

SAS/STAT[®] 14.2 User's Guide

The HPLMIXED

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 54

The HPLMIXED Procedure

Contents

Overview: HPLMIXED Procedure	4272
PROC HPLMIXED Features	4273
Notation for the Mixed Model	4274
PROC HPLMIXED Contrasted with Other SAS Procedures	4275
Getting Started: HPLMIXED Procedure	4275
Mixed Model Analysis of Covariance with Many Groups	4275
Syntax: HPLMIXED Procedure	4278
PROC HPLMIXED Statement	4279
CLASS Statement	4284
ID Statement	4285
MODEL Statement	4285
OUTPUT Statement	4286
PARMS Statement	4287
PERFORMANCE Statement	4289
RANDOM Statement	4289
REPEATED Statement	4296
Details: HPLMIXED Procedure	4297
Linear Mixed Models Theory	4297
Matrix Notation	4297
Formulation of the Mixed Model	4298
Estimating Covariance Parameters in the Mixed Model	4302
Estimating Fixed and Random Effects in the Mixed Model	4303
Statistical Properties	4304
Computational Method	4305
Distributed Computing	4305
Multithreading	4305
Displayed Output	4306
Performance Information	4306
Model Information	4306
Class Level Information	4306
Dimensions	4306
Number of Observations	4306
Optimization Information	4307
Iteration History	4307
Convergence Status	4308
Covariance Parameter Estimates	4308

Fit Statistics	4308
Timing Information	4309
ODS Table Names	4309
Examples: HPLMIXED Procedure	4310
Example 54.1: Computing BLUPs for a Large Number of Subjects	4310
References	4312

Overview: HPLMIXED Procedure

The HPLMIXED procedure fits a variety of mixed linear models to data and enables you to use these fitted models to make statistical inferences about the data. A *mixed linear model* is a generalization of the standard linear model used in the GLM procedure in SAS/STAT software; the generalization is that the data are permitted to exhibit correlation and nonconstant variability. Therefore, the mixed linear model provides you with the flexibility of modeling not only the means of your data (as in the standard linear model) but also their variances and covariances.

The primary assumptions underlying the analyses performed by PROC HPLMIXED are as follows:

- The data are normally distributed (Gaussian).
- The means (expected values) of the data are linear in terms of a certain set of parameters.
- The variances and covariances of the data are in terms of a different set of parameters, and they exhibit a structure that matches one of those available in PROC HPLMIXED.

Because Gaussian data can be modeled entirely in terms of their means and variances/covariances, the two sets of parameters in a mixed linear model actually specify the complete probability distribution of the data. The parameters of the mean model are referred to as *fixed-effects parameters*, and the parameters of the variance-covariance model are referred to as *covariance parameters*.

The fixed-effects parameters are associated with known explanatory variables, as in the standard linear model. These variables can be either qualitative (as in the traditional analysis of variance) or quantitative (as in standard linear regression). However, the covariance parameters are what distinguishes the mixed linear model from the standard linear model.

The need for covariance parameters arises quite frequently in applications; the following scenarios are the most typical:

- The experimental units on which the data are measured can be grouped into clusters, and the data from a common cluster are correlated. This scenario can be generalized to include one set of clusters nested within another. For example, if students are the experimental unit, they can be clustered into classes, which in turn can be clustered into schools. Each level of this hierarchy can introduce an additional source of variability and correlation.
- Repeated measurements are taken on the same experimental unit, and these repeated measurements are correlated or exhibit variability that changes. This scenario occurs in longitudinal studies, where repeated measurements are taken over time. Alternatively, the repeated measures could be spatial or multivariate in nature.

PROC HPLMIXED provides a variety of covariance structures to handle these two scenarios. The most common covariance structures arise from the use of *random-effects parameters*, which are additional unknown random variables that are assumed to affect the variability of the data. The variances of the random-effects parameters, commonly known as *variance components*, become the covariance parameters for this particular structure. Traditional mixed linear models contain both fixed- and random-effects parameters; in fact, it is the combination of these two types of effects that led to the name *mixed model*. PROC HPLMIXED fits not only these traditional variance component models but also numerous other covariance structures.

PROC HPLMIXED fits the structure you select to the data by using the method of *restricted maximum likelihood (REML)*, also known as *residual maximum likelihood*. It is here that the Gaussian assumption for the data is exploited.

PROC HPLMIXED runs in either single-machine mode or distributed mode.

NOTE: Distributed mode requires SAS High-Performance Statistics.

