.157287	9.05	<.0001
disease	2	415.873046	207.936523	1.88	0.1637
drug*disease	6	707.266259	117.877710	1.07	0.3958

Source	DF	Type IV SS	Mean Square	F Value	Pr > F
drug	3	2997.471860	999.157287	9.05	<.0001
disease	2	415.873046	207.936523	1.88	0.1637
drug*disease	6	707.266259	117.877710	1.07	0.3958

Note the differences among the four types of sums of squares. The Type I sum of squares for drug essentially tests for differences between the expected values of the arithmetic mean response for different drugs, unadjusted for the effect of disease. By contrast, the Type II sum of squares for drug measures the differences between arithmetic means for each drug after adjusting for disease. The Type III sum of squares measures the differences between predicted drug means over a balanced drug×disease population—that is, between the LS-means for drug. Finally, the Type IV sum of squares is the same as the Type III sum of squares in this case, since there are data for every drug-by-disease combination.

No matter which sum of squares you prefer to use, this analysis shows a significant difference among the four drugs, while the disease effect and the drug-by-disease interaction are not significant. As the previous discussion indicates, Type III sums of squares correspond to differences between LS-means, so you can follow up the Type III tests with a multiple-comparison analysis of the drug LS-means. Since the GLM procedure is interactive, you can accomplish this by submitting the following statements after the previous ones that performed the ANOVA.

```
lsmeans drug / pdiff=all adjust=tukey;
run;
```

Both the LS-means themselves and a matrix of adjusted p -values for pairwise differences between them are displayed; see [Output 47.3.3](#) and [Output 47.3.4](#).

Output 47.3.3 LS-Means for Unbalanced ANOVA**Unbalanced Two-Way Analysis of Variance**

The GLM Procedure
Least Squares Means
Adjustment for Multiple Comparisons: Tukey-Kramer

		LSMEAN
drug	y	LSMEAN
		Number
1	25.9944444	1
2	26.5555556	2
3	9.7444444	3
4	13.5444444	4

Output 47.3.4 Adjusted p -Values for Pairwise LS-Mean Differences

Least Squares Means for effect drug Pr > t for H0: LSMean(i)=LSMean(j)				
Dependent Variable: y				
i/j	1	2	3	4
1		0.9989	0.0016	0.0107
2	0.9989		0.0011	0.0071
3	0.0016	0.0011		0.7870
4	0.0107	0.0071	0.7870	

The multiple-comparison analysis shows that drugs 1 and 2 have very similar effects, and that drugs 3 and 4 are also insignificantly different from each other. Evidently, the main contribution to the significant drug effect is the difference between the 1/2 pair and the 3/4 pair.

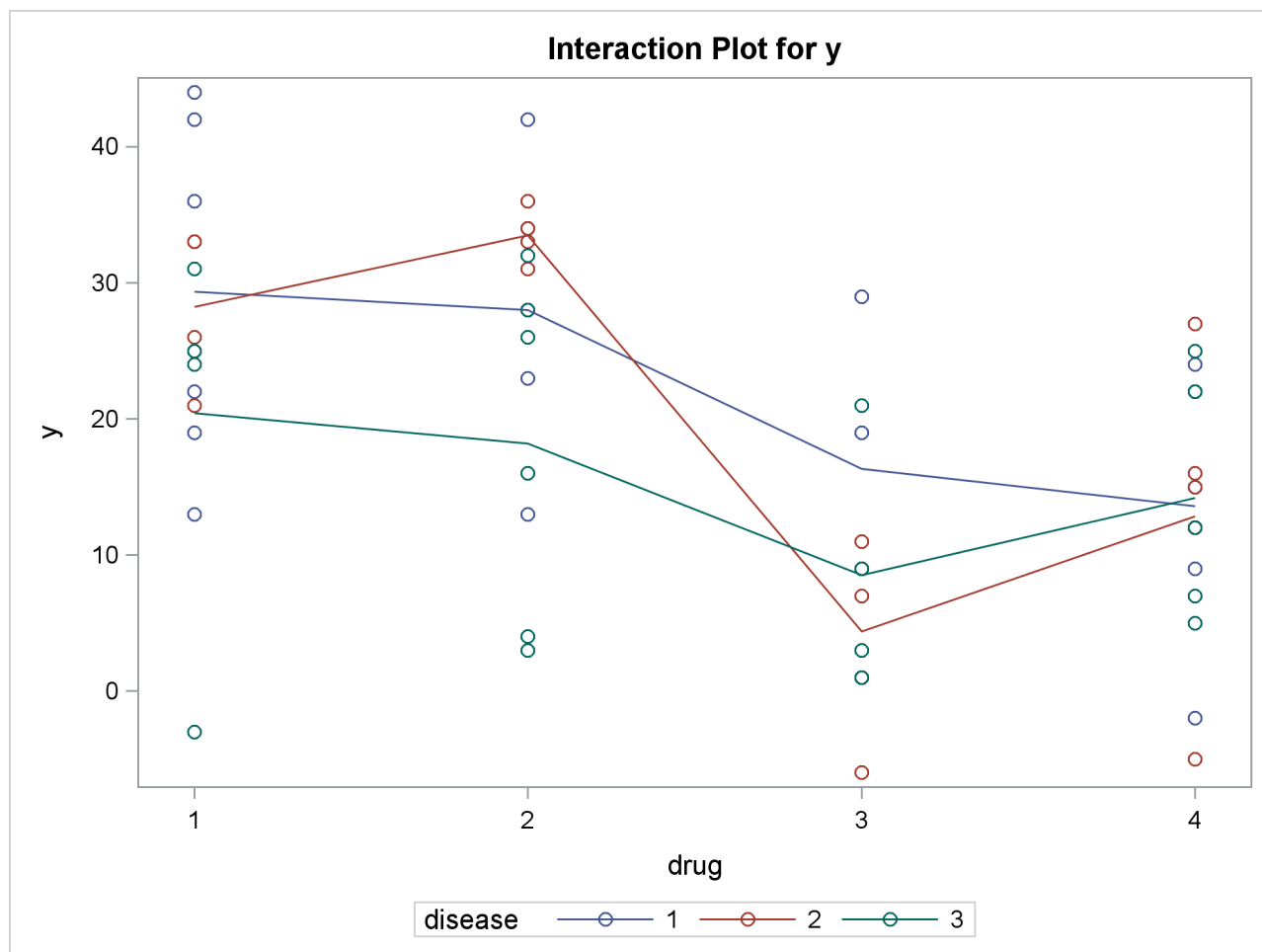
If ODS Graphics is enabled for the previous analysis, GLM also displays three additional plots by default:

- an interaction plot for the effects of disease and drug
- a mean plot of the drug LS-means
- a plot of the adjusted pairwise differences and their significance levels

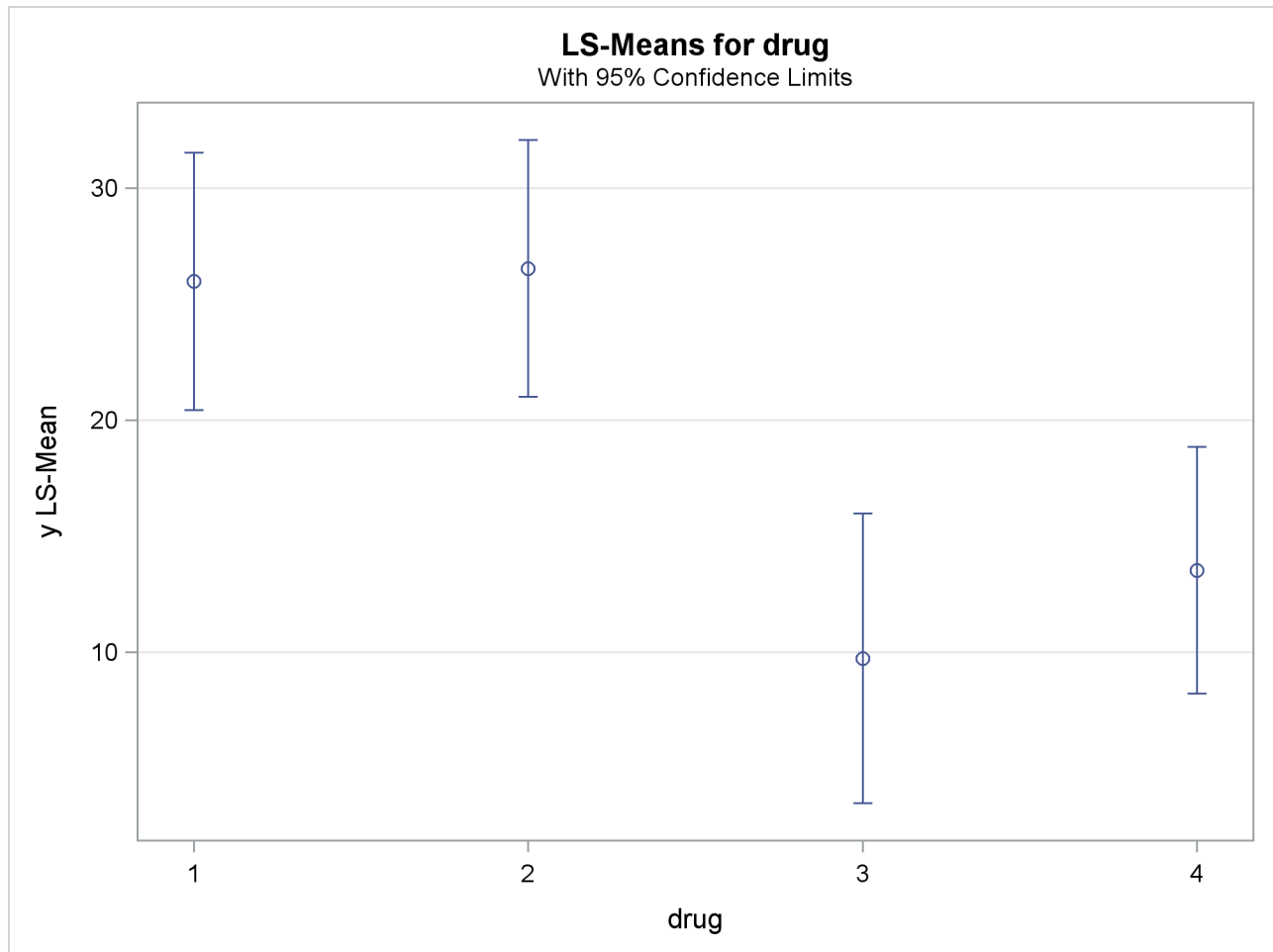
The following statements reproduce the previous analysis with ODS Graphics enabled. Additionally, the **PLOTS=MEANPLOT(CL)** option specifies that confidence limits for the LS-means should also be displayed in the mean plot. The graphical results are shown in [Output 47.3.5](#) through [Output 47.3.7](#).

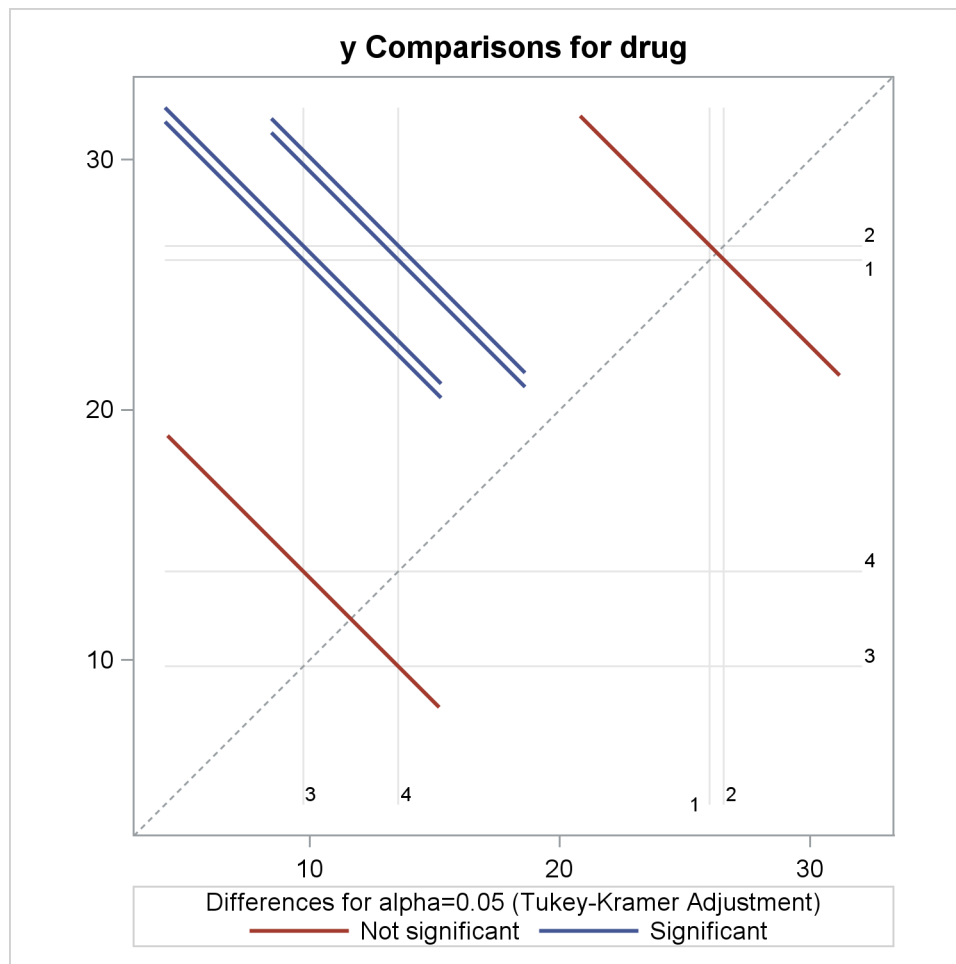
```
ods graphics on;
proc glm plot=meanplot(cl);
  class drug disease;
  model y=drug disease drug*disease;
  lsmeans drug / pdiff=all adjust=tukey;
run;
ods graphics off;
```

Output 47.3.5 Plot of Response by Drug and Disease



Output 47.3.6 Plot of Response LS-Means for Drug



Output 47.3.7 Plot of Response LS-Mean Differences for Drug

The significance of the drug differences is difficult to discern in the original data, as displayed in [Output 47.3.5](#), but the plot of just the LS-means and their individual confidence limits in [Output 47.3.6](#) makes it clearer. Finally, [Output 47.3.7](#) indicates conclusively that the significance of the effect of drug is due to the difference between the two drug pairs (1, 2) and (3, 4).

Example 47.4: Analysis of Covariance

Analysis of covariance combines some of the features of both regression and analysis of variance. Typically, a continuous variable (the covariate) is introduced into the model of an analysis-of-variance experiment.

Data in the following example are selected from a larger experiment on the use of drugs in the treatment of leprosy (Snedecor and Cochran 1967, p. 422).

Variables in the study are as follows:

Drug two antibiotics (A and D) and a control (F)
 PreTreatment a pretreatment score of leprosy bacilli
 PostTreatment a posttreatment score of leprosy bacilli

Ten patients are selected for each treatment (Drug), and six sites on each patient are measured for leprosy bacilli.

The covariate (a pretreatment score) is included in the model for increased precision in determining the effect of drug treatments on the posttreatment count of bacilli.

The following statements create the data set, perform a parallel-slopes analysis of covariance with PROC GLM, and compute Drug LS-means. These statements produce [Output 47.4.1](#) and [Output 47.4.2](#).

```
data DrugTest;
  input Drug $ PreTreatment PostTreatment @@;
  datalines;
A 11 6   A 8 0   A 5 2   A 14 8   A 19 11
A 6 4   A 10 13  A 6 1   A 11 8   A 3 0
D 6 0   D 6 2   D 7 3   D 8 1   D 18 18
D 8 4   D 19 14  D 8 9   D 5 1   D 15 9
F 16 13  F 13 10  F 11 18  F 9 5   F 21 23
F 16 12  F 12 5   F 12 16  F 7 1   F 12 20
;

proc glm data=DrugTest;
  class Drug;
  model PostTreatment = Drug PreTreatment / solution;
  lsmeans Drug / stderr pdiff cov out=adjmeans;
run;

proc print data=adjmeans;
run;
```

Output 47.4.1 Classes and Levels

The GLM Procedure

Class Level Information	
Class	Levels Values
Drug	3 A D F

Number of Observations Read	30
Number of Observations Used	30

Output 47.4.2 Overall Analysis of Variance**The GLM Procedure****Dependent Variable: PostTreatment**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	871.497403	290.499134	18.10	<.0001
Error	26	417.202597	16.046254		
Corrected Total	29	1288.700000			

R-Square	Coeff Var	Root MSE	PostTreatment Mean
0.676261	50.70604	4.005778	7.900000

This model assumes that the slopes relating posttreatment scores to pretreatment scores are parallel for all drugs. You can check this assumption by including the class-by-covariate interaction, Drug*PreTreatment, in the model and examining the ANOVA test for the significance of this effect. This extra test is omitted in this example, but it is insignificant, justifying the equal-slopes assumption.

In [Output 47.4.3](#), the Type I SS for Drug (293.6) gives the between-drug sums of squares that are obtained for the analysis-of-variance model PostTreatment=Drug. This measures the difference between arithmetic means of posttreatment scores for different drugs, disregarding the covariate. The Type III SS for Drug (68.5537) gives the Drug sum of squares adjusted for the covariate. This measures the differences between Drug LS-means, controlling for the covariate. The Type I test is highly significant ($p = 0.001$), but the Type III test is not. This indicates that, while there is a statistically significant difference between the arithmetic drug means, this difference is reduced to below the level of background noise when you take the pretreatment scores into account. From the table of parameter estimates, you can derive the least squares predictive formula model for estimating posttreatment score based on pretreatment score and drug:

$$\text{post} = \begin{cases} (-0.435 + -3.446) + 0.987 \cdot \text{pre}, & \text{if Drug=A} \\ (-0.435 + -3.337) + 0.987 \cdot \text{pre}, & \text{if Drug=D} \\ -0.435 + 0.987 \cdot \text{pre}, & \text{if Drug=F} \end{cases}$$

Output 47.4.3 Tests and Parameter Estimates

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Drug	2	293.6000000	146.8000000	9.15	0.0010
PreTreatment	1	577.8974030	577.8974030	36.01	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Drug	2	68.5537106	34.2768553	2.14	0.1384
PreTreatment	1	577.8974030	577.8974030	36.01	<.0001

Parameter	Estimate	Standard Error	t Value	Pr > t
Intercept	-0.434671164	2.47135356	-0.18	0.8617
Drug A	-3.446138280	1.88678065	-1.83	0.0793
Drug D	-3.337166948	1.85386642	-1.80	0.0835
Drug F	0.000000000	.	.	.
PreTreatment	0.987183811	0.16449757	6.00	<.0001

Output 47.4.4 displays the LS-means, which are, in a sense, the means adjusted for the covariate. The **STDERR** option in the **LSMEANS** statement causes the standard error of the LS-means and the probability of getting a larger t value under the hypothesis $H_0: \text{LS-mean} = 0$ to be included in this table as well. Specifying the **PDIF** option causes all probability values for the hypothesis $H_0: \text{LS-mean}(i) = \text{LS-mean}(j)$ to be displayed, where the indexes i and j are numbered treatment levels.

Output 47.4.4 LS-Means

The GLM Procedure Least Squares Means

Drug	PostTreatment LSMEAN	Standard Error	Pr > t	LSMEAN Number
A	6.7149635	1.2884943	<.0001	1
D	6.8239348	1.2724690	<.0001	2
F	10.1611017	1.3159234	<.0001	3

Least Squares Means for effect Drug Pr > |t| for $H_0: \text{LSMean}(i)=\text{LSMean}(j)$

Dependent Variable: PostTreatment

i/j	1	2	3
1		0.9521	0.0793
2	0.9521		0.0835
3	0.0793	0.0835	

The **OUT=** and **COV** options in the **LSMEANS** statement create a data set of the estimates, their standard errors, and the variances and covariances of the LS-means, which is displayed in Output 47.4.5.

Output 47.4.5 LS-Means Output Data Set

Obs	_NAME_	Drug	LSMEAN	STDERR	NUMBER	COV1	COV2	COV3
1	PostTreatment A		6.7150	1.28849	1	1.66022	0.02844	-0.08403
2	PostTreatment D		6.8239	1.27247	2	0.02844	1.61918	-0.04299
3	PostTreatment F		10.1611	1.31592	3	-0.08403	-0.04299	1.73165

The new graphical features of PROC GLM enable you to visualize the fitted analysis of covariance model. The following statements enable ODS Graphics by specifying the ODS GRAPHICS statement and then fit an analysis-of-covariance model with LS-means for Drug.

```
ods graphics on;

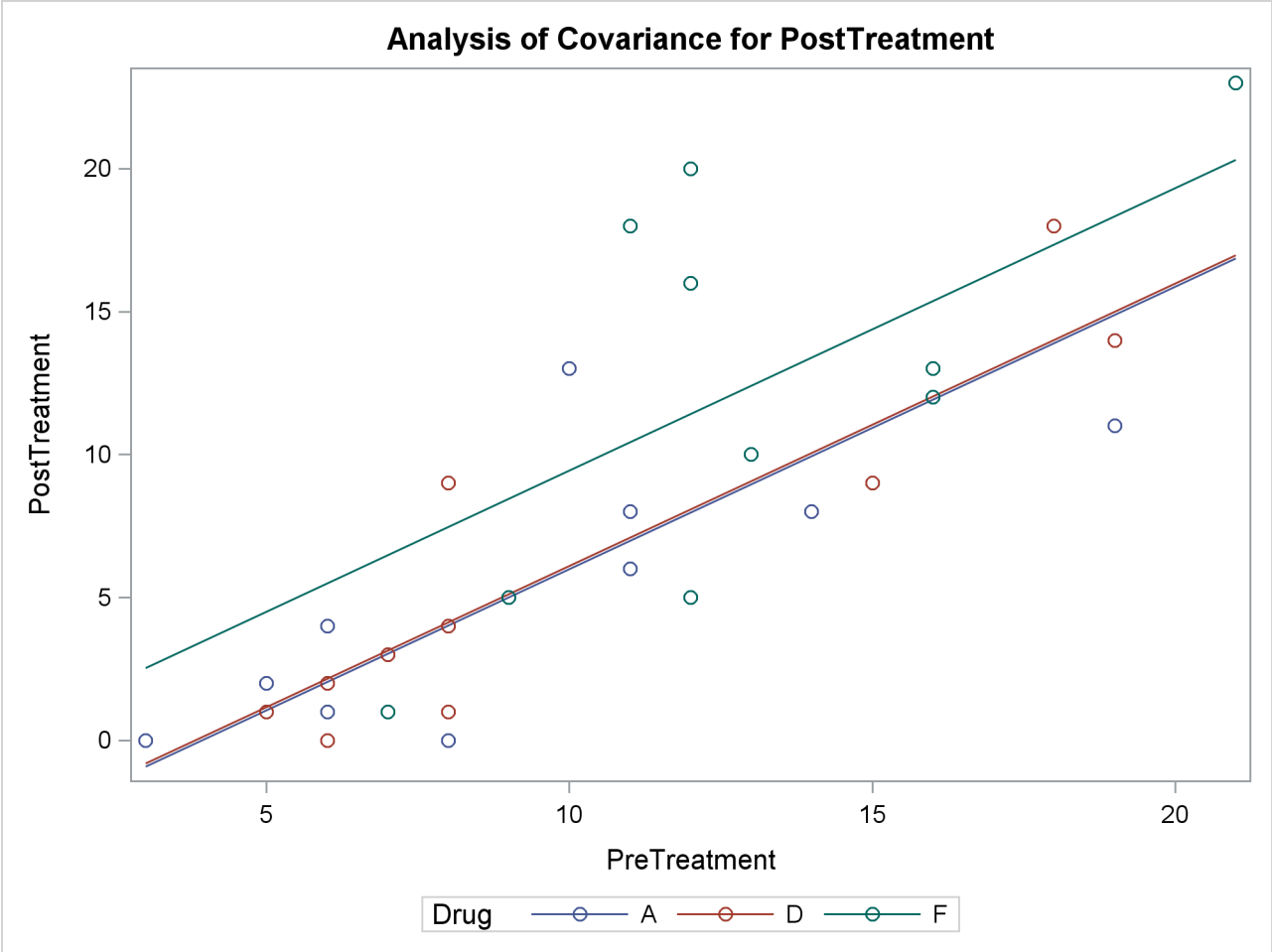
proc glm data=DrugTest plot=meanplot(c1);
  class Drug;
  model PostTreatment = Drug PreTreatment;
  lsmeans Drug / pdiff;
run;

ods graphics off;
```

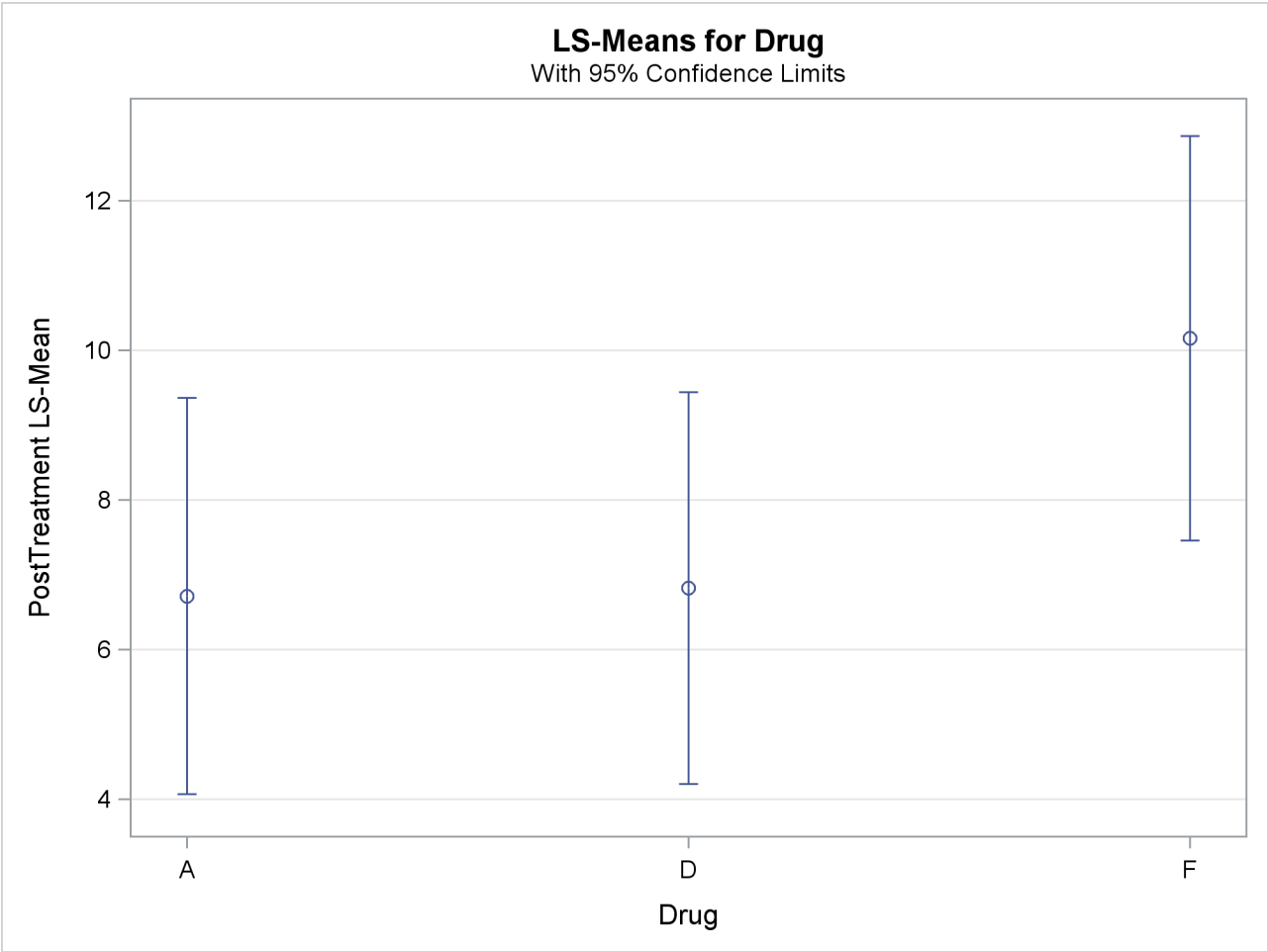
With graphics enabled, the GLM procedure output includes an analysis-of-covariance plot, as in Output 47.4.6. The **LSMEANS** statement produces a plot of the LS-means; the SAS statements previously shown use

the `PLOTS=MEANPLOT(CL)` option to add confidence limits for the individual LS-means, shown in [Output 47.4.7](#). If you also specify the `PDIFF` option in the `LSMEANS` statement, the output also includes a plot appropriate for the type of LS-mean differences computed. In this case, the default is to compare all LS-means with each other pairwise, so the plot is a “diffogram” or “mean-mean scatter plot” (Hsu 1996), as in [Output 47.4.8](#). For general information about ODS Graphics, see Chapter 21, “Statistical Graphics Using ODS.” For specific information about the graphics available in the GLM procedure, see the section “ODS Graphics” on page 3732.

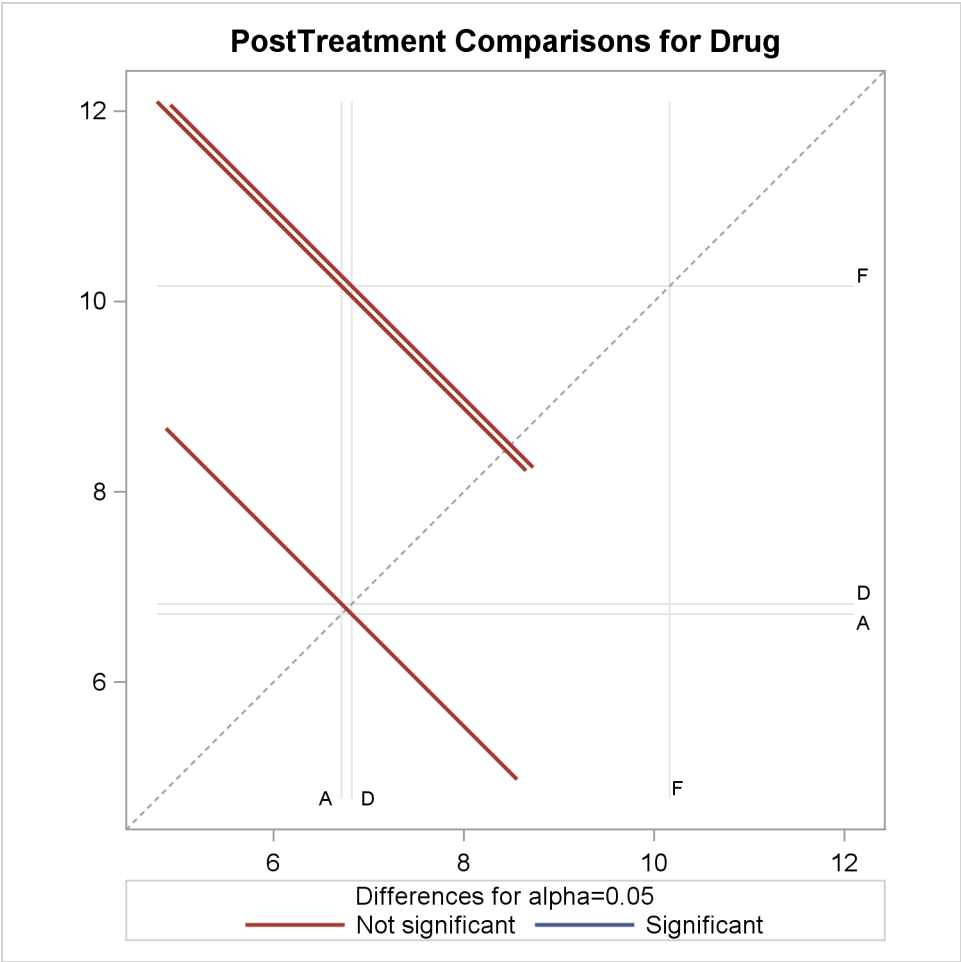
Output 47.4.6 Analysis of Covariance Plot of PostTreatment Score by Drug and PreTreatment Score



Output 47.4.7 LS-Means for PostTreatment Score by Drug



Output 47.4.8 Plot of Differences between Drug LS-Means for PostTreatment Scores



The analysis of covariance plot [Output 47.4.6](#) makes it clear that the control (drug F) has higher posttreatment scores across the range of pretreatment scores, while the fitted models for the two antibiotics (drugs A and D) nearly coincide. Similarly, while the diffogram [Output 47.4.8](#) indicates that none of the LS-mean differences are significant at the 5% level, the difference between the LS-means for the two antibiotics is much closer to zero than the differences between either one and the control.

Example 47.5: Three-Way Analysis of Variance with Contrasts

This example uses data from Cochran and Cox (1957, p. 176) to illustrate the analysis of a three-way factorial design with replication, including the use of the **CONTRAST** statement with interactions, the **OUTSTAT=** data set, and the **SLICE=** option in the **LSMEANS** statement.

The object of the study is to determine the effects of electric current on denervated muscle. The variables are as follows:

- Rep the replicate number, 1 or 2
- Time the length of time the current is applied to the muscle, ranging from 1 to 4

Current the level of electric current applied, ranging from 1 to 4
 Number the number of treatments per day, ranging from 1 to 3
 MuscleWeight the weight of the denervated muscle

The following statements produce [Output 47.5.1](#) through [Output 47.5.4](#).

```
data muscles;
  do Rep=1 to 2;
    do Time=1 to 4;
      do Current=1 to 4;
        do Number=1 to 3;
          input MuscleWeight @@;
          output;
        end;
      end;
    end;
  end;
datalines;
72 74 69 61 61 65 62 65 70 85 76 61
67 52 62 60 55 59 64 65 64 67 72 60
57 66 72 72 43 43 63 66 72 56 75 92
57 56 78 60 63 58 61 79 68 73 86 71
46 74 58 60 64 52 71 64 71 53 65 66
44 58 54 57 55 51 62 61 79 60 78 82
53 50 61 56 57 56 56 56 71 56 58 69
46 55 64 56 55 57 64 66 62 59 58 88
;

proc glm outstat=summary;
  class Rep Current Time Number;
  model MuscleWeight = Rep Current|Time|Number;
  contrast 'Time in Current 3'
    Time 1 0 0 -1 Current*Time 0 0 0 0 0 0 0 0 0 1 0 0 -1,
    Time 0 1 0 -1 Current*Time 0 0 0 0 0 0 0 0 0 0 1 0 -1,
    Time 0 0 1 -1 Current*Time 0 0 0 0 0 0 0 0 0 0 0 1 -1;
  contrast 'Current 1 versus 2' Current 1 -1;
  lsmeans Current*Time / slice=Current;
run;

proc print data=summary;
run;
```

The first **CONTRAST** statement examines the effects of Time within level 3 of Current. This is also called the *simple effect* of Time within Current*Time. Note that, since there are three degrees of freedom, it is necessary to specify three rows in the **CONTRAST** statement, separated by commas. Since the parameterization that PROC GLM uses is determined in part by the ordering of the variables in the **CLASS** statement, Current is specified before Time so that the Time parameters are nested within the Current*Time parameters; thus, the Current*Time contrast coefficients in each row are simply the Time coefficients of that row within the appropriate level of Current.

The second **CONTRAST** statement isolates a single-degree-of-freedom effect corresponding to the difference between the first two levels of **Current**. You can use such a contrast in a large experiment where certain preplanned comparisons are important, but you want to take advantage of the additional error degrees of freedom available when all levels of the factors are considered.

The **LSMEANS** statement with the **SLICE=** option is an alternative way to test for the simple effect of **Time** within **Current*Time**. In addition to listing the LS-means for each current strength and length of time, it gives a table of *F* tests for differences between the LS-means across **Time** within each **Current** level. In some cases, this can be a way to disentangle a complex interaction.

Output 47.5.1 Overall Analysis

The GLM Procedure

Class Level Information		
Class	Levels	Values
Rep	2	1 2
Current	4	1 2 3 4
Time	4	1 2 3 4
Number	3	1 2 3

Number of Observations Read 96

Number of Observations Used 96

The GLM Procedure

Dependent Variable: MuscleWeight

Source	DF	Sum of		F Value	Pr > F
		Squares	Mean Square		
Model	48	5782.916667	120.477431	1.77	0.0261
Error	47	3199.489583	68.074246		
Corrected Total	95	8982.406250			

R-Square	Coeff Var	Root MSE	MuscleWeight Mean
0.643805	13.05105	8.250712	63.21875

The output, shown in [Output 47.5.2](#) and [Output 47.5.3](#), indicates that the main effects for **Rep**, **Current**, and **Number** are significant (with *p*-values of 0.0045, <0.0001, and 0.0461, respectively), but the main effect for **Time** is not significant, indicating that, in general, it does not matter how long the current is applied. None of the interaction terms are significant, nor are the contrasts significant. Notice that the row in the sliced ANOVA table corresponding to level 3 of current matches the “Time in Current 3” contrast.

Output 47.5.2 Individual Effects and Contrasts

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Rep	1	605.010417	605.010417	8.89	0.0045
Current	3	2145.447917	715.149306	10.51	<.0001
Time	3	223.114583	74.371528	1.09	0.3616
Current*Time	9	298.677083	33.186343	0.49	0.8756
Number	2	447.437500	223.718750	3.29	0.0461
Current*Number	6	644.395833	107.399306	1.58	0.1747
Time*Number	6	367.979167	61.329861	0.90	0.5023
Current*Time*Number	18	1050.854167	58.380787	0.86	0.6276

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Rep	1	605.010417	605.010417	8.89	0.0045
Current	3	2145.447917	715.149306	10.51	<.0001
Time	3	223.114583	74.371528	1.09	0.3616
Current*Time	9	298.677083	33.186343	0.49	0.8756
Number	2	447.437500	223.718750	3.29	0.0461
Current*Number	6	644.395833	107.399306	1.58	0.1747
Time*Number	6	367.979167	61.329861	0.90	0.5023
Current*Time*Number	18	1050.854167	58.380787	0.86	0.6276

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Time in Current 3	3	34.83333333	11.61111111	0.17	0.9157
Current 1 versus 2	1	99.18750000	99.18750000	1.46	0.2334

Output 47.5.3 Simple Effects of Time

The GLM Procedure
Least Squares Means

Current*Time Effect Sliced by Current for MuscleWeight					
Current	DF	Sum of Squares	Mean Square	F Value	Pr > F
1	3	271.458333	90.486111	1.33	0.2761
2	3	120.666667	40.222222	0.59	0.6241
3	3	34.833333	11.611111	0.17	0.9157
4	3	94.833333	31.611111	0.46	0.7085

The SS, F statistics, and p -values can be stored in an `OUTSTAT=` data set, as shown in [Output 47.5.4](#).