PROC HPLMIXED Features

PROC HPLMIXED provides easy accessibility to numerous mixed linear models that are useful in many common statistical analyses.

Here are some basic features of PROC HPLMIXED:

- covariance structures, including variance components, compound symmetry, unstructured, AR(1), Toeplitz, and factor analytic
- **MODEL**, **RANDOM**, and **REPEATED** statements for model specification as in the HPLMIXED procedure
- appropriate standard errors, *t* tests, and *F* tests for all specified estimable linear combinations of fixed and random effects
- a subject effect that enables blocking
- REML and ML (maximum likelihood) estimation methods implemented with a variety of optimization algorithms
- capacity to handle unbalanced data
- special dense and sparse algorithms that take advantage of distributed and multicore computing environments

Because the HPLMIXED procedure is a high-performance analytical procedure, it also does the following:

- enables you to run in distributed mode on a cluster of machines that distribute the data and the computations
- enables you to run in single-machine mode on the server where SAS is installed
- exploits all the available cores and concurrent threads, regardless of execution mode

For more information, see the section “Processing Modes” (Chapter 2, *SAS/STAT User’s Guide: High-Performance Procedures*).

PROC HPLMIXED uses the Output Delivery System (ODS), a SAS subsystem that provides capabilities for displaying and controlling the output from SAS procedures. ODS enables you to convert any output from PROC HPLMIXED into a SAS data set. See the section “ODS Table Names” on page 4309.

Notation for the Mixed Model

This section introduces the mathematical notation used throughout this chapter to describe the mixed linear model and assumes familiarity with basic matrix algebra (for an overview, see Searle 1982). A more detailed description of the mixed model is contained in the section “Linear Mixed Models Theory” on page 4297.

A statistical model is a mathematical description of how data are generated. The standard linear model, as used by the GLM procedure, is one of the most common statistical models:

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

In this expression, $\mathbf{y}$ represents a vector of observed data, $\boldsymbol{\beta}$ is an unknown vector of fixed-effects parameters with a known design matrix $\mathbf{X}$, and $\boldsymbol{\epsilon}$ is an unknown random error vector that models the statistical noise around $\mathbf{X}\boldsymbol{\beta}$. The focus of the standard linear model is to model the mean of $\mathbf{y}$ by using the fixed-effects parameters $\boldsymbol{\beta}$. The residual errors $\boldsymbol{\epsilon}$ are assumed to be independent and identically distributed Gaussian random variables with mean 0 and variance σ^2 .

The mixed model generalizes the standard linear model as follows:

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

Here, $\boldsymbol{\gamma}$ is an unknown vector of random-effects parameters with a known design matrix $\mathbf{Z}$, and $\boldsymbol{\epsilon}$ is an unknown random error vector whose elements are no longer required to be independent and homogeneous.

To further develop this notion of variance modeling, assume that $\boldsymbol{\gamma}$ and $\boldsymbol{\epsilon}$ are Gaussian random variables that are uncorrelated, have expectations $\mathbf{0}$, and have variances $\mathbf{G}$ and $\mathbf{R}$, respectively. The variance of $\mathbf{y}$ is thus

$$\mathbf{V} = \mathbf{Z}\mathbf{G}\mathbf{Z}' + \mathbf{R}$$

Note that when $\mathbf{R} = \sigma^2\mathbf{I}$ and $\mathbf{Z} = \mathbf{0}$, the mixed model reduces to the standard linear model.

You can model the variance of the data $\mathbf{y}$ by specifying the structure of $\mathbf{Z}$, $\mathbf{G}$, and $\mathbf{R}$. The model matrix $\mathbf{Z}$ is set up in the same fashion as $\mathbf{X}$, the model matrix for the fixed-effects parameters. For $\mathbf{G}$ and $\mathbf{R}$, you must select some covariance structure. Possible covariance structures include the following:

- variance components
- compound symmetry (common covariance plus diagonal)
- unstructured (general covariance)
- autoregressive
- spatial

- general linear
- factor analytic

By appropriately defining the model matrices **X** and **Z** in addition to the covariance structure matrices **G** and **R**, you can perform numerous mixed model analyses.

PROC HPLMIXED Contrasted with Other SAS Procedures

The **RANDOM** and **REPEATED** statements of the HPLMIXED procedure follow the convention of the same statements in the MIXED procedure in SAS/STAT software. For information about how these statements differ from **RANDOM** and **REPEATED** statements in the MIXED procedure, see the documentation for the MIXED procedure in the *SAS/STAT User's Guide*.