Output 47.5.4 Contents of the OUTSTAT= Data Set

Obs	_NAME_	_SOURCE_	_TYPE_	DF	SS	F	PROB
1	MuscleWeight	ERROR	ERROR	47	3199.49	.	.
2	MuscleWeight	Rep	SS1	1	605.01	8.8875	0.00454
3	MuscleWeight	Current	SS1	3	2145.45	10.5054	0.00002
4	MuscleWeight	Time	SS1	3	223.11	1.0925	0.36159
5	MuscleWeight	Current*Time	SS1	9	298.68	0.4875	0.87562
6	MuscleWeight	Number	SS1	2	447.44	3.2864	0.04614
7	MuscleWeight	Current*Number	SS1	6	644.40	1.5777	0.17468
8	MuscleWeight	Time*Number	SS1	6	367.98	0.9009	0.50231
9	MuscleWeight	Current*Time*Number	SS1	18	1050.85	0.8576	0.62757
10	MuscleWeight	Rep	SS3	1	605.01	8.8875	0.00454
11	MuscleWeight	Current	SS3	3	2145.45	10.5054	0.00002
12	MuscleWeight	Time	SS3	3	223.11	1.0925	0.36159
13	MuscleWeight	Current*Time	SS3	9	298.68	0.4875	0.87562
14	MuscleWeight	Number	SS3	2	447.44	3.2864	0.04614
15	MuscleWeight	Current*Number	SS3	6	644.40	1.5777	0.17468
16	MuscleWeight	Time*Number	SS3	6	367.98	0.9009	0.50231
17	MuscleWeight	Current*Time*Number	SS3	18	1050.85	0.8576	0.62757
18	MuscleWeight	Time in Current 3	CONTRAST	3	34.83	0.1706	0.91574
19	MuscleWeight	Current 1 versus 2	CONTRAST	1	99.19	1.4570	0.23344

Example 47.6: Multivariate Analysis of Variance

This example employs multivariate analysis of variance (MANOVA) to measure differences in the chemical characteristics of ancient pottery found at four kiln sites in Great Britain. The data are from Tubb, Parker, and Nickless (1980), as reported in Hand et al. (1994).

For each of 26 samples of pottery, the percentages of oxides of five metals are measured. The following statements create the data set and invoke the GLM procedure to perform a one-way MANOVA. Additionally, it is of interest to know whether the pottery from one site in Wales (Llanederyn) differs from the samples from other sites; a `CONTRAST` statement is used to test this hypothesis.

```

title "Romano-British Pottery";
data pottery;
  input Site $12. Al Fe Mg Ca Na;
  datalines;
Llanederyn 14.4 7.00 4.30 0.15 0.51
Llanederyn 13.8 7.08 3.43 0.12 0.17
Llanederyn 14.6 7.09 3.88 0.13 0.20
Llanederyn 11.5 6.37 5.64 0.16 0.14
Llanederyn 13.8 7.06 5.34 0.20 0.20
Llanederyn 10.9 6.26 3.47 0.17 0.22

```

```

Llanederyn    10.1  4.26  4.26  0.20  0.18
Llanederyn    11.6  5.78  5.91  0.18  0.16
Llanederyn    11.1  5.49  4.52  0.29  0.30
Llanederyn    13.4  6.92  7.23  0.28  0.20
Llanederyn    12.4  6.13  5.69  0.22  0.54
Llanederyn    13.1  6.64  5.51  0.31  0.24
Llanederyn    12.7  6.69  4.45  0.20  0.22
Llanederyn    12.5  6.44  3.94  0.22  0.23
Caldicot      11.8  5.44  3.94  0.30  0.04
Caldicot      11.6  5.39  3.77  0.29  0.06
IslandThorns  18.3  1.28  0.67  0.03  0.03
IslandThorns  15.8  2.39  0.63  0.01  0.04
IslandThorns  18.0  1.50  0.67  0.01  0.06
IslandThorns  18.0  1.88  0.68  0.01  0.04
IslandThorns  20.8  1.51  0.72  0.07  0.10
AshleyRails   17.7  1.12  0.56  0.06  0.06
AshleyRails   18.3  1.14  0.67  0.06  0.05
AshleyRails   16.7  0.92  0.53  0.01  0.05
AshleyRails   14.8  2.74  0.67  0.03  0.05
AshleyRails   19.1  1.64  0.60  0.10  0.03
;

proc glm data=pottery;
  class Site;
  model Al Fe Mg Ca Na = Site;
  contrast 'Llanederyn vs. the rest' Site 1 1 1 -3;
  manova h=_all_ / printe printh;
run;

```

After the summary information, displayed in [Output 47.6.1](#), PROC GLM produces the univariate analyses for each of the dependent variables, as shown in [Output 47.6.2](#) through [Output 47.6.6](#). These analyses show that sites are significantly different for all oxides individually. You can suppress these univariate analyses by specifying the **NOUNI** option in the **MODEL** statement.

Output 47.6.1 Summary Information about Groups

Romano-British Pottery

The GLM Procedure

Class Level Information	
Class Levels	Values
Site	4 AshleyRails Caldicot IslandThorns Llanederyn

Number of Observations Read	26
Number of Observations Used	26

Output 47.6.2 Univariate Analysis of Variance for Aluminum Oxide**Romano-British Pottery****The GLM Procedure****Dependent Variable: Al**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	175.6103187	58.5367729	26.67	<.0001
Error	22	48.2881429	2.1949156		
Corrected Total	25	223.8984615			

R-Square	Coeff Var	Root MSE	Al Mean
0.784330	10.22284	1.481525	14.49231

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Site	3	175.6103187	58.5367729	26.67	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Site	3	175.6103187	58.5367729	26.67	<.0001

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Llanederyn vs. the rest	1	58.58336640	58.58336640	26.69	<.0001

Output 47.6.3 Univariate Analysis of Variance for Iron Oxide**Romano-British Pottery****The GLM Procedure****Dependent Variable: Fe**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	134.2216158	44.7405386	89.88	<.0001
Error	22	10.9508457	0.4977657		
Corrected Total	25	145.1724615			

R-Square	Coeff Var	Root MSE	Fe Mean
0.924567	15.79171	0.705525	4.467692

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Site	3	134.2216158	44.7405386	89.88	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Site	3	134.2216158	44.7405386	89.88	<.0001

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Llanederyn vs. the rest	1	71.15144132	71.15144132	142.94	<.0001

Output 47.6.4 Univariate Analysis of Variance for Calcium Oxide**Romano-British Pottery****The GLM Procedure****Dependent Variable: Ca**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	0.20470275	0.06823425	29.16	<.0001
Error	22	0.05148571	0.00234026		
Corrected Total	25	0.25618846			

R-Square	Coeff Var	Root MSE	Ca Mean
0.799032	33.01265	0.048376	0.146538

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Site	3	0.20470275	0.06823425	29.16	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Site	3	0.20470275	0.06823425	29.16	<.0001

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Llanederyn vs. the rest	1	0.03531688	0.03531688	15.09	0.0008

Output 47.6.5 Univariate Analysis of Variance for Magnesium Oxide**Romano-British Pottery****The GLM Procedure****Dependent Variable: Mg**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	103.3505270	34.4501757	49.12	<.0001
Error	22	15.4296114	0.7013460		
Corrected Total	25	118.7801385			

R-Square	Coeff Var	Root MSE	Mg Mean
0.870099	26.65777	0.837464	3.141538

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Site	3	103.3505270	34.4501757	49.12	<.0001

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Site	3	103.3505270	34.4501757	49.12	<.0001

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Llanederyn vs. the rest	1	56.59349339	56.59349339	80.69	<.0001

Output 47.6.6 Univariate Analysis of Variance for Sodium Oxide**Romano-British Pottery****The GLM Procedure****Dependent Variable: Na**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	3	0.25824560	0.08608187	9.50	0.0003
Error	22	0.19929286	0.00905877		
Corrected Total	25	0.45753846			

R-Square	Coeff Var	Root MSE	Na Mean
0.564424	60.06350	0.095178	0.158462

Source	DF	Type I SS	Mean Square	F Value	Pr > F
Site	3	0.25824560	0.08608187	9.50	0.0003

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Site	3	0.25824560	0.08608187	9.50	0.0003

Contrast	DF	Contrast SS	Mean Square	F Value	Pr > F
Llanederyn vs. the rest	1	0.23344446	0.23344446	25.77	<.0001

The **PRINTE** option in the **MANOVA** statement displays the elements of the error matrix, also called the Error Sums of Squares and Crossproducts matrix. (See [Output 47.6.7](#).) The diagonal elements of this matrix are the error sums of squares from the corresponding univariate analyses.

The **PRINTE** option also displays the partial correlation matrix associated with the E matrix. In this example, none of the oxides are very strongly correlated; the strongest correlation ($r = 0.488$) is between magnesium oxide and calcium oxide.

Output 47.6.7 Error SSCP Matrix and Partial Correlations**Romano-British Pottery****The GLM Procedure**
Multivariate Analysis of Variance

E = Error SSCP Matrix					
	Al	Fe	Mg	Ca	Na
Al	48.288142857	7.0800714286	0.6080142857	0.1064714286	0.5889571429
Fe	7.0800714286	10.950845714	0.5270571429	-0.155194286	0.0667585714
Mg	0.6080142857	0.5270571429	15.429611429	0.4353771429	0.0276157143
Ca	0.1064714286	-0.155194286	0.4353771429	0.0514857143	0.0100785714
Na	0.5889571429	0.0667585714	0.0276157143	0.0100785714	0.1992928571

Output 47.6.7 *continued*

Partial Correlation Coefficients from the Error SSCP Matrix / Prob > r					
DF = 22	Al	Fe	Mg	Ca	Na
Al	1.000000	0.307889 0.1529	0.022275 0.9196	0.067526 0.7595	0.189853 0.3856
Fe	0.307889 0.1529	1.000000	0.040547 0.8543	-0.206685 0.3440	0.045189 0.8378
Mg	0.022275 0.9196	0.040547 0.8543	1.000000	0.488478 0.0180	0.015748 0.9431
Ca	0.067526 0.7595	-0.206685 0.3440	0.488478 0.0180	1.000000	0.099497 0.6515
Na	0.189853 0.3856	0.045189 0.8378	0.015748 0.9431	0.099497 0.6515	1.000000

The **PRINTH** option produces the SSCP matrix for the hypotheses being tested (Site and the contrast); see [Output 47.6.8](#) and [Output 47.6.9](#). Since the Type III SS are the highest-level SS produced by PROC GLM by default, and since the **HTYPE=** option is not specified, the SSCP matrix for Site gives the Type III **H** matrix. The diagonal elements of this matrix are the model sums of squares from the corresponding univariate analyses.

Four multivariate tests are computed, all based on the characteristic roots and vectors of $\mathbf{E}^{-1}\mathbf{H}$. These roots and vectors are displayed along with the tests. All four tests can be transformed to variates that have F distributions under the null hypothesis. Note that the four tests all give the same results for the contrast, since it has only one degree of freedom. In this case, the multivariate analysis matches the univariate results: there is an overall difference between the chemical composition of samples from different sites, and the samples from Llanederyn are different from the average of the other sites.

Output 47.6.8 Hypothesis SSCP Matrix and Multivariate Tests for Overall Site Effect**Romano-British Pottery****The GLM Procedure
Multivariate Analysis of Variance**

H = Type III SSCP Matrix for Site					
	Al	Fe	Mg	Ca	Na
Al	175.61031868	-149.295533	-130.8097066	-5.889163736	-5.372264835
Fe	-149.295533	134.22161582	117.74503516	4.8217865934	5.3259491209
Mg	-130.8097066	117.74503516	103.35052703	4.2091613187	4.7105458242
Ca	-5.889163736	4.8217865934	4.2091613187	0.2047027473	0.154782967
Na	-5.372264835	5.3259491209	4.7105458242	0.154782967	0.2582456044

Output 47.6.8 continued

Characteristic Roots and Vectors of: E Inverse * H, where H = Type III SSCP Matrix for Site E = Error SSCP Matrix Characteristic Vector V'EV=1						
Characteristic Root	Percent	Al	Fe	Mg	Ca	Na
34.1611140	96.39	0.09562211	-0.26330469	-0.05305978	-1.87982100	-0.47071123
1.2500994	3.53	0.02651891	-0.01239715	0.17564390	-4.25929785	1.23727668
0.0275396	0.08	0.09082220	0.13159869	0.03508901	-0.15701602	-1.39364544
0.0000000	0.00	0.03673984	-0.15129712	0.20455529	0.54624873	-0.17402107
0.0000000	0.00	0.06862324	0.03056912	-0.10662399	2.51151978	1.23668841

MANOVA Test Criteria and F Approximations for the Hypothesis of No Overall Site Effect H = Type III SSCP Matrix for Site E = Error SSCP Matrix					
S=3 M=0.5 N=8					
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.01230091	13.09	15	50.091	<.0001
Pillai's Trace	1.55393619	4.30	15	60	<.0001
Hotelling-Lawley Trace	35.43875302	40.59	15	29.13	<.0001
Roy's Greatest Root	34.16111399	136.64	5	20	<.0001

NOTE: F Statistic for Roy's Greatest Root is an upper bound.

Output 47.6.9 Hypothesis SSCP Matrix and Multivariate Tests for Differences between Llanederyn and the Other Sites

H = Contrast SSCP Matrix for Llanederyn vs. the rest					
	Al	Fe	Mg	Ca	Na
Al	58.583366402	-64.56230291	-57.57983466	-1.438395503	-3.698102513
Fe	-64.56230291	71.151441323	63.456352116	1.5851961376	4.0755256878
Mg	-57.57983466	63.456352116	56.593493386	1.4137558201	3.6347541005
Ca	-1.438395503	1.5851961376	1.4137558201	0.0353168783	0.0907993915
Na	-3.698102513	4.0755256878	3.6347541005	0.0907993915	0.2334444577

Characteristic Roots and Vectors of: E Inverse * H, where H = Contrast SSCP Matrix for Llanederyn vs. the rest E = Error SSCP Matrix Characteristic Vector V'EV=1						
Characteristic Root	Percent	Al	Fe	Mg	Ca	Na
16.1251646	100.00	-0.08883488	0.25458141	0.08723574	0.98158668	0.71925759
0.0000000	0.00	-0.00503538	0.03825743	-0.17632854	5.16256699	-0.01022754
0.0000000	0.00	0.00162771	-0.08885364	-0.01774069	-0.83096817	2.17644566
0.0000000	0.00	0.04450136	-0.15722494	0.22156791	0.00000000	0.00000000
0.0000000	0.00	0.11939206	0.10833549	0.00000000	0.00000000	0.00000000

Output 47.6.9 *continued*

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of No Overall Llanederyn vs. the rest Effect					
H = Contrast SSCP Matrix for Llanederyn vs. the rest					
E = Error SSCP Matrix					
		S=1	M=1.5	N=8	
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.05839360	58.05	5	18	<.0001
Pillai's Trace	0.94160640	58.05	5	18	<.0001
Hotelling-Lawley Trace	16.12516462	58.05	5	18	<.0001
Roy's Greatest Root	16.12516462	58.05	5	18	<.0001

Example 47.7: Repeated Measures Analysis of Variance

This example uses data from Cole and Grizzle (1966) to illustrate a commonly occurring repeated measures ANOVA design. Sixteen dogs are randomly assigned to four groups. (One animal is removed from the analysis due to a missing value for one dependent variable.) Dogs in each group receive either morphine or trimethaphan (variable Drug) and have either depleted or intact histamine levels (variable Depleted) before receiving the drugs. The dependent variable is the blood concentration of histamine at 0, 1, 3, and 5 minutes after injection of the drug. Logarithms are applied to these concentrations to minimize correlation between the mean and the variance of the data.

The following SAS statements perform both univariate and multivariate repeated measures analyses and produce [Output 47.7.1](#) through [Output 47.7.7](#).

```
data dogs;
  input Drug $12. Depleted $ Histamine0 Histamine1
        Histamine3 Histamine5;
  LogHistamine0=log(Histamine0);
  LogHistamine1=log(Histamine1);
  LogHistamine3=log(Histamine3);
  LogHistamine5=log(Histamine5);
  datalines;
Morphine      N   .04   .20   .10   .08
Morphine      N   .02   .06   .02   .02
Morphine      N   .07  1.40   .48   .24
Morphine      N   .17   .57   .35   .24
Morphine      Y   .10   .09   .13   .14
Morphine      Y   .12   .11   .10   .
Morphine      Y   .07   .07   .06   .07
Morphine      Y   .05   .07   .06   .07
Trimethaphan  N   .03   .62   .31   .22
Trimethaphan  N   .03  1.05   .73   .60
Trimethaphan  N   .07   .83  1.07   .80
Trimethaphan  N   .09  3.13  2.06  1.23
Trimethaphan  Y   .10   .09   .09   .08
Trimethaphan  Y   .08   .09   .09   .10
Trimethaphan  Y   .13   .10   .12   .12
Trimethaphan  Y   .06   .05   .05   .05
;
```

```

proc glm;
  class Drug Depleted;
  model LogHistamine0--LogHistamine5 =
    Drug Depleted Drug*Depleted / nouni;
  repeated Time 4 (0 1 3 5) polynomial / summary printe;
run;

```

The **NOUNI** option in the **MODEL** statement suppresses the individual ANOVA tables for the original dependent variables. These analyses are usually of no interest in a repeated measures analysis. The **POLYNOMIAL** option in the **REPEATED** statement indicates that the transformation used to implement the repeated measures analysis is an orthogonal polynomial transformation, and the **SUMMARY** option requests that the univariate analyses for the orthogonal polynomial contrast variables be displayed. The parenthetical numbers (0 1 3 5) determine the spacing of the orthogonal polynomials used in the analysis.

Output 47.7.1 Summary Information about Groups

The GLM Procedure

Class Level Information		
Class	Levels	Values
Drug	2	Morphine Trimethaphan
Depleted	2	N Y

Number of Observations Read	16
Number of Observations Used	15

The “Repeated Measures Level Information” table gives information about the repeated measures effect; it is displayed in [Output 47.7.2](#). In this example, the within-subject (within-dog) effect is Time, which has the levels 0, 1, 3, and 5.

Output 47.7.2 Repeated Measures Levels

The GLM Procedure Repeated Measures Analysis of Variance

Repeated Measures Level Information				
Dependent Variable	LogHistamine0	LogHistamine1	LogHistamine3	LogHistamine5
Level of Time	0	1	3	5

The multivariate analyses for within-subject effects and related interactions are displayed in [Output 47.7.3](#). For the example, the first table displayed shows that the **TIME** effect is significant. In addition, the **Time*Drug*Depleted** interaction is significant, as shown in the fourth table. This means that the effect of Time on the blood concentration of histamine is different for the four **Drug*Depleted** combinations studied.

Output 47.7.3 Multivariate Tests of Within-Subject Effects

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Time Effect
H = Type III SSCP Matrix for Time
E = Error SSCP Matrix

	S=1	M=0.5	N=3.5		
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.11097706	24.03	3	9	0.0001
Pillai's Trace	0.88902294	24.03	3	9	0.0001
Hotelling-Lawley Trace	8.01087137	24.03	3	9	0.0001
Roy's Greatest Root	8.01087137	24.03	3	9	0.0001

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Time*Drug Effect
H = Type III SSCP Matrix for Time*Drug
E = Error SSCP Matrix

	S=1	M=0.5	N=3.5		
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.34155984	5.78	3	9	0.0175
Pillai's Trace	0.65844016	5.78	3	9	0.0175
Hotelling-Lawley Trace	1.92774470	5.78	3	9	0.0175
Roy's Greatest Root	1.92774470	5.78	3	9	0.0175

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Time*Depleted Effect
H = Type III SSCP Matrix for Time*Depleted
E = Error SSCP Matrix

	S=1	M=0.5	N=3.5		
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.12339988	21.31	3	9	0.0002
Pillai's Trace	0.87660012	21.31	3	9	0.0002
Hotelling-Lawley Trace	7.10373567	21.31	3	9	0.0002
Roy's Greatest Root	7.10373567	21.31	3	9	0.0002

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Time*Drug*Depleted Effect
H = Type III SSCP Matrix for Time*Drug*Depleted
E = Error SSCP Matrix

	S=1	M=0.5	N=3.5		
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.19383010	12.48	3	9	0.0015
Pillai's Trace	0.80616990	12.48	3	9	0.0015
Hotelling-Lawley Trace	4.15915732	12.48	3	9	0.0015
Roy's Greatest Root	4.15915732	12.48	3	9	0.0015

Output 47.7.4 displays tests of hypotheses for between-subject (between-dog) effects. This section tests the hypotheses that the different Drugs, Depleteds, and their interactions have no effects on the dependent variables, while ignoring the within-dog effects. From this analysis, there is a significant between-dog effect for Depleted (p -value=0.0229). The interaction and the main effect for Drug are not significant (p -values=0.1734 and 0.1281, respectively).

Output 47.7.4 Tests of Between-Subject Effects

The GLM Procedure
Repeated Measures Analysis of Variance
Tests of Hypotheses for Between Subjects Effects

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Drug	1	5.99336243	5.99336243	2.71	0.1281
Depleted	1	15.44840703	15.44840703	6.98	0.0229
Drug*Depleted	1	4.69087508	4.69087508	2.12	0.1734
Error	11	24.34683348	2.21334850		

Univariate analyses for within-subject (within-dog) effects and related interactions are displayed in [Output 47.7.6](#). The results for this example are the same as for the multivariate analyses; this is not always the case. In addition, before the univariate analyses are used to make conclusions about the data, the result of the sphericity test (requested with the [PRINTE](#) option in the [REPEATED](#) statement and displayed in [Output 47.7.5](#)) should be examined. If the sphericity test is rejected, consider using the adjusted G-G or H-F-L probabilities. See the section “[Repeated Measures Analysis of Variance](#)” on page 3712 for more information.

Output 47.7.5 Sphericity Test

Sphericity Tests				
Variables	DF	Mauchly's		
		Criterion	Chi-Square	Pr > ChiSq
Transformed Variates	5	0.1752641	16.930873	0.0046
Orthogonal Components	5	0.1752641	16.930873	0.0046

Output 47.7.6 Univariate Tests of Within-Subject Effects

The GLM Procedure
Repeated Measures Analysis of Variance
Univariate Tests of Hypotheses for Within Subject Effects

Source	DF	Type III SS	Mean Square	F Value	Pr > F	Adj Pr > F	
						G - G	H-F-L
Time	3	12.05898677	4.01966226	53.44	<.0001	<.0001	<.0001
Time*Drug	3	1.84429514	0.61476505	8.17	0.0003	0.0039	0.0023
Time*Depleted	3	12.08978557	4.02992852	53.57	<.0001	<.0001	<.0001
Time*Drug*Depleted	3	2.93077939	0.97692646	12.99	<.0001	0.0005	0.0002
Error(Time)	33	2.48238887	0.07522391				

Greenhouse-Geisser Epsilon	0.5694
Huynh-Feldt-Lecoutre Epsilon	0.6636

[Output 47.7.7](#) is produced by the [SUMMARY](#) option in the [REPEATED](#) statement. If the [POLYNOMIAL](#) option is not used, a similar table is displayed using the default [CONTRAST](#) transformation. The linear, quadratic, and cubic trends for Time, labeled as ‘Time_1’, ‘Time_2’, and ‘Time_3’, are displayed, and in each case, the Source labeled ‘Mean’ gives a test for the respective trend.

Output 47.7.7 Tests of Between-Subject Effects for Transformed Variables

The GLM Procedure
Repeated Measures Analysis of Variance
Analysis of Variance of Contrast Variables

Time_N represents the nth degree polynomial contrast for Time

Contrast Variable: Time_1

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Mean	1	2.00963483	2.00963483	34.99	0.0001
Drug	1	1.18069076	1.18069076	20.56	0.0009
Depleted	1	1.36172504	1.36172504	23.71	0.0005
Drug*Depleted	1	2.04346848	2.04346848	35.58	<.0001
Error	11	0.63171161	0.05742833		

Contrast Variable: Time_2

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Mean	1	5.40988418	5.40988418	57.15	<.0001
Drug	1	0.59173192	0.59173192	6.25	0.0295
Depleted	1	5.94945506	5.94945506	62.86	<.0001
Drug*Depleted	1	0.67031587	0.67031587	7.08	0.0221
Error	11	1.04118707	0.09465337		

Contrast Variable: Time_3

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Mean	1	4.63946776	4.63946776	63.04	<.0001
Drug	1	0.07187246	0.07187246	0.98	0.3443
Depleted	1	4.77860547	4.77860547	64.94	<.0001
Drug*Depleted	1	0.21699504	0.21699504	2.95	0.1139
Error	11	0.80949018	0.07359002		

Example 47.8: Mixed Model Analysis of Variance with the RANDOM Statement

Milliken and Johnson (1984) present an example of an unbalanced mixed model. Three machines, which are considered as a fixed effect, and six employees, which are considered a random effect, are studied. Each employee operates each machine for either one, two, or three different times. The dependent variable is an overall rating, which takes into account the number and quality of components produced.

The following statements form the data set and perform a mixed model analysis of variance by requesting the **TEST** option in the **RANDOM** statement. Note that the machine*person interaction is declared as a random effect; in general, when an interaction involves a random effect, it too should be declared as random. The results of the analysis are shown in [Output 47.8.1](#) through [Output 47.8.4](#).

```
data machine;
  input machine person rating @@;
  datalines;
1 1 52.0 1 2 51.8 1 2 52.8 1 3 60.0 1 4 51.1 1 4 52.3 1 5 50.9
1 5 51.8 1 5 51.4 1 6 46.4 1 6 44.8 1 6 49.2 2 1 64.0 2 2 59.7
```

```

2 2 60.0  2 2 59.0  2 3 68.6  2 3 65.8  2 4 63.2  2 4 62.8  2 4 62.2
2 5 64.8  2 5 65.0  2 6 43.7  2 6 44.2  2 6 43.0  3 1 67.5  3 1 67.2
3 1 66.9  3 2 61.5  3 2 61.7  3 2 62.3  3 3 70.8  3 3 70.6  3 3 71.0
3 4 64.1  3 4 66.2  3 4 64.0  3 5 72.1  3 5 72.0  3 5 71.1  3 6 62.0
3 6 61.4  3 6 60.5
;

proc glm data=machine;
  class machine person;
  model rating=machine person machine*person;
  random person machine*person / test;
run;

```

The **TEST** option in the **RANDOM** statement requests that PROC GLM determine the appropriate F tests based on person and machine*person being treated as random effects. As you can see in [Output 47.8.4](#), this requires that a linear combination of mean squares be constructed to test both the machine and person hypotheses; thus, F tests that use Satterthwaite approximations are needed.

Output 47.8.1 Summary Information about Groups

The GLM Procedure

Class Level Information

Class	Levels	Values
machine	3	1 2 3
person	6	1 2 3 4 5 6

Number of Observations Read 44

Number of Observations Used 44

Output 47.8.2 Fixed-Effect Model Analysis of Variance

The GLM Procedure

Dependent Variable: rating

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	17	3061.743333	180.102549	206.41	<.0001
Error	26	22.686667	0.872564		
Corrected Total	43	3084.430000			

R-Square	Coeff Var	Root MSE	rating Mean
0.992645	1.560754	0.934111	59.85000

Source	DF	Type I SS	Mean Square	F Value	Pr > F
machine	2	1648.664722	824.332361	944.72	<.0001
person	5	1008.763583	201.752717	231.22	<.0001
machine*person	10	404.315028	40.431503	46.34	<.0001

Output 47.8.2 *continued*

Source	DF	Type III SS	Mean Square	F Value	Pr > F
machine	2	1238.197626	619.098813	709.52	<.0001
person	5	1011.053834	202.210767	231.74	<.0001
machine*person	10	404.315028	40.431503	46.34	<.0001

Output 47.8.3 Expected Values of Type III Mean Squares

Source	Type III Expected Mean Square
machine	$\text{Var}(\text{Error}) + 2.137 \text{ Var}(\text{machine*person}) + \text{Q}(\text{machine})$
person	$\text{Var}(\text{Error}) + 2.2408 \text{ Var}(\text{machine*person}) + 6.7224 \text{ Var}(\text{person})$
machine*person	$\text{Var}(\text{Error}) + 2.3162 \text{ Var}(\text{machine*person})$

Output 47.8.4 Mixed Model Analysis of Variance

The GLM Procedure
Tests of Hypotheses for Mixed Model Analysis of Variance

Dependent Variable: rating

Source	DF	Type III SS	Mean Square	F Value	Pr > F
machine	2	1238.197626	619.098813	16.57	0.0007
Error	10.036	375.057436	37.370384		
Error: 0.9226*MS(machine*person) + 0.0774*MS(Error)					
Source	DF	Type III SS	Mean Square	F Value	Pr > F
person	5	1011.053834	202.210767	5.17	0.0133
Error	10.015	392.005726	39.143708		
Error: 0.9674*MS(machine*person) + 0.0326*MS(Error)					
Source	DF	Type III SS	Mean Square	F Value	Pr > F
machine*person	10	404.315028	40.431503	46.34	<.0001
Error: MS(Error)	26	22.686667	0.872564		

Note that you can also use the MIXED procedure to analyze mixed models. The following statements use PROC MIXED to reproduce the mixed model analysis of variance; the relevant part of the PROC MIXED results is shown in [Output 47.8.5](#).

```
proc mixed data=machine method=type3;
  class machine person;
  model rating = machine;
  random person machine*person;
run;
```

Output 47.8.5 PROC MIXED Mixed Model Analysis of Variance (Partial Output)

The Mixed Procedure

Type 3 Analysis of Variance						
Source	DF	Sum of Squares	Mean Square	Expected Mean Square	Error Term	Error DF
machine	2	1238.197626	619.098813	Var(Residual) + 2.137 Var(machine*person) + Q(machine)	0.9226 MS(machine*person) + 0.0774 MS(Residual)	10.036
person	5	1011.053834	202.210767	Var(Residual) + 2.2408 Var(machine*person) + 6.7224 Var(person)	0.9674 MS(machine*person) + 0.0326 MS(Residual)	10.015
machine*person	10	404.315028	40.431503	Var(Residual) + 2.3162 Var(machine*person)	MS(Residual)	26
Residual	26	22.686667	0.872564	Var(Residual)	.	

Type 3 Analysis of Variance		
Source	F Value	Pr > F
machine	16.57	0.0007
person	5.17	0.0133
machine*person	46.34	<.0001
Residual		

The advantage of PROC MIXED is that it offers more versatility for mixed models; the disadvantage is that it can be less computationally efficient for large data sets. See Chapter 78, “[The MIXED Procedure](#),” for more details.

Example 47.9: Analyzing a Doubly Multivariate Repeated Measures Design

This example shows how to analyze a doubly multivariate repeated measures design by using PROC GLM with an IDENTITY factor in the REPEATED statement. Note that this differs from previous releases of PROC GLM, in which you had to use a MANOVA statement to get a doubly repeated measures analysis.

Two responses, Y1 and Y2, are each measured three times for each subject (pretreatment, posttreatment, and in a later follow-up). Each subject receives one of three treatments; A, B, or the control. In PROC GLM, you use a **REPEATED** factor of type **IDENTITY** to identify the different responses and another repeated factor to identify the different measurement times. The repeated measures analysis includes multivariate tests for time and treatment main effects, as well as their interactions, across responses. The following statements produce **Output 47.9.1** through **Output 47.9.3**.

[illegible]

```

A      1  3  13  9  0  0  9
A      2  0  14 10  6  6  3
A      3  4   6 17  8  2  6
A      4  7   7 13  7  6  4
A      5  3  12 11  6 12  6
A      6 10  14  8 13  3  8
B      1  9  11 17  8 11 27
B      2  4  16 13  9  3 26
B      3  8  10  9 12  0 18
B      4  5   9 13  3  0 14
B      5  0  15 11  3  0 25
B      6  4  11 14  4  2  9
Control 1 10  12 15  4  3  7
Control 2  2   8 12  8  7 20
Control 3  4   9 10  2  0 10
Control 4 10   8  8  5  8 14
Control 5 11  11 11  1  0 11
Control 6  1  5  15  8  9 10
;

proc glm data=Trials;
  class Treatment;
  model PreY1 PostY1 FollowY1
        PreY2 PostY2 FollowY2 = Treatment / nouni;
  repeated Response 2 identity, Time 3;
run;

```

Output 47.9.1 A Doubly Multivariate Repeated Measures Design**The GLM Procedure**

Class Level Information		
Class	Levels	Values
Treatment	3	A B Control

Number of Observations Read	18
Number of Observations Used	18

The levels of the repeated factors are displayed in [Output 47.9.2](#). Note that RESPONSE is 1 for all the Y1 measurements and 2 for all the Y2 measurements, while the three levels of Time identify the pretreatment, posttreatment, and follow-up measurements within each response. The multivariate tests for within-subject effects are displayed in [Output 47.9.3](#).

Output 47.9.2 Repeated Factor Levels**The GLM Procedure
Repeated Measures Analysis of Variance**

Repeated Measures Level Information						
Dependent Variable	PreY1	PostY1	FollowY1	PreY2	PostY2	FollowY2
Level of Response	1	1	1	2	2	2
Level of Time	1	2	3	1	2	3

Output 47.9.3 Within-Subject Tests

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Response Effect
 H = Type III SSCP Matrix for Response
 E = Error SSCP Matrix

S=1 M=0 N=6					
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.02165587	316.24	2	14	<.0001
Pillai's Trace	0.97834413	316.24	2	14	<.0001
Hotelling-Lawley Trace	45.17686368	316.24	2	14	<.0001
Roy's Greatest Root	45.17686368	316.24	2	14	<.0001

MANOVA Test Criteria and F Approximations for the Hypothesis of no Response*Treatment Effect
 H = Type III SSCP Matrix for Response*Treatment
 E = Error SSCP Matrix

S=2 M=-0.5 N=6					
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.72215797	1.24	4	28	0.3178
Pillai's Trace	0.27937444	1.22	4	30	0.3240
Hotelling-Lawley Trace	0.38261660	1.31	4	15.818	0.3074
Roy's Greatest Root	0.37698780	2.83	2	15	0.0908

NOTE: F Statistic for Roy's Greatest Root is an upper bound.

NOTE: F Statistic for Wilks' Lambda is exact.

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of no Response*Time Effect
 H = Type III SSCP Matrix for Response*Time
 E = Error SSCP Matrix

S=1 M=1 N=5					
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.14071380	18.32	4	12	<.0001
Pillai's Trace	0.85928620	18.32	4	12	<.0001
Hotelling-Lawley Trace	6.10662362	18.32	4	12	<.0001
Roy's Greatest Root	6.10662362	18.32	4	12	<.0001

MANOVA Test Criteria and F Approximations for the Hypothesis of no Response*Time*Treatment Effect
 H = Type III SSCP Matrix for Response*Time*Treatment
 E = Error SSCP Matrix

S=2 M=0.5 N=5					
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.22861451	3.27	8	24	0.0115
Pillai's Trace	0.96538785	3.03	8	26	0.0151
Hotelling-Lawley Trace	2.52557514	3.64	8	15	0.0149
Roy's Greatest Root	2.12651905	6.91	4	13	0.0033

NOTE: F Statistic for Roy's Greatest Root is an upper bound.

NOTE: F Statistic for Wilks' Lambda is exact.

The table for Response*Treatment tests for an overall treatment effect across the two responses; likewise, the tables for Response*Time and Response*Treatment*Time test for time and the treatment-by-time interaction, respectively. In this case, there is a strong main effect for time and possibly for the interaction, but not for treatment.

In previous releases (before the IDENTITY transformation was introduced), in order to perform a doubly repeated measures analysis, you had to use a **MANOVA** statement with a customized transformation matrix M. You might still want to use this approach to see details of the analysis, such as the univariate ANOVA for each transformed variate. The following statements demonstrate this approach by using the **MANOVA** statement to test for the overall main effect of time and specifying the **SUMMARY** option.

```
proc glm data=Trial;
  class Treatment;
  model PreY1 PostY1 FollowY1
        PreY2 PostY2 FollowY2 = Treatment / noint;
  manova h=intercept m=prey1 - posty1,
        prey1 - followy1,
        prey2 - posty2,
        prey2 - followy2 / summary;
run;
```

The M matrix used to perform the test for time effects is displayed in [Output 47.9.4](#), while the results of the multivariate test are given in [Output 47.9.5](#). Note that the test results are the same as for the Response*Time effect in [Output 47.9.3](#).

Output 47.9.4 M Matrix to Test for Time Effect (Repeated Measure)

The GLM Procedure
Multivariate Analysis of Variance

M Matrix Describing Transformed Variables						
	PreY1	PostY1	FollowY1	PreY2	PostY2	FollowY2
MVAR1	1	-1	0	0	0	0
MVAR2	1	0	-1	0	0	0
MVAR3	0	0	0	1	-1	0
MVAR4	0	0	0	1	0	-1

Output 47.9.5 Tests for Time Effect (Repeated Measure)

The GLM Procedure
Multivariate Analysis of Variance

Characteristic Roots and Vectors of: E Inverse * H, where H = Type III SSCP Matrix for Intercept E = Error SSCP Matrix						
Variables have been transformed by the M Matrix						
Characteristic Vector V'EV=1						
Characteristic Root	Percent	MVAR1	MVAR2	MVAR3	MVAR4	
6.10662362	100.00	-0.00157729	0.04081620	-0.04210209	0.03519437	
0.00000000	0.00	0.00796367	0.00493217	0.05185236	0.00377940	
0.00000000	0.00	-0.03534089	-0.01502146	-0.00283074	0.04259372	
0.00000000	0.00	-0.05672137	0.04500208	0.00000000	0.00000000	

Output 47.9.5 *continued*

MANOVA Test Criteria and Exact F Statistics for the Hypothesis of No Overall Intercept Effect on the Variables Defined by the M Matrix Transformation H = Type III SSCP Matrix for Intercept E = Error SSCP Matrix					
		S=1	M=1	N=5	
Statistic	Value	F Value	Num DF	Den DF	Pr > F
Wilks' Lambda	0.14071380	18.32	4	12	<.0001
Pillai's Trace	0.85928620	18.32	4	12	<.0001
Hotelling-Lawley Trace	6.10662362	18.32	4	12	<.0001
Roy's Greatest Root	6.10662362	18.32	4	12	<.0001

The **SUMMARY** option in the **MANOVA** statement creates an ANOVA table for each transformed variable as defined by the M matrix. MVAR1 and MVAR2 contrast the pretreatment measurement for Y1 with the posttreatment and follow-up measurements for Y1, respectively; MVAR3 and MVAR4 are the same contrasts for Y2. **Output 47.9.6** displays these univariate ANOVA tables and shows that the contrasts are all strongly significant except for the pre-versus-post difference for Y2.