The GLIMMIX procedure in SAS/STAT software fits generalized linear mixed models. Linear mixed models—where the data are normally distributed, given the random effects—are in the class of generalized linear mixed models. Therefore, PROC GLIMMIX accommodates nonnormal data with random effects.

Generalized linear mixed models have intrinsically nonlinear features because a nonlinear mapping (the link function) connects the conditional mean of the data (given the random effects) to the explanatory variables. The NLMIXED procedure also accommodates nonlinear structures in the conditional mean, but places no restrictions on the nature of the nonlinearity.

The HPMIXED procedure in SAS/STAT software is also termed a “high-performance” procedure, but it does not follow the general pattern of high-performance analytical procedures. The HPMIXED procedure does not take advantage of distributed or multicore computing environments; it derives high performance from applying sparse techniques to solving the mixed model equations. The HPMIXED procedure fits a small subset of the statistical models you can fit with the MIXED or HPLMIXED procedures and is particularly suited for problems in which the $[XZ]'[XZ]$ crossproducts matrix is sparse.

The HPLMIXED procedure employs algorithms that are specialized for distributed and multicore computing environments. The HPLMIXED procedure does not support BY processing.

Getting Started: HPLMIXED Procedure

Mixed Model Analysis of Covariance with Many Groups

Suppose you are an educational researcher who studies how student scores on math tests change over time. Students are tested four times, and you want to estimate the overall rise or fall, accounting for correlation between test response behaviors of students in the same neighborhood and school. One way to model this correlation is by using a random-effects analysis of covariance, where the scores for students from the same neighborhood and school are all assumed to share the same quadratic mean test response function, the parameters of this response function being random. The following statements simulate a data set with this structure:

```

data SchoolSample;
  do SchoolID = 1 to 300;
    do nID = 1 to 25;
      Neighborhood = (SchoolID-1)*5 + nID;
      bInt    = 5*ranuni(1);
      bTime   = 5*ranuni(1);
      bTime2  =   ranuni(1);
      do sID = 1 to 2;
        do Time = 1 to 4;
          Math = bInt + bTime*Time + bTime2*Time*Time + rannor(2);
          output;
        end;
      end;
    end;
  end;
run;

```

In this data, there are 300 schools and about 1,500 neighborhoods; neighborhoods are associated with more than one school and vice versa. The following statements use PROC HPLMIXED to fit a mixed analysis of covariance model to this data. To run these statements successfully, you need to set the macro variables GRIDHOST and GRIDINSTALLLOC to resolve to appropriate values, or you can replace the references to macro variables with appropriate values.

```

proc hplmixed data=SchoolSample;
  performance host="&GRIDHOST" install="&GRIDINSTALLLOC" nodes=20 nthreads=4;
  class Neighborhood SchoolID;
  model Math = Time Time*Time / solution;
  random int Time Time*Time / sub=Neighborhood(SchoolID) type=un;
run;

```

This model fits a quadratic mean response model with an unstructured covariance matrix to model the covariance between the random parameters of the response model. With 7,500 neighborhood/school combinations, this model can be computationally daunting to fit, but PROC HPLMIXED finishes quickly and displays the results shown in [Figure 54.1](#).

Figure 54.1 Mixed Model Analysis of Covariance

Performance Information			
Host Node	<< your grid host >>		
Install Location	<< your grid install location >>		
Execution Mode	Distributed		
Number of Compute Nodes	20		
Number of Threads per Node	4		
Data Access Information			
Data	Engine	Role	Path
WORK.SCHOOLSAMPLE	V9	Input	From Client

Figure 54.1 continued

Model Information	
Data Set	WORK.SCHOOLSAMPLE
Dependent Variable	Math
Covariance Structure	Unstructured
Subject Effect	Neighborho(SchoolID)
Estimation Method	Restricted Maximum Likelihood
Residual Variance Method	Profile
Fixed Effects SE Method	Model-Based
Degrees of Freedom Method	Residual

Class Level Information	
Class	Levels Values
Neighborhood	1520 1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20 ...
SchoolID	300 1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20 ...

Dimensions	
Covariance Parameters	7
Columns in X	3
Columns in Z per Subject	3
Subjects	7500
Max Obs per Subject	8

Number of Observations Read	60000
Number of Observations Used	60000
Number of Observations Not Used	0

Optimization Information	
Optimization Technique	Newton-Raphson with Ridging
Parameters in Optimization	6
Lower Boundaries	3
Upper Boundaries	0
Starting Values From	Data