Output 47.9.6 Summary Output for the Test for Time Effect

The GLM Procedure
Multivariate Analysis of Variance

Dependent Variable: MVAR1

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Intercept	1	512.0000000	512.0000000	22.65	0.0003
Error	15	339.0000000	22.6000000		

The GLM Procedure
Multivariate Analysis of Variance

Dependent Variable: MVAR2

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Intercept	1	813.3888889	813.3888889	32.87	<.0001
Error	15	371.1666667	24.7444444		

The GLM Procedure
Multivariate Analysis of Variance

Dependent Variable: MVAR3

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Intercept	1	68.0555556	68.0555556	3.49	0.0814
Error	15	292.5000000	19.5000000		

Output 47.9.6 *continued*The GLM Procedure
Multivariate Analysis of Variance

Dependent Variable: MVAR4

Source	DF	Type III SS	Mean Square	F Value	Pr > F
Intercept	1	800.0000000	800.0000000	26.43	0.0001
Error	15	454.0000000	30.2666667		

Example 47.10: Testing for Equal Group Variances

This example demonstrates how you can test for equal group variances in a one-way design. The data come from the University of Pennsylvania Smell Identification Test (UPSIT), reported in O'Brien and Heft (1995). The study is undertaken to explore how age and gender are related to sense of smell. A total of 180 subjects 20 to 89 years old are exposed to 40 different odors: for each odor, subjects are asked to choose which of four words best describes the odor. The Freeman-Tukey modified arcsine transformation (Bishop, Fienberg, and Holland 1975) is applied to the proportion of correctly identified odors to arrive at an olfactory index. For the following analysis, subjects are divided into five age groups:

$$\text{agegroup} = \begin{cases} 1 & \text{if } \text{age} \leq 25 \\ 2 & \text{if } 25 < \text{age} \leq 40 \\ 3 & \text{if } 40 < \text{age} \leq 55 \\ 4 & \text{if } 55 < \text{age} \leq 70 \\ 5 & \text{if } 70 < \text{age} \end{cases}$$

The following statements create a data set named `upsit`, containing the age group and olfactory index for each subject.

```
data upsit;
  input agegroup smell @@;
  datalines;
1 1.381 1 1.322 1 1.162 1 1.275 1 1.381 1 1.275 1 1.322
1 1.492 1 1.322 1 1.381 1 1.162 1 1.013 1 1.322 1 1.322
1 1.275 1 1.492 1 1.322 1 1.322 1 1.492 1 1.322 1 1.381
1 1.234 1 1.162 1 1.381 1 1.381 1 1.381 1 1.322 1 1.381
1 1.322 1 1.381 1 1.275 1 1.492 1 1.275 1 1.322 1 1.275
1 1.381 1 1.234 1 1.105
2 1.234 2 1.234 2 1.381 2 1.322 2 1.492 2 1.234 2 1.381
2 1.381 2 1.492 2 1.492 2 1.275 2 1.492 2 1.381 2 1.492
2 1.322 2 1.275 2 1.275 2 1.275 2 1.322 2 1.492 2 1.381
2 1.322 2 1.492 2 1.196 2 1.322 2 1.275 2 1.234 2 1.322
2 1.098 2 1.322 2 1.381 2 1.275 2 1.492 2 1.492 2 1.381
2 1.196
3 1.381 3 1.381 3 1.492 3 1.492 3 1.492 3 1.098 3 1.492
3 1.381 3 1.234 3 1.234 3 1.129 3 1.069 3 1.234 3 1.322
3 1.275 3 1.230 3 1.234 3 1.234 3 1.322 3 1.322 3 1.381
4 1.322 4 1.381 4 1.381 4 1.322 4 1.234 4 1.234 4 1.234
4 1.381 4 1.322 4 1.275 4 1.275 4 1.492 4 1.234 4 1.098
4 1.322 4 1.129 4 0.687 4 1.322 4 1.322 4 1.234 4 1.129
```

```

4 1.492 4 0.810 4 1.234 4 1.381 4 1.040 4 1.381 4 1.381
4 1.129 4 1.492 4 1.129 4 1.098 4 1.275 4 1.322 4 1.234
4 1.196 4 1.234 4 0.585 4 0.785 4 1.275 4 1.322 4 0.712
4 0.810
5 1.322 5 1.234 5 1.381 5 1.275 5 1.275 5 1.322 5 1.162
5 0.909 5 0.502 5 1.234 5 1.322 5 1.196 5 0.859 5 1.196
5 1.381 5 1.322 5 1.234 5 1.275 5 1.162 5 1.162 5 0.585
5 1.013 5 0.960 5 0.662 5 1.129 5 0.531 5 1.162 5 0.737
5 1.098 5 1.162 5 1.040 5 0.558 5 0.960 5 1.098 5 0.884
5 1.162 5 1.098 5 0.859 5 1.275 5 1.162 5 0.785 5 0.859
;

```

Older people are more at risk for problems with their sense of smell, and this should be reflected in significant differences in the mean of the olfactory index across the different age groups. However, many older people also have an excellent sense of smell, which implies that the older age groups should have greater variability. In order to test this hypothesis and to compute a one-way ANOVA for the olfactory index that is robust to the possibility of unequal group variances, you can use the **HOVTEST** and **WELCH** options in the **MEANS** statement for the GLM procedure, as shown in the following statements.

```

proc glm data=upsit;
  class agegroup;
  model smell = agegroup;
  means agegroup / hovtest welch;
run;

```

Output 47.10.1, Output 47.10.2, and Output 47.10.3 display the usual ANOVA test for equal age group means, Levene's test for equal age group variances, and Welch's test for equal age group means, respectively. The hypotheses of age effects for mean and variance of the olfactory index are both confirmed.

Output 47.10.1 Usual ANOVA Test for Age Group Differences in Mean Olfactory Index

The GLM Procedure

Dependent Variable: smell

Source	DF	Type III SS	Mean Square	F Value	Pr > F
agegroup	4	2.13878141	0.53469535	16.65	<.0001

Output 47.10.2 Levene's Test for Age Group Differences in Olfactory Variability

The GLM Procedure

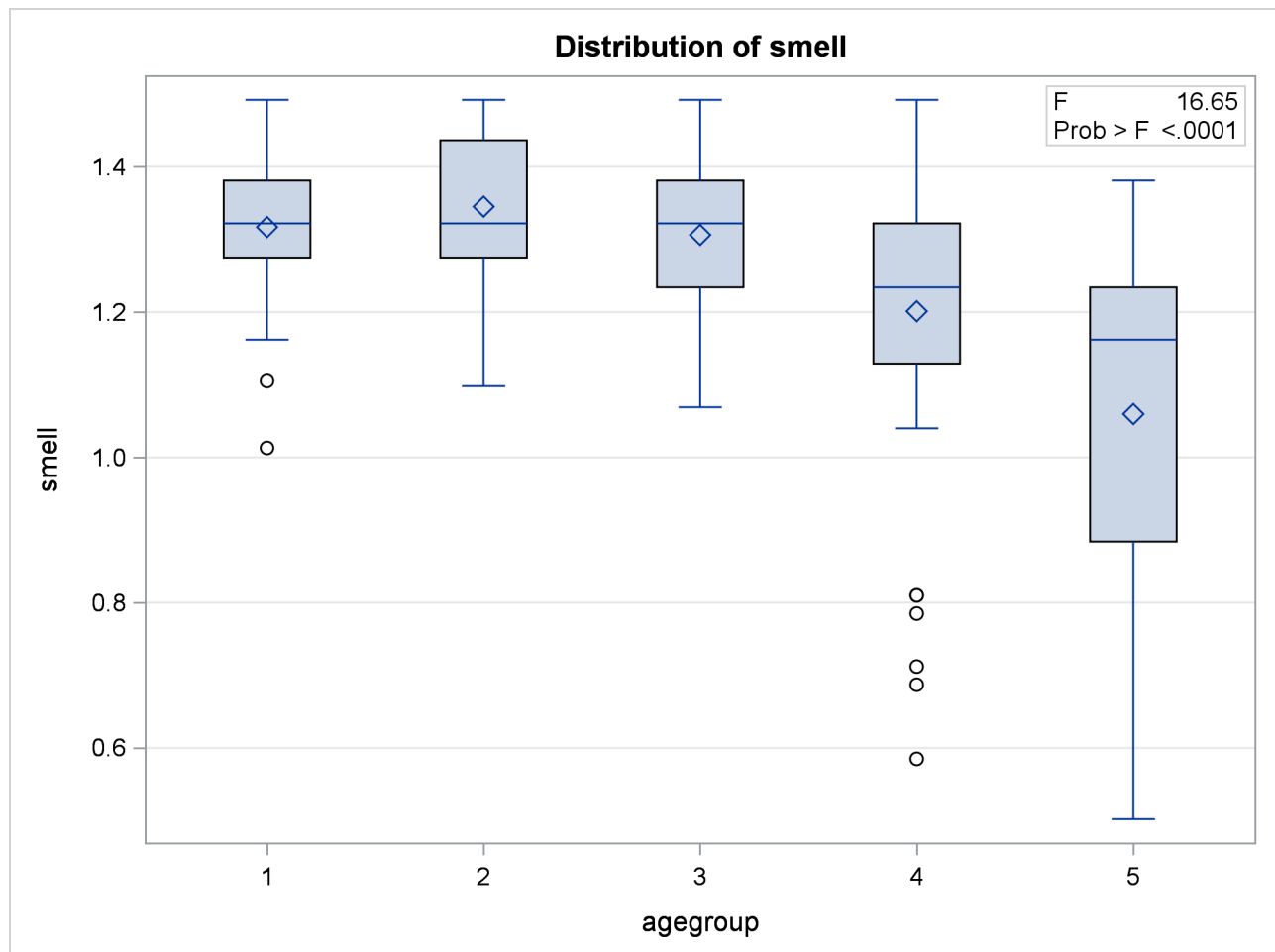
Levene's Test for Homogeneity of smell Variance ANOVA of Squared Deviations from Group Means					
Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
agegroup	4	0.0799	0.0200	6.35	<.0001
Error	175	0.5503	0.00314		

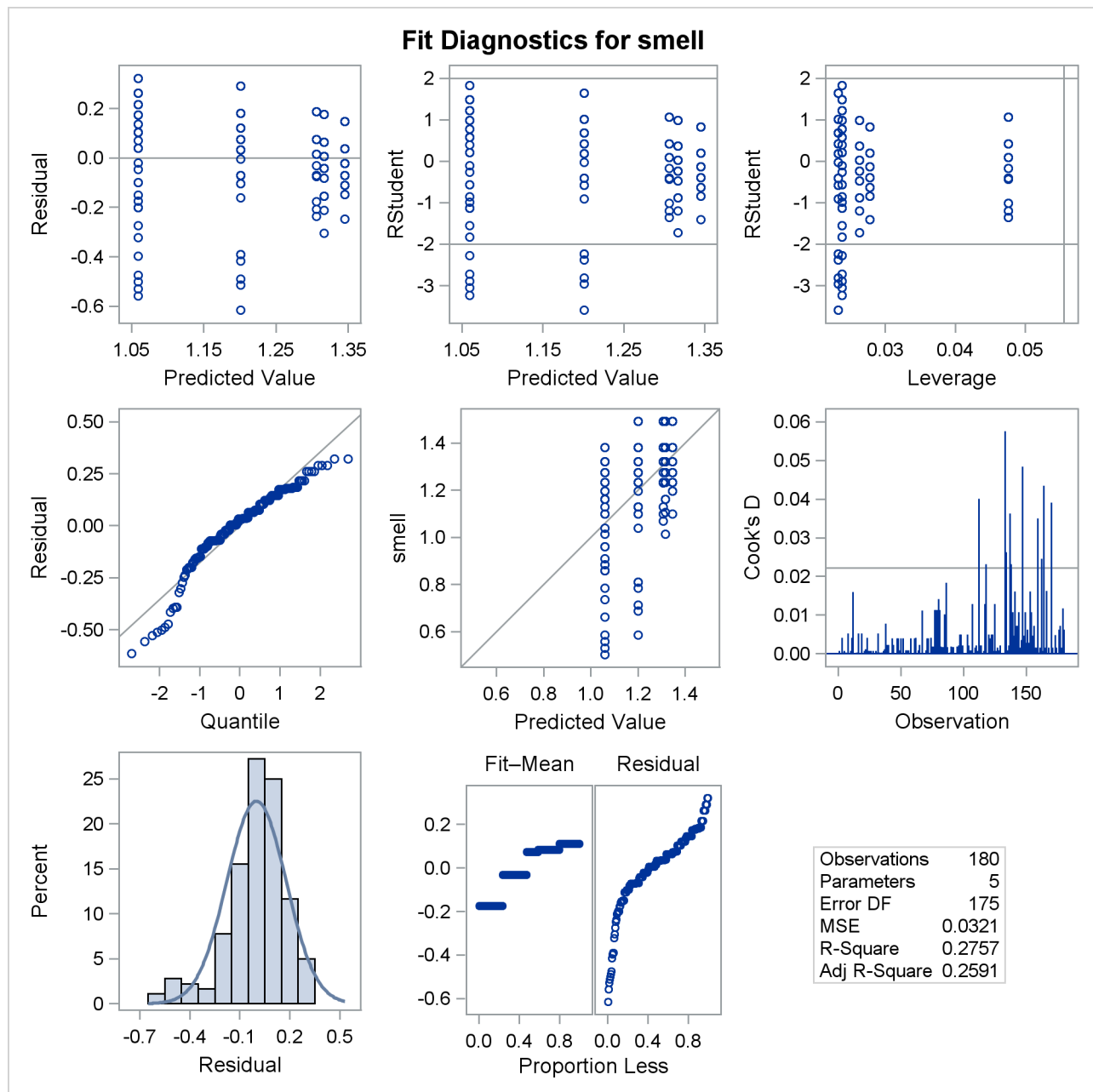
Output 47.10.3 Welch's Test for Age Group Differences in Mean Olfactory Index

Welch's ANOVA for smell			
Source	DF	F Value	Pr > F
agegroup	4.0000	13.72	<.0001
Error	78.7489		

As discussed in “[Homogeneity of Variance in One-Way Models](#)” on page 3706, Levene’s test or any other test for homogeneity of variance should not be used as a diagnostic for the assumption of equal group variances that underlies the usual analysis of variance. However, graphical diagnostics can be a useful informal tool for monitoring whether your data meet the assumptions of a GLM analysis. The following statements perform a one-way ANOVA as before, but with ODS Graphics enabled. In addition to the box plot that is produced by default, the `PLOTS=DIAGNOSTICS` option requests a panel of summary diagnostics for the fit. These additional plots are shown in [Output 47.10.4](#) and [Output 47.10.5](#).

```
ods graphics on;
proc glm data=upsit plot=diagnostics;
  class agegroup;
  model smell = agegroup;
run;
ods graphics off;
```

Output 47.10.4 Box Plot of Olfactory Index by Age Group

Output 47.10.5 Diagnostics for One-Way ANOVA of Olfactory Index by Age Group

Output 47.10.4 clearly shows different degrees of variability for olfactory index within different age groups, with the variability generally rising with age. Likewise, several of the plots in the diagnostics panel shown in Output 47.10.5 indicate a relationship between olfactory variability and mean olfactory index. Also, note that the plot of Cook's D statistic indicates that observations in the higher, more variable age groups are overly influential on the analysis of group means. The overall inference from these plots is that an assumption of equal group variances is probably untenable and that the analysis of the group means should thus take this into account.

Example 47.11: Analysis of a Screening Design

Yin and Jillie (1987) describe an experiment performed on a nitride etch process for a single wafer plasma etcher. The experiment is run using four factors: cathode power (power), gas flow (flow), reactor chamber pressure (pressure), and electrode gap (gap). Of interest are the main effects and interaction effects of the factors on the nitride etch rate (rate). The following statements create a SAS data set named HalfFraction, containing the factor settings and the observed etch rate for each of eight experimental runs.

```
data HalfFraction;
  input power flow pressure gap rate;
  datalines;
0.8   4.5 125 275      550
0.8   4.5 200 325      650
0.8 550.0 125 325      642
0.8 550.0 200 275      601
1.2   4.5 125 325      749
1.2   4.5 200 275     1052
1.2 550.0 125 275     1075
1.2 550.0 200 325      729
;
```

Notice that each of the factors has just two values. This is a common experimental design when the intent is to screen from the many factors that *might* affect the response the few that actually *do*. Since there are $2^4 = 16$ different possible settings of four two-level factors, this design with only eight runs is called a “half fraction.” The eight runs are chosen specifically to provide unambiguous information on main effects at the cost of confounding interaction effects with each other.

One way to analyze these data is simply to use PROC GLM to compute an analysis of variance, including both main effects and interactions in the model. The following statements demonstrate this approach.

```
proc glm data=HalfFraction;
  class power flow pressure gap;
  model rate=power|flow|pressure|gap@2;
run;
```

The “@2” notation in the **MODEL** statement includes all main effects and two-factor interactions between the factors. The output is shown in [Output 47.11.1](#).

Output 47.11.1 Analysis of Variance for Nitride Etch Process Half Fraction

The GLM Procedure

Class Level Information		
Class	Levels	Values
power	2	0.8 1.2
flow	2	4.5 550
pressure	2	125 200
gap	2	275 325

Number of Observations Read	8
Number of Observations Used	8

Output 47.11.1 *continued***The GLM Procedure****Dependent Variable: rate**

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	7	280848.0000	40121.1429	.	.
Error	0	0.0000	.	.	.
Corrected Total	7	280848.0000			

R-Square	Coeff Var	Root MSE	rate Mean
1.000000	.	.	756.0000

Source	DF	Type I SS	Mean Square	F Value	Pr > F
power	1	168780.5000	168780.5000	.	.
flow	1	264.5000	264.5000	.	.
power*flow	1	200.0000	200.0000	.	.
pressure	1	32.0000	32.0000	.	.
power*pressure	1	1300.5000	1300.5000	.	.
flow*pressure	1	78012.5000	78012.5000	.	.
gap	1	32258.0000	32258.0000	.	.
power*gap	0	0.0000	.	.	.
flow*gap	0	0.0000	.	.	.
pressure*gap	0	0.0000	.	.	.

Source	DF	Type III SS	Mean Square	F Value	Pr > F
power	1	168780.5000	168780.5000	.	.
flow	1	264.5000	264.5000	.	.
power*flow	0	0.0000	.	.	.
pressure	1	32.0000	32.0000	.	.
power*pressure	0	0.0000	.	.	.
flow*pressure	0	0.0000	.	.	.
gap	1	32258.0000	32258.0000	.	.
power*gap	0	0.0000	.	.	.
flow*gap	0	0.0000	.	.	.
pressure*gap	0	0.0000	.	.	.

Notice that there are no error degrees of freedom. This is because there are 10 effects in the model (4 main effects plus 6 interactions) but only 8 observations in the data set. This is another cost of using a fractional design: not only is it impossible to estimate all the main effects and interactions, but there is also no information left to estimate the underlying error rate in order to measure the significance of the effects that are estimable.

Another thing to notice in [Output 47.11.1](#) is the difference between the Type I and Type III ANOVA tables. The rows corresponding to main effects in each are the same, but no Type III interaction tests are estimable, while some Type I interaction tests are estimable. This indicates that there is *aliasing* in the design: some interactions are completely confounded with each other.

In order to analyze this confounding, you should examine the aliasing structure of the design by using the **ALIASING** option in the **MODEL** statement. Before doing so, however, it is advisable to *code* the design,

replacing low and high levels of each factor with the values -1 and $+1$, respectively. This puts each factor on an equal footing in the model and makes the aliasing structure much more interpretable. The following statements code the data, creating a new data set named Coded.

```
data Coded; set HalfFraction;
  power = -1*(power =0.80) + 1*(power =1.20);
  flow  = -1*(flow  =4.50) + 1*(flow  =550 );
  pressure = -1*(pressure=125 ) + 1*(pressure=200 );
  gap    = -1*(gap    =275 ) + 1*(gap    =325 );
run;
```

The following statements use the GLM procedure to reanalyze the coded design, displaying the parameter estimates as well as the functions of the parameters that they each estimate.

```
proc glm data=Coded;
  model rate=power|flow|pressure|gap@2 / solution aliasing;
run;
```

The parameter estimates table is shown in [Output 47.11.2](#).

Output 47.11.2 Parameter Estimates and Aliases for Nitride Etch Process Half Fraction

The GLM Procedure

Dependent Variable: rate

Parameter	Estimate	Standard Error	t Value	Pr > t	Expected Value
Intercept	756.0000000	.	.	.	Intercept
power	145.2500000	.	.	.	power
flow	5.7500000	.	.	.	flow
power*flow	-5.0000000 B	.	.	.	power*flow + pressure*gap
pressure	2.0000000	.	.	.	pressure
power*pressure	-12.7500000 B	.	.	.	power*pressure + flow*gap
flow*pressure	-98.7500000 B	.	.	.	flow*pressure + power*gap
gap	-63.5000000	.	.	.	gap
power*gap	0.0000000 B	.	.	.	.
flow*gap	0.0000000 B	.	.	.	.
pressure*gap	0.0000000 B	.	.	.	.

In the “Expected Value” column, notice that, while each of the main effects is unambiguously estimated by its associated term in the model, the expected values of the interaction estimates are more complicated. For example, the relatively large effect (-98.75) corresponding to $\text{flow}*\text{pressure}$ actually estimates the combined effect of $\text{flow}*\text{pressure}$ and $\text{power}*\text{gap}$. Without further information, it is impossible to disentangle these aliased interactions; however, since the main effects of both **power** and **gap** are large and those for **flow** and **pressure** are small, it is reasonable to suspect that $\text{power}*\text{gap}$ is the more “active” of the two interactions.

Fortunately, eight more runs are available for this experiment (the other half fraction). The following statements create a data set containing these extra runs and add it to the previous eight, resulting in a full $2^4 = 16$ run replicate. Then PROC GLM displays the analysis of variance again.

```

data OtherHalf;
  input power flow pressure gap rate;
  datalines;
0.8   4.5 125 325      669
0.8   4.5 200 275      604
0.8 550.0 125 275      633
0.8 550.0 200 325      635
1.2   4.5 125 275     1037
1.2   4.5 200 325      868
1.2 550.0 125 325      860
1.2 550.0 200 275     1063
;
data FullRep;
  set HalfFraction OtherHalf;
run;

proc glm data=FullRep;
  class power flow pressure gap;
  model rate=power|flow|pressure|gap@2;
run;

```

The results are displayed in [Output 47.11.3](#).

Output 47.11.3 Analysis of Variance for Nitride Etch Process Full Replicate

The GLM Procedure

Class Level Information		
Class	Levels	Values
power	2	0.8 1.2
flow	2	4.5 550
pressure	2	125 200
gap	2	275 325

Number of Observations Read	16
Number of Observations Used	16

The GLM Procedure

Dependent Variable: rate

Source	DF	Sum of Squares	Mean Square	F Value	Pr > F
Model	10	521234.1250	52123.4125	25.58	0.0011
Error	5	10186.8125	2037.3625		
Corrected Total	15	531420.9375			

R-Square	Coeff Var	Root MSE	rate Mean
0.980831	5.816175	45.13715	776.0625

Output 47.11.3 continued

Source	DF	Type I SS	Mean Square	F Value	Pr > F
power	1	374850.0625	374850.0625	183.99	<.0001
flow	1	217.5625	217.5625	0.11	0.7571
power*flow	1	18.0625	18.0625	0.01	0.9286
pressure	1	10.5625	10.5625	0.01	0.9454
power*pressure	1	1.5625	1.5625	0.00	0.9790
flow*pressure	1	7700.0625	7700.0625	3.78	0.1095
gap	1	41310.5625	41310.5625	20.28	0.0064
power*gap	1	94402.5625	94402.5625	46.34	0.0010
flow*gap	1	2475.0625	2475.0625	1.21	0.3206
pressure*gap	1	248.0625	248.0625	0.12	0.7414

Source	DF	Type III SS	Mean Square	F Value	Pr > F
power	1	374850.0625	374850.0625	183.99	<.0001
flow	1	217.5625	217.5625	0.11	0.7571
power*flow	1	18.0625	18.0625	0.01	0.9286
pressure	1	10.5625	10.5625	0.01	0.9454
power*pressure	1	1.5625	1.5625	0.00	0.9790
flow*pressure	1	7700.0625	7700.0625	3.78	0.1095
gap	1	41310.5625	41310.5625	20.28	0.0064
power*gap	1	94402.5625	94402.5625	46.34	0.0010
flow*gap	1	2475.0625	2475.0625	1.21	0.3206
pressure*gap	1	248.0625	248.0625	0.12	0.7414

With 16 runs, the analysis of variance tells the whole story: all effects are estimable and there are five degrees of freedom left over to estimate the underlying error. The main effects of power and gap and their interaction are all significant, and no other effects are. Notice that the Type I and Type III ANOVA tables are the same; this is because the design is orthogonal and all effects are estimable.

This example illustrates the use of the GLM procedure for the model analysis of a screening experiment. Typically, there is much more involved in performing an experiment of this type, from selecting the design points to be studied to graphically assessing significant effects, optimizing the final model, and performing subsequent experimentation. Specialized tools for this are available in SAS/QC software, in particular the ADX Interface and the FACTEX and OPTTEX procedures. See *SAS/QC User's Guide* for more information.

References

- Afifi, A. A., and Azen, S. P. (1972). *Statistical Analysis: A Computer-Oriented Approach*. New York: Academic Press.
- Anderson, T. W. (1958). *An Introduction to Multivariate Statistical Analysis*. New York: John Wiley & Sons.
- Bartlett, M. S. (1937). "Properties of Sufficiency and Statistical Tests." *Proceedings of the Royal Society of London, Series A* 160:268–282.

- Begun, J. M., and Gabriel, K. R. (1981). "Closure of the Newman-Keuls Multiple Comparisons Procedure." *Journal of the American Statistical Association* 76:374.
- Belsley, D. A., Kuh, E., and Welsch, R. E. (1980). *Regression Diagnostics: Identifying Influential Data and Sources of Collinearity*. New York: John Wiley & Sons.
- Bishop, Y. M. M., Fienberg, S. E., and Holland, P. W. (1975). *Discrete Multivariate Analysis: Theory and Practice*. Cambridge, MA: MIT Press.
- Box, G. E. P. (1953). "Non-normality and Tests on Variances." *Biometrika* 40:318–335.
- Box, G. E. P. (1954). "Some Theorems on Quadratic Forms Applied in the Study of Analysis of Variance Problems, Part 2: Effects of Inequality of Variance and of Correlation between Errors in the Two-Way Classification." *Annals of Mathematical Statistics* 25:484–498.
- Brown, M. B., and Forsythe, A. B. (1974). "Robust Tests for Equality of Variances." *Journal of the American Statistical Association* 69:364–367.
- Carmer, S. G., and Swanson, M. R. (1973). "Evaluation of Ten Pairwise Multiple Comparison Procedures by Monte Carlo Methods." *Journal of the American Statistical Association* 68:66–74.
- 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.
- Cochran, W. G., and Cox, G. M. (1957). *Experimental Designs*. 2nd ed. New York: John Wiley & Sons.
- Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Cohen, R. (2002). "SAS Meets Big Iron: High Performance Computing in SAS Analytical Procedures." In *Proceedings of the Twenty-Seventh Annual SAS Users Group International Conference*. Cary, NC: SAS Institute Inc. <http://www2.sas.com/proceedings/sugi27/p246-27.pdf>.
- Cole, J. W. L., and Grizzle, J. E. (1966). "Applications of Multivariate Analysis of Variance to Repeated Measures Experiments." *Biometrics* 22:810–828.
- Conover, W. J., Johnson, M. E., and Johnson, M. M. (1981). "A Comparative Study of Tests for Homogeneity of Variances, with Applications to the Outer Continental Shelf Bidding Data." *Technometrics* 23:351–361.
- Cornfield, J., and Tukey, J. W. (1956). "Average Values of Mean Squares in Factorials." *Annals of Mathematical Statistics* 27:907–949.
- Draper, N. R., and Smith, H. (1966). *Applied Regression Analysis*. New York: John Wiley & Sons.
- Duncan, D. B. (1975). "t Tests and Intervals for Comparisons Suggested by the Data." *Biometrics* 31:339–359.
- Dunnett, C. W. (1955). "A Multiple Comparisons Procedure for Comparing Several Treatments with a Control." *Journal of the American Statistical Association* 50:1096–1121.
- Dunnett, C. W. (1980). "Pairwise Multiple Comparisons in the Homogeneous Variance, Unequal Sample Size Case." *Journal of the American Statistical Association* 75:789–795.
- Edwards, D., and Berry, J. J. (1987). "The Efficiency of Simulation-Based Multiple Comparisons." *Biometrics* 43:913–928.

- Einot, I., and Gabriel, K. R. (1975). "A Study of the Powers of Several Methods of Multiple Comparisons." *Journal of the American Statistical Association* 70:351.
- Fidler, F., and Thompson, B. (2001). "Computing Correct Confidence Intervals for ANOVA Fixed- and Random-Effects Effect Sizes." *Educational and Psychological Measurement* 61:575–604.
- Freund, R. J., Littell, R. C., and Spector, P. C. (1986). *SAS System for Linear Models*. 1986 ed. Cary, NC: SAS Institute Inc.
- Gabriel, K. R. (1978). "A Simple Method of Multiple Comparisons of Means." *Journal of the American Statistical Association* 73:724–729.
- Games, P. A. (1977). "An Improved t Table for Simultaneous Control on g Contrasts." *Journal of the American Statistical Association* 72:531–534.
- Goodnight, J. H. (1976). "General Linear Models Procedure." In *Proceedings of the First Annual SAS Users Group International Conference*, 1–39. Cary, NC: SAS Institute Inc. <http://www.sascommunity.org/sugi/SUGI76/Sugi-76-02%20Goodnight.pdf>.
- Goodnight, J. H. (1978a). *The SWEEP Operator: Its Importance in Statistical Computing*. Technical Report R-106, SAS Institute Inc., Cary, NC.
- Goodnight, J. H. (1978b). *Tests of Hypotheses in Fixed-Effects Linear Models*. Technical Report R-101, SAS Institute Inc., Cary, NC.
- Goodnight, J. H. (1979). "A Tutorial on the Sweep Operator." *American Statistician* 33:149–158.
- Goodnight, J. H., and Harvey, W. R. (1978). *Least-Squares Means in the Fixed-Effects General Linear Models*. Technical Report R-103, SAS Institute Inc., Cary, NC.
- Goodnight, J. H., and Speed, F. M. (1978). *Computing Expected Mean Squares*. Technical Report R-102, SAS Institute Inc., Cary, NC.
- Graybill, F. A. (1961). *An Introduction to Linear Statistical Models*. Vol. 1. New York: McGraw-Hill.
- Greenhouse, S. W., and Geisser, S. (1959). "On Methods in the Analysis of Profile Data." *Psychometrika* 32:95–112.
- Gribbin, M. J. (2007). "Better Power Methods for the Univariate Approach to Repeated Measures." Ph.D. diss., Department of Biostatistics, University of North Carolina at Chapel Hill.
- Guirguis, G. H., and Tobias, R. D. (2004). "On the Computation of the Distribution for the Analysis of Means." *Communications in Statistics—Simulation and Computation* 33:861–888.
- Hand, D. J., Daly, F., Lunn, A. D., McConway, K. J., and Ostrowski, E. (1994). *A Handbook of Small Data Sets*. London: Chapman & Hall.
- Hartley, H. O., and Searle, S. R. (1969). "On Interaction Variance Components in Mixed Models." *Biometrics* 25:573–576.
- Harvey, W. R. (1975). *Least-Squares Analysis of Data with Unequal Subclass Numbers*. Technical Report ARS H-4, Agriculture Research Service, U.S. Department of Agriculture.

- Hayter, A. J. (1984). "A Proof of the Conjecture That the Tukey-Kramer Method Is Conservative." *Annals of Statistics* 12:61–75.
- Hayter, A. J. (1989). "Pairwise Comparisons of Generally Correlated Means." *Journal of the American Statistical Association* 84:208–213.
- Heck, D. L. (1960). "Charts of Some Upper Percentage Points of the Distribution of the Largest Characteristic Root." *Annals of Mathematical Statistics* 31:625–642.
- Hochberg, Y. (1974). "Some Conservative Generalizations of the T-Method in Simultaneous Inference." *Journal of Multivariate Analysis* 4:224–234.
- Hocking, R. R. (1973). "A Discussion of the Two-Way Mixed Model." *American Statistician* 27:148–152.
- Hocking, R. R. (1976). "The Analysis and Selection of Variables in a Linear Regression." *Biometrics* 32:1–50.
- Hocking, R. R. (1985). *The Analysis of Linear Models*. Monterey, CA: Brooks/Cole.
- Hsu, J. C. (1992). "The Factor Analytic Approach to Simultaneous Inference in the General Linear Model." *Journal of Computational and Graphical Statistics* 1:151–168.
- Hsu, J. C. (1996). *Multiple Comparisons: Theory and Methods*. London: Chapman & Hall.
- Hsu, J. C., and Nelson, B. L. (1998). "Multiple Comparisons in the General Linear Model." *Journal of Computational and Graphical Statistics* 7:23–41.
- Hsu, J. C., and Peruggia, M. (1994). "Graphical Representation of Tukey's Multiple Comparison Method." *Journal of Computational and Graphical Statistics* 3:143–161.
- Huynh, H., and Feldt, L. S. (1970). "Conditions Under Which Mean Square Ratios in Repeated Measurements Designs Have Exact F-Distributions." *Journal of the American Statistical Association* 65:1582–1589.
- 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.
- Johnson, N. L., Kotz, S., and Balakrishnan, N. (1994). *Continuous Univariate Distributions*. 2nd ed. Vol. 1. New York: John Wiley & Sons.
- Kennedy, W. J., Jr., and Gentle, J. E. (1980). *Statistical Computing*. New York: Marcel Dekker.
- Kramer, C. Y. (1956). "Extension of Multiple Range Tests to Group Means with Unequal Numbers of Replications." *Biometrics* 12:307–310.
- Krishnaiah, P. R., and Armitage, J. V. (1966). "Tables for Multivariate t Distribution." *Sankhyā, Series B* 28:31–56.
- Kutner, M. H. (1974). "Hypothesis Testing in Linear Models (Eisenhart Model)." *American Statistician* 28:98–100.
- LaTour, S. A., and Miniard, P. W. (1983). "The Misuse of Repeated Measures Analysis in Marketing Research." *Journal of Marketing Research* 20:45–57.
- 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.

- Levene, H. (1960). "Robust Tests for the Equality of Variance." In *Contributions to Probability and Statistics: Essays in Honor of Harold Hotelling*, edited by I. Olkin, S. G. Ghurye, W. Hoeffding, W. G. Madow, and H. B. Mann, 278–292. Palo Alto, CA: Stanford University Press.
- Marcus, R., Peritz, E., and Gabriel, K. R. (1976). "On Closed Testing Procedures with Special Reference to Ordered Analysis of Variance." *Biometrika* 63:655–660.
- Maxwell, S. E. (2000). "Sample Size and Multiple Regression Analysis." *Psychological Methods* 5:434–458.
- McLean, R. A., Sanders, W. L., and Stroup, W. W. (1991). "A Unified Approach to Mixed Linear Models." *American Statistician* 45:54–64.
- Miller, R. G., Jr. (1981). *Simultaneous Statistical Inference*. New York: Springer-Verlag.
- Milliken, G. A., and Johnson, D. E. (1984). *Designed Experiments*. Vol. 1 of Analysis of Messy Data. Belmont, CA: Lifetime Learning Publications.
- Morrison, D. F. (1976). *Multivariate Statistical Methods*. 2nd ed. New York: McGraw-Hill.
- Nelder, J. A. (1994). "The Statistics of Linear Models: Back to Basics." *Statistics and Computing* 4:221–234.
- Nelson, P. R. (1982). "Exact Critical Points for the Analysis of Means." *Communications in Statistics—Theory and Methods* 11:699–709.
- Nelson, P. R. (1991). "Numerical Evaluation of Multivariate Normal Integrals with Correlations $\rho_{lj} = -\alpha_l \alpha_j$." In *Frontiers of Statistical Scientific Theory and Industrial Applications: Proceedings of the ICOSCO I Conference*, edited by A. Öztürk and E. C. van der Meulen, 97–114. Columbus, OH: American Sciences Press.
- Nelson, P. R. (1993). "Additional Uses for the Analysis of Means and Extended Tables of Critical Values." *Technometrics* 35:61–71.
- O'Brien, R. G. (1979). "A General ANOVA Method for Robust Tests of Additive Models for Variances." *Journal of the American Statistical Association* 74:877–880.
- O'Brien, R. G. (1981). "A Simple Test for Variance Effects in Experimental Designs." *Psychological Bulletin* 89:570–574.
- O'Brien, R. G., and Heft, M. W. (1995). "New Discrimination Indexes and Models for Studying Sensory Functioning in Aging." *Journal of Applied Statistics* 22:9–27.
- Olejnik, S. F., and Algina, J. (1987). "Type I Error Rates and Power Estimates of Selected Parametric and Non-parametric Tests of Scale." *Journal of Educational Statistics* 12:45–61.
- Ott, E. R. (1967). "Analysis of Means: A Graphical Procedure." *Industrial Quality Control* 24:101–109. Reprinted in *Journal of Quality Technology* 15 (1983): 10–18.
- Perlman, M. D., and Rasmussen, U. A. (1975). "Some Remarks on Estimating a Noncentrality Parameter." *Communications in Statistics—Theory and Methods* 4:455–468.
- Petrinovich, L. F., and Hardyck, C. D. (1969). "Error Rates for Multiple Comparison Methods: Some Evidence Concerning the Frequency of Erroneous Conclusions." *Psychological Bulletin* 71:43–54.

- Pillai, K. C. S. (1960). *Statistical Table for Tests of Multivariate Hypotheses*. Manila: University of Philippines Statistical Center.
- Pringle, R. M., and Rayner, A. A. (1971). *Generalized Inverse Matrices with Applications to Statistics*. New York: Hafner Publishing.
- Ramsey, P. H. (1978). "Power Differences between Pairwise Multiple Comparisons." *Journal of the American Statistical Association* 73:479–485.
- Rao, C. R. (1965). *Linear Statistical Inference and Its Applications*. New York: John Wiley & Sons.
- Rodriguez, R. N., Tobias, R. D., and Wolfinger, R. D. (1995). "Comments on J. A. Nelder, 'The Statistics of Linear Models: Back to Basics'." *Statistics and Computing* 5:97–101.
- Ryan, T. A. (1959). "Multiple Comparisons in Psychological Research." *Psychological Bulletin* 56:26–47.
- Ryan, T. A. (1960). "Significance Tests for Multiple Comparison of Proportions, Variances, and Other Statistics." *Psychological Bulletin* 57:318–328.
- Satterthwaite, F. E. (1946). "An Approximate Distribution of Estimates of Variance Components." *Biometrics Bulletin* 2:110–114.
- Schatzoff, M. (1966). "Exact Distributions of Wilks's Likelihood Ratio Criterion." *Biometrika* 53:347–358.
- Scheffé, H. (1953). "A Method for Judging All Contrasts in the Analysis of Variance." *Biometrika* 40:87–104.
- Scheffé, H. (1959). *The Analysis of Variance*. New York: John Wiley & Sons.
- Searle, S. R. (1971). *Linear Models*. New York: John Wiley & Sons.
- Searle, S. R. (1987). *Linear Models for Unbalanced Data*. New York: John Wiley & Sons.
- Searle, S. R. (1995). "Comments on J. A. Nelder, 'The Statistics of Linear Models: Back to Basics'." *Statistics and Computing* 5:103–107.
- Searle, S. R., Casella, G., and McCulloch, C. E. (1992). *Variance Components*. New York: John Wiley & Sons.
- Searle, S. R., Speed, F. M., and Milliken, G. A. (1980). "Population Marginal Means in the Linear Model: An Alternative to Least Squares Means." *American Statistician* 34:216–221.
- Šidák, Z. (1967). "Rectangular Confidence Regions for the Means of Multivariate Normal Distributions." *Journal of the American Statistical Association* 62:626–633.
- Smithson, M. (2003). *Confidence Intervals*. Thousand Oaks, CA: Sage Publications.
- Smithson, M. (2004). Personal communication.
- Snedecor, G. W., and Cochran, W. G. (1967). *Statistical Methods*. 6th ed. Ames: Iowa State University Press.
- Steel, R. G. D., and Torrie, J. H. (1960). *Principles and Procedures of Statistics*. New York: McGraw-Hill.
- Steiger, J. H., and Fouladi, R. T. (1997). "Noncentrality Interval Estimation and the Evaluation of Statistical Models." In *What If There Were No Significance Tests?*, edited by L. Harlow, S. Mulaik, and J. H. Steiger, 222–257. Hillsdale, NJ: Lawrence Erlbaum Associates.

- Stenstrom, F. H. (1940). "The Growth of Snapdragons, Stocks, Cinerarias, and Carnations on Six Iowa Soils." Master's thesis, Iowa State College.
- Tubb, A., Parker, A. J., and Nickless, G. (1980). "The Analysis of Romano-British Pottery by Atomic Absorption Spectrophotometry." *Archaeometry* 22:153–171.
- Tukey, J. W. (1952). "Allowances for Various Types of Error Rates." Invited address to Blacksburg meeting of Institute of Mathematical Studies.
- Tukey, J. W. (1953). "The Problem of Multiple Comparisons." In *Multiple Comparisons, 1948–1983*, edited by H. I. Braun, vol. 8 of *The Collected Works of John W. Tukey* (published 1994), 1–300. London: Chapman & Hall. Unpublished manuscript.
- Urquhart, N. S. (1968). "Computation of Generalized Inverse Matrices Which Satisfy Specific Conditions." *SIAM Review* 10:216–218.
- Waller, R. A., and Duncan, D. B. (1969). "A Bayes Rule for the Symmetric Multiple Comparison Problem." *Journal of the American Statistical Association* 64:1484–1499.
- Waller, R. A., and Duncan, D. B. (1972). "Corrigenda to 'A Bayes Rule for the Symmetric Multiple Comparison Problem'." *Journal of the American Statistical Association* 67:253–255.
- Waller, R. A., and Kemp, K. E. (1976). "Computations of Bayesian *t*-Values for Multiple Comparisons." *Journal of Statistical Computation and Simulation* 75:169–172.
- Welch, B. L. (1951). "On the Comparison of Several Mean Values: An Alternative Approach." *Biometrika* 38:330–336.
- Welsch, R. E. (1977). "Stepwise Multiple Comparison Procedures." *Journal of the American Statistical Association* 72:566–575.
- Westfall, P. H., and Young, S. S. (1993). *Resampling-Based Multiple Testing: Examples and Methods for p-Value Adjustment*. New York: John Wiley & Sons.
- Winer, B. J. (1971). *Statistical Principles in Experimental Design*. 2nd ed. New York: McGraw-Hill.
- Wolfinger, R. D., and Chang, M. (1995). "Comparing the SAS GLM and MIXED Procedures for Repeated Measures." In *Proceedings of the Twentieth Annual SAS Users Group Conference*, 1172–1182. Cary, NC: SAS Institute Inc. <http://www.sascommunity.org/sugi/SUGI95/Sugi-95-198%20Wolfinger%20Chang.pdf>.
- Yin, G. Z., and Jillie, D. W. (1987). "Orthogonal Design for Process Optimization and Its Application in Plasma Etching." *Solid State Technology* 30:127–132.

Subject Index

- absorption of effects
 - GLM procedure, 3630, 3687
- adjusted means
 - See least squares means, 3637
- aliasing structure
 - GLM procedure, 3657
- aliasing structure (GLM), 3782
- alpha level
 - GLM procedure, 3640, 3652, 3657, 3662
- analysis of covariance
 - examples (GLM), 3748
 - MODEL statements (GLM), 3672
- analysis of means
 - comparing LS-means (GLM), 3700
- analysis of variance
 - mixed models (GLM), 3769
 - MODEL statements (GLM), 3672
 - multivariate (GLM), 3645, 3710, 3711, 3758
 - one-way, variance-weighted, 3655
 - repeated measures (GLM), 3712, 3765, 3772
 - three-way design (GLM), 3754
 - unbalanced (GLM), 3613, 3692, 3742
- ANOVA procedure
 - compared to other procedures, 3612
- at sign (@) operator
 - GLM procedure, 3672
- balanced data
 - example, complete block, 3735
- bar (l) operator
 - GLM procedure, 3672
- Bartlett's test
 - GLM procedure, 3653, 3706
- between-subject factors
 - repeated measures, 3664, 3714
- Bonferroni adjustment
 - GLM procedure, 3638
- Bonferroni *t* test, 3652, 3697
- Box's epsilon, 3715
- Brown and Forsythe's test
 - GLM procedure, 3653, 3706
- canonical analysis
 - GLM procedure, 3647
- CATMOD procedure
 - compared to other procedures, 3677
- characteristic roots and vectors
 - GLM procedure, 3646
- classification variables
 - GLM procedure, 3671
- comparing
 - groups (GLM), 3691
- comparisonwise error rate (GLM), 3697
- complete block design
 - example (GLM), 3735
- continuous variables, 3671
- continuous-by-class effects
 - model parameterization (GLM), 3676
 - specifying (GLM), 3671
- continuous-nesting-class effects
 - model parameterization (GLM), 3675
 - specifying (GLM), 3671
- contrasts
 - GLM procedure, 3633
 - repeated measures (GLM), 3665
- control
 - comparing treatments to (GLM), 3695, 3699
- Cook's *D* influence statistic, 3660
- covariates
 - model parameterization (GLM), 3674
- crossed effects
 - model parameterization (GLM), 3674
 - specifying (GLM), 3671
- degrees of freedom
 - models with classification variables (GLM), 3677
- DFFITS statistic
 - GLM procedure, 3660
- discrete variables, *see* classification variables
- Duncan's multiple range test, 3652, 3702
- Duncan-Waller test, 3655, 3703
 - error seriousness ratio, 3654
- Dunnett's adjustment
 - GLM procedure, 3638
- Dunnett's test, 3652, 3653, 3699, 3700
 - one-tailed lower, 3653
- effect
 - definition, 3670
 - specification (GLM), 3670
- effect sizes
 - GLM procedure, 3683
 - MODEL statement (GLM), 3658
- Einot and Gabriel's multiple range test
 - examples (GLM), 3738
 - GLM procedure, 3654, 3703
- error seriousness ratio
 - Waller-Duncan test, 3654

- estimability
 - GLM procedure, 3634
- estimable functions
 - checking (GLM), 3634
 - displaying (GLM), 3657, 3658
 - example (GLM), 3678
 - general form of, 3679
 - GLM procedure, 3635, 3636, 3645, 3659, 3678, 3679, 3681, 3682, 3691
 - printing (GLM), 3657
- estimated population marginal means, *see* least squares means
- expected mean squares
 - computing, types (GLM), 3722
 - random effects, 3720
- experimental design, 3785
 - aliasing structure (GLM), 3782
- experimentwise error rate (GLM), 3697
- Fisher's LSD test, 3655
- fixed effects
 - sum-to-zero assumptions, 3721
- Forward-Dolittle transformation, 3680
- G-G epsilon, 3715
- Gabriel's multiple-comparison procedure
 - GLM procedure, 3653, 3699
- GLM procedure
 - absorption of effects, 3630, 3687
 - aliasing structure, 3657, 3782
 - alpha level, 3640, 3652, 3657, 3662
 - Bartlett's test, 3653, 3706
 - Bonferroni adjustment, 3638
 - Brown and Forsythe's test, 3653, 3706
 - canonical analysis, 3647
 - characteristic roots and vectors, 3646
 - compared to other procedures, 3612, 3663, 3720, 3771, 3785
 - comparing groups, 3691
 - computational method, 3726
 - computational resources, 3723
 - contrasts, 3633, 3666
 - covariate values for least squares means, 3640
 - diffogram, 3643
 - disk space, 3625
 - Dunnett's adjustment, 3638
 - effect sizes, 3683
 - effect specification, 3670
 - error effect, 3647
 - estimability, 3634–3636, 3645, 3659, 3679, 3691
 - estimable functions, 3678
 - ESTIMATE specification, 3690
 - homogeneity of variance tests, 3653, 3706
 - Hsu's adjustment, 3638
 - hypothesis tests, 3668, 3678
 - interactive use, 3673
 - interactivity and BY statement, 3631
 - interactivity and missing values, 3673, 3723
 - introductory example, 3613
 - least squares means (LS-means), 3637
 - Levene's test for homogeneity of variance, 3653, 3706
 - means, 3650
 - means versus least squares means, 3692
 - memory requirements, reduction of, 3630
 - missing values, 3625, 3646, 3723
 - model specification, 3670
 - multiple comparisons, least squares means, 3638, 3641, 3693, 3696
 - multiple comparisons, means, 3652–3655, 3693, 3696
 - multiple comparisons, procedures, 3650
 - multivariate analysis of variance, 3625, 3645, 3710
 - Nelson's adjustment, 3638
 - nonstandard weights for least squares means, 3641
 - O'Brien's test, 3653
 - observed margins for least squares means, 3641
 - ODS graph names, 3734
 - ODS table names, 3729
 - ordering of effects, 3625
 - output data sets, 3660, 3726–3728
 - parameterization, 3673
 - positional requirements for statements, 3622
 - predicted population margins, 3638
 - Q effects, 3721
 - random effects, 3662, 3719, 3720
 - regression, quadratic, 3616
 - repeated measures, 3664, 3712
 - Sidak's adjustment, 3638
 - simple effects, 3645
 - simulation-based adjustment, 3638
 - singularity checking, 3634, 3636, 3645, 3658
 - sphericity tests, 3667, 3715
 - SSCP matrix for multivariate tests, 3646
 - statistical assumptions, 3670
 - summary of features, 3611
 - tests, hypothesis, 3633
 - transformations for MANOVA, 3647
 - transformations for repeated measures, 3665
 - Tukey's adjustment, 3638
 - types of least squares means comparisons, 3641
 - unbalanced analysis of variance, 3613, 3692, 3742
 - unbalanced design, 3613, 3692, 3720, 3742, 3769
 - weighted analysis, 3669
 - weighted means, 3707

- Welch's ANOVA, 3655
- WHERE statement, 3673
- Greenhouse-Geisser epsilon, 3715
- GT2 multiple-comparison method, 3655, 3699
- H-F epsilon, 3716
- half-fraction design, analysis, 3781
- Hochberg's GT2 multiple-comparison method, 3655, 3699
- homogeneity of variance tests, 3653, 3706
 - Bartlett's test (GLM), 3653, 3706
 - Brown and Forsythe's test (GLM), 3653, 3706
 - examples, 3778
 - Levene's test (GLM), 3653, 3706
 - O'Brien's test (GLM), 3653
 - Welch's ANOVA, 3706
- honestly significant difference test, 3655, 3698, 3700
- Hotelling-Lawley trace, 3646, 3714
- HSD test, 3655, 3698, 3700
- Hsu's adjustment
 - GLM procedure, 3638
- Huynh-Feldt
 - epsilon (GLM), 3716
 - structure (GLM), 3715
- hypothesis tests
 - comparing adjusted means (GLM), 3645
 - contrasts (GLM), 3633
 - contrasts, examples (GLM), 3737, 3755, 3763
 - customized (GLM), 3668
 - for intercept (GLM), 3658
 - GLM procedure, 3678
 - MANOVA (GLM), 3711
 - random effects (GLM), 3663, 3720
 - repeated measures (GLM), 3714
 - Type I sum of squares (GLM), 3679
 - Type II sum of squares (GLM), 3681
 - Type III sum of squares (GLM), 3682
 - Type IV sum of squares (GLM), 3682
- interaction effects
 - model parameterization (GLM), 3674
 - specifying (GLM), 3671
- intercept
 - hypothesis tests for (GLM), 3658
 - model parameterization (GLM), 3674
- least squares means
 - Bonferroni adjustment (GLM), 3638
 - coefficient adjustment, 3709
 - compared to means (GLM), 3692
 - comparison types (GLM), 3641
 - construction of, 3707
 - covariate values (GLM), 3640
 - Dunnett's adjustment (GLM), 3638
 - examples (GLM), 3744, 3756
 - GLM procedure, 3637
 - Hsu's adjustment (GLM), 3638
 - multiple comparisons adjustment (GLM), 3638, 3641
 - Nelson's adjustment (GLM), 3638
 - nonstandard weights (GLM), 3641
 - observed margins (GLM), 3641
 - Sidak's adjustment (GLM), 3638
 - simple effects (GLM), 3645, 3704
 - simulation-based adjustment (GLM), 3638
 - Tukey's adjustment (GLM), 3638
 - least-significant-difference test, 3655
 - levels, of classification variable, 3671
 - Levene's test for homogeneity of variance
 - GLM procedure, 3653, 3706, 3778
 - leverage, 3660
 - LS-means, *see* least squares means
 - LSD test, 3655
- main effects
 - model parameterization (GLM), 3674
 - specifying (GLM), 3671
- MANOVA, *see* multivariate analysis of variance, *see* multivariate analysis of variance
- mean separation tests, *see* multiple-comparison procedures
- means
 - compared to least squares means (GLM), 3692
 - GLM procedure, 3650
 - weighted (GLM), 3707
- memory requirements
 - reduction of (GLM), 3630
- missing values
 - and interactivity (GLM), 3673
- mixed model
 - unbalanced (GLM), 3769
- MIXED procedure
 - contrasted SAS procedures, 3612, 3720, 3771
- model
 - parameterization (GLM), 3673
 - specification (GLM), 3670
- multiple comparison procedures
 - GLM procedure, 3650
 - multiple-stage tests, 3701
 - pairwise (GLM), 3695
 - recommendations, 3704
 - with a control (GLM), 3699
 - with the average (GLM), 3700
- multiple comparisons of least squares means, *see also* multiple-comparison procedures
 - GLM procedure, 3638, 3641, 3696
 - interpretation, 3704
- multiple comparisons of means, *see also* multiple-comparison procedures

- Bonferroni t test, 3652
- Duncan's multiple range test, 3652
- Dunnett's test, 3652, 3653
- error mean square, 3653
- examples, 3735
- Fisher's LSD test, 3655
- Gabriel's procedure, 3653
- GLM procedure, 3693, 3696
- GT2 method, 3655
- interpretation, 3704
- Ryan-Einot-Gabriel-Welsch test, 3654
- Scheffé's procedure, 3654
- Sidak's adjustment, 3654
- SMM, 3655
- Student-Newman-Keuls test, 3655
- Tukey's studentized range test, 3655
- Waller-Duncan test, 3655
- multiple-comparison procedures, 3693, *see also*
 - multiple comparisons of least squares means, *see also* multiple comparisons of means
 - pairwise (GLM), 3694
 - with a control (GLM), 3695
- multiple-stage tests, *see* multiple comparison procedures, 3701
- multivariate analysis of variance
 - examples (GLM), 3758
 - GLM procedure, 3625, 3645, 3710
 - hypothesis tests (GLM), 3711
 - partial correlations, 3711
- multivariate general linear hypothesis, 3711
- multivariate tests
 - repeated measures, 3714, 3715
- Nelson's adjustment
 - GLM procedure, 3638
- nested effects
 - model parameterization (GLM), 3675
 - specifying (GLM), 3671
- NESTED procedure
 - compared to other procedures, 3612
- Newman-Keuls' multiple range test, 3655, 3702
- nominal variables, *see also* classification variables
- NPAR1WAY procedure
 - compared to other procedures, 3612
- O'Brien's test for homogeneity of variance
 - GLM procedure, 3653
- ODS graph names
 - GLM procedure, 3734
- orthonormalizing transformation matrix
 - GLM procedure, 3648
- pairwise comparisons
 - GLM procedure, 3694, 3695
- parameterization
 - of models (GLM), 3673
- partial correlations
 - multivariate analysis of variance, 3711
- Pillai's trace, 3646, 3714
- polynomial effects
 - model parameterization (GLM), 3674
 - specifying (GLM), 3671
- predicted population margins
 - GLM procedure, 3638
- PRESS statistic, 3661
- quadratic forms for fixed effects
 - displaying (GLM), 3663
- quadratic regression, 3616
- qualitative variables, *see* classification variables
- R-notation, 3679
- random effects
 - expected mean squares, 3720
 - GLM procedure, 3662, 3719
- randomized complete block design
 - example, 3735
- reduction notation, 3679
- REG procedure
 - compared to other procedures, 3612
- regression
 - examples (GLM), 3739
 - MODEL statements (GLM), 3672
 - quadratic (GLM), 3616
- regression effects
 - model parameterization (GLM), 3674
 - specifying (GLM), 3671
- repeated measures
 - contrasts (GLM), 3665
 - data organization (GLM), 3712
 - doubly multivariate design, 3772
 - examples (GLM), 3667, 3765
 - GLM procedure, 3664, 3712
 - hypothesis tests (GLM), 3714, 3717
 - more than one factor (GLM), 3717
 - transformations, 3717–3719
- response variable, 3671
- Roy's greatest root, 3646, 3714
- RSREG procedure
 - compared to other procedures, 3612
- Ryan's multiple range test, 3654, 3703
 - examples, 3738
- Satterthwaite's approximation
 - testing random effects, 3721
- Scheffé's multiple-comparison procedure, 3698
 - p, 3654
- screening design, analysis, 3781
- Sidak's adjustment
 - GLM procedure, 3638

- Sidak's t test, 3654, 3697
- simple effects
 - GLM procedure, 3645, 3704
- simulation-based adjustment
 - GLM procedure, 3638
- singularity checking
 - GLM procedure, 3634, 3636, 3645, 3658
- SMM multiple-comparison method, 3655, 3699
- sphericity tests, 3667, 3768
- SSCP matrix
 - for multivariate tests (GLM), 3646, 3648
- statistical
 - assumptions (GLM), 3670
- stepdown methods
 - GLM procedure, 3701
- Student's multiple range test, 3655, 3702
- Studentized maximum modulus
 - pairwise comparisons, 3655, 3699
- studentized residual, 3661
- sum-to-zero assumptions, 3721
- sums of squares
 - GLM procedure, 3659
 - Type II (GLM), 3659
- tests, hypothesis
 - examples (GLM), 3737
 - GLM procedure, 3633
- transformation matrix
 - orthonormalizing, 3648
- transformations
 - for multivariate ANOVA, 3647
 - repeated measures, 3717–3719
- transformations for repeated measures
 - GLM procedure, 3665
- TTEST procedure
 - compared to other procedures, 3613
- Tukey's adjustment
 - GLM procedure, 3638
- Tukey's studentized range test, 3655, 3698, 3700
- Tukey-Kramer test, 3655, 3698, 3700
- Type I error rate
 - repeated multiple comparisons, 3696
- Type H covariance structure, 3715
- Type I sum of squares
 - computing in GLM, 3722
 - displaying (GLM), 3659
 - estimable functions for, 3657
 - estimable functions for (GLM), 3679
 - examples, 3744
- Type II sum of squares
 - computing in GLM, 3722
 - displaying (GLM), 3659
 - estimable functions for, 3657
 - estimable functions for (GLM), 3681
 - examples, 3744
- Type III sum of squares
 - displaying (GLM), 3659
 - estimable functions for, 3657
 - estimable functions for (GLM), 3682
 - examples, 3744
- Type IV sum of squares
 - computing in GLM, 3722
 - displaying (GLM), 3659
 - estimable functions for, 3658
 - estimable functions for (GLM), 3682
 - examples, 3744
- unbalanced design
 - GLM procedure, 3613, 3692, 3720, 3742, 3769
- univariate tests
 - repeated measures, 3714, 3715
- VARCOMP procedure
 - compared to other procedures, 3613
- Waller-Duncan test, 3655, 3703
 - error seriousness ratio, 3654
 - examples, 3738
- weighted least squares
 - normal equations (GLM), 3669
- weighted means
 - GLM procedure, 3707
- Welch's ANOVA, 3655
 - homogeneity of variance tests, 3706
 - using homogeneity of variance tests, 3778
- Welsch's multiple range test, 3654, 3703
 - examples, 3738
- WHERE statement
 - GLM procedure, 3673
- Wilks' lambda, 3646, 3714
- within-subject factors
 - repeated measures, 3664, 3714

Syntax Index

- ABSORB statement
 - GLM procedure, [3630](#)
- ADJUST= option
 - LSMEANS statement (GLM), [3638](#)
- ALIASING option
 - MODEL statement (GLM), [3657](#), [3782](#)
- _ALL_ effect
 - MANOVA statement, H= option (GLM), [3646](#)
- ALPHA= option
 - LSMEANS statement (GLM), [3640](#)
 - MEANS statement (GLM), [3652](#)
 - MODEL statement (GLM), [3657](#)
 - OUTPUT statement (GLM), [3662](#)
 - PROC GLM statement, [3625](#)
- AT option
 - LSMEANS statement (GLM), [3640](#), [3708](#)
- BON option
 - MEANS statement (GLM), [3652](#)
- BY statement
 - GLM procedure, [3630](#)
- BYLEVEL option
 - LSMEANS statement (GLM), [3640](#), [3710](#)
- CANONICAL option
 - MANOVA statement (GLM), [3647](#)
 - REPEATED statement (GLM), [3666](#)
- CL option
 - LSMEANS statement (GLM), [3640](#)
- CLASS statement
 - GLM procedure, [3631](#)
- CLDIFF option
 - MEANS statement (GLM), [3652](#)
- CLI option
 - MODEL statement (GLM), [3657](#)
- CLM option
 - MEANS statement (GLM), [3652](#)
 - MODEL statement (GLM), [3657](#)
- CLPARM option
 - MODEL statement (GLM), [3657](#)
- CODE statement
 - GLM procedure, [3632](#)
- CONTRAST option
 - REPEATED statement (GLM), [3666](#), [3717](#)
- CONTRAST statement
 - GLM procedure, [3633](#)
- COOKD keyword
 - OUTPUT statement (GLM), [3660](#)
- COV option
 - LSMEANS statement (GLM), [3640](#)
- COVRATIO keyword
 - OUTPUT statement (GLM), [3660](#)
- DATA= option
 - PROC GLM statement, [3625](#)
- DEPONLY option
 - MEANS statement (GLM), [3652](#)
- DFFITS keyword
 - OUTPUT statement (GLM), [3660](#)
- DIVISOR= option
 - ESTIMATE statement (GLM), [3636](#)
- DUNCAN option
 - MEANS statement (GLM), [3652](#)
- DUNNETT option
 - MEANS statement (GLM), [3652](#)
- DUNNETTL option
 - MEANS statement (GLM), [3653](#)
- DUNNETTU option
 - MEANS statement (GLM), [3653](#)
- E option
 - CONTRAST statement (GLM), [3634](#)
 - ESTIMATE statement (GLM), [3636](#)
 - LSMEANS statement (GLM), [3640](#)
 - MODEL statement (GLM), [3657](#)
- E1 option
 - MODEL statement (GLM), [3657](#)
- E2 option
 - MODEL statement (GLM), [3657](#)
- E3 option
 - MODEL statement (GLM), [3657](#)
- E4 option
 - MODEL statement (GLM), [3658](#)
- E= option
 - CONTRAST statement (GLM), [3634](#)
 - MANOVA statement (GLM), [3647](#)
 - MEANS statement (GLM), [3653](#)
 - REPEATED statement (GLM), [3669](#)
- EFFECTSIZE option
 - MODEL statement (GLM), [3658](#)
- ESTIMATE statement
 - GLM procedure, [3635](#), [3690](#)
- ETYPE option
 - LSMEANS statement (GLM), [3641](#)
- ETYPE= option
 - CONTRAST statement (GLM), [3634](#)
 - MANOVA statement (GLM), [3648](#)
 - MEANS statement (GLM), [3653](#)

- TEST statement (GLM), 3669
- factor specification
 - REPEATED statement (GLM), 3665
- FREQ statement
 - GLM procedure, 3636
- GABRIEL option
 - MEANS statement (GLM), 3653
- GLM procedure
 - syntax, 3622
 - GLM procedure, ABSORB statement, 3630
 - GLM procedure, BY statement, 3630
 - GLM procedure, CLASS statement, 3631
 - REF= option, 3632
 - REF= variable option, 3632
 - TRUNCATE option, 3632
 - GLM procedure, CODE statement, 3632
 - GLM procedure, CONTRAST statement, 3633
 - E option, 3634
 - E= option, 3634
 - ETYPE= option, 3634
 - INTERCEPT effect, 3633, 3635
 - SINGULAR= option, 3634
 - GLM procedure, ESTIMATE statement, 3635
 - DIVISOR= option, 3636
 - E option, 3636
 - SINGULAR= option, 3636
 - GLM procedure, FREQ statement, 3636
 - GLM procedure, ID statement, 3637
 - GLM procedure, LSMEANS statement, 3637
 - ADJUST= option, 3638
 - ALPHA= option, 3640
 - AT option, 3640, 3708
 - BYLEVEL option, 3710
 - BYLEVEL options, 3640
 - CL option, 3640
 - COV option, 3640
 - E option, 3640
 - ETYPE option, 3641
 - LINES option, 3641
 - NOPRINT option, 3641
 - OBSMARGINS option, 3641, 3709
 - OM option, 3641, 3709
 - OUT= option, 3641
 - PDIFF option, 3641
 - PLOTS= option, 3643
 - SINGULAR option, 3645
 - SLICE= option, 3645
 - STDERR option, 3645
 - TDIFF option, 3645
 - GLM procedure, MANOVA statement, 3645
 - _ALL_ effect (H= option), 3646
 - CANONICAL option, 3647
 - E= option, 3647
 - ETYPE= option, 3648
 - H= option, 3646
 - HTYPE= option, 3648
 - INTERCEPT effect (H= option), 3646
 - M= option, 3647
 - MNAMES= option, 3647
 - MSTAT= option, 3648
 - ORTH option, 3648
 - PREFIX= option, 3647
 - PRINTE option, 3648
 - PRINTH option, 3648
 - SUMMARY option, 3648
 - GLM procedure, MEANS statement, 3650, 3707
 - ALPHA= option, 3652
 - BON option, 3652
 - CLDIFF option, 3652
 - CLM option, 3652
 - DEPONLY option, 3652
 - DUNCAN option, 3652
 - DUNNETT option, 3652
 - DUNNETTL option, 3653
 - DUNNETTU option, 3653
 - E= option, 3653
 - ETYPE= option, 3653
 - GABRIEL option, 3653
 - GT2 option, 3653
 - HOVTEST option, 3653, 3706
 - HTYPE= option, 3654
 - KRATIO= option, 3654
 - LINES option, 3654
 - LSD option, 3654
 - NOSORT option, 3654
 - REGWQ option, 3654
 - SCHEFFE option, 3654
 - SIDAK option, 3654
 - SMM option, 3655
 - SNK option, 3655
 - T option, 3655
 - TUKEY option, 3655
 - WALLER option, 3655
 - WELCH option, 3655
 - GLM procedure, MODEL statement, 3655
 - ALIASING option, 3657, 3782
 - ALPHA= option, 3657
 - CLI option, 3657
 - CLM option, 3657
 - CLPARM option, 3657
 - E option, 3657
 - E1 option, 3657
 - E2 option, 3657
 - E3 option, 3657
 - E4 option, 3658
 - EFFECTSIZE option, 3658

- I option, 3658
- INTERCEPT option, 3658
- INVERSE option, 3658
- NOINT option, 3658
- NOUNI option, 3658
- P option, 3658
- SINGULAR= option, 3658
- SOLUTION option, 3659
- SS1 option, 3659
- SS2 option, 3659
- SS3 option, 3659
- SS4 option, 3659
- TOLERANCE option, 3659
- XPX option, 3659
- ZETA= option, 3659
- GLM procedure, OUTPUT statement, 3660
 - ALPHA= option, 3662
 - COOKD keyword, 3660
 - COVRATIO keyword, 3660
 - DFFITS keyword, 3660
 - H keyword, 3660
 - keyword= option, 3660
 - LCL keyword, 3660
 - LCLM keyword, 3661
 - OUT= option, 3662
 - PREDICTED keyword, 3661
 - PRESS keyword, 3661
 - RESIDUAL keyword, 3661
 - RSTUDENT keyword, 3661
 - STDI keyword, 3661
 - STDP keyword, 3661
 - STDR keyword, 3661
 - STUDENT keyword, 3661
 - UCL keyword, 3661
 - UCLM keyword, 3661
- GLM procedure, PROC GLM statement, 3624
 - ALPHA= option, 3625
 - DATA= option, 3625
 - MANOVA option, 3625
 - MULTIPASS option, 3625
 - NAMELEN= option, 3625
 - NOPRINT option, 3625
 - ORDER= option, 3625, 3691
 - OUTSTAT= option, 3626
 - PLOTS= option, 3626
- GLM procedure, RANDOM statement, 3662
 - Q option, 3663, 3721
 - TEST option, 3663
- GLM procedure, REPEATED statement, 3664
 - CANONICAL option, 3666
 - CONTRAST option, 3666, 3717
 - E= option, 3669
 - factor specification, 3665
 - H= option, 3669
 - HELMERT option, 3666, 3718
 - HTYPE= option, 3666
 - IDENTITY option, 3666, 3772
 - MEAN option, 3666, 3719
 - MSTAT= option, 3666
 - NOM option, 3667
 - NOU option, 3667
 - POLYNOMIAL option, 3666, 3718, 3766
 - PRINTE option, 3667, 3715
 - PRINTH option, 3667
 - PRINTM option, 3667
 - PRINTRV option, 3667
 - PROFILE option, 3666, 3719
 - SUMMARY option, 3667
 - UEPSDEF option, 3667
- GLM procedure, STORE statement, 3668
- GLM procedure, TEST statement, 3668
 - ETYPE= option, 3669
 - HTYPE= option, 3669
- GLM procedure, WEIGHT statement, 3669
- GT2 option
 - MEANS statement (GLM), 3653
- H keyword
 - OUTPUT statement (GLM), 3660
- H= option
 - MANOVA statement (GLM), 3646
 - REPEATED statement (GLM), 3669
- HELMERT option
 - REPEATED statement (GLM), 3666, 3718
- HOVTEST option
 - MEANS statement (GLM), 3653, 3706
- HTYPE= option
 - MANOVA statement (GLM), 3648
 - MEANS statement (GLM), 3654
 - REPEATED statement (GLM), 3666
 - TEST statement (GLM), 3669
- I option
 - MODEL statement (GLM), 3658
- ID statement
 - GLM procedure, 3637
- IDENTITY option
 - REPEATED statement (GLM), 3666, 3772
- INTERCEPT effect
 - CONTRAST statement (GLM), 3633, 3635
 - MANOVA statement, H= option (GLM), 3646
- INTERCEPT option
 - MODEL statement (GLM), 3658
- INVERSE option
 - MODEL statement (GLM), 3658
- keyword= option
 - OUTPUT statement (GLM), 3660
- KRATIO= option

- MEANS statement (GLM), 3654
- LCL keyword
 - OUTPUT statement (GLM), 3660
- LCLM keyword
 - OUTPUT statement (GLM), 3661
- LINES option
 - LSMEANS statement (GLM), 3641
 - MEANS statement (GLM), 3654
- LSD option
 - MEANS statement (GLM), 3654
- LSMEANS statement
 - GLM procedure, 3637
- M= option
 - MANOVA statement (GLM), 3647
- MANOVA option
 - PROC GLM statement, 3625
- MANOVA statement
 - GLM procedure, 3645
- MEAN option
 - REPEATED statement (GLM), 3666, 3719
- MEANS statement
 - GLM procedure, 3650
- MNAMES= option
 - MANOVA statement (GLM), 3647
- MODEL statement
 - GLM procedure, 3655
- MSTAT= option
 - MANOVA statement (GLM), 3648
 - REPEATED statement (GLM), 3666
- MULTIPASS option
 - PROC GLM statement, 3625
- NAMELEN= option
 - PROC GLM statement, 3625
- NOINT option
 - MODEL statement (GLM), 3658
- NOM option
 - REPEATED statement (GLM), 3667
- NOPRINT option
 - LSMEANS statement (GLM), 3641
 - PROC GLM statement, 3625
- NOSORT option
 - MEANS statement (GLM), 3654
- NOU option
 - REPEATED statement (GLM), 3667
- NOUNI option
 - MODEL statement (GLM), 3658
- OBSMARGINS option
 - LSMEANS statement (GLM), 3641, 3709
- OM option
 - LSMEANS statement (GLM), 3641, 3709
- ORDER= option
 - PROC GLM statement, 3625, 3691
- ORTH option
 - MANOVA statement (GLM), 3648
- OUT= option
 - LSMEANS statement (GLM), 3641
 - OUTPUT statement (GLM), 3662
- OUTPUT statement
 - GLM procedure, 3660
- OUTSTAT= option
 - PROC GLM statement, 3626
- P option
 - MODEL statement (GLM), 3658
- PDIF= option
 - LSMEANS statement (GLM), 3641
- PLOTS= option
 - LSMEANS statement (GLM), 3643
 - PROC GLM statement, 3626
- POLYNOMIAL option
 - REPEATED statement (GLM), 3666, 3718, 3766
- PREDICTED keyword
 - OUTPUT statement (GLM), 3661
- PREFIX= option
 - MANOVA statement (GLM), 3647
- PRESS keyword
 - OUTPUT statement (GLM), 3661
- PRINTE option
 - MANOVA statement (GLM), 3648
 - REPEATED statement (GLM), 3667, 3715
- PRINTH option
 - MANOVA statement (GLM), 3648
 - REPEATED statement (GLM), 3667
- PRINTM option
 - REPEATED statement (GLM), 3667
- PRINTRV option
 - REPEATED statement (GLM), 3667
- PROC GLM statement, *see* GLM procedure
- PROFILE option
 - REPEATED statement (GLM), 3666, 3719
- Q option
 - RANDOM statement (GLM), 3663, 3721
- RANDOM statement
 - GLM procedure, 3662
- REF= option
 - CLASS statement (GLM), 3632
- REGWQ option
 - MEANS statement (GLM), 3654
- REPEATED statement
 - GLM procedure, 3664
- RESIDUAL keyword
 - OUTPUT statement (GLM), 3661
- RSTUDENT keyword
 - OUTPUT statement (GLM), 3661

- SCHEFFE option
 - MEANS statement (GLM), 3654
- SIDAK option
 - MEANS statement (GLM), 3654
- SINGULAR option
 - CONTRAST statement (GLM), 3634
 - LSMEANS statement (GLM), 3645
- SINGULAR= option
 - ESTIMATE statement (GLM), 3636
 - MODEL statement (GLM), 3658
- SLICE= option
 - LSMEANS statement (GLM), 3645
- SMM option
 - MEANS statement (GLM), 3655
- SNK option
 - MEANS statement (GLM), 3655
- SOLUTION option
 - MODEL statement (GLM), 3659
- SS1 option
 - MODEL statement (GLM), 3659
- SS2 option
 - MODEL statement (GLM), 3659
- SS3 option
 - MODEL statement (GLM), 3659
- SS4 option
 - MODEL statement (GLM), 3659
- STDERR option
 - LSMEANS statement (GLM), 3645
- STDI keyword
 - OUTPUT statement (GLM), 3661
- STDP keyword
 - OUTPUT statement (GLM), 3661
- STDR keyword
 - OUTPUT statement (GLM), 3661
- STORE statement
 - GLM procedure, 3668
- STUDENT keyword
 - OUTPUT statement (GLM), 3661
- SUMMARY option
 - MANOVA statement (GLM), 3648
 - REPEATED statement (GLM), 3667
- T option
 - MEANS statement (GLM), 3655
- TDIFF option
 - LSMEANS statement (GLM), 3645
- TEST option
 - RANDOM statement (GLM), 3663
- TEST statement
 - GLM procedure, 3668
- TOLERANCE option
 - MODEL statement (GLM), 3659
- TRUNCATE option
 - CLASS statement (GLM), 3632

- TUKEY option
 - MEANS statement (GLM), 3655
- UCL keyword
 - OUTPUT statement (GLM), 3661
- UCLM keyword
 - OUTPUT statement (GLM), 3661
- UEPSDEF option
 - REPEATED statement (GLM), 3667
- WALLER option
 - MEANS statement (GLM), 3655
- WEIGHT statement
 - GLM procedure, 3669
- WELCH option
 - MEANS statement (GLM), 3655
- WHERE statement
 - GLM procedure, 3673
- XPX option
 - MODEL statement (GLM), 3659
- ZETA= option
 - MODEL statement (GLM), 3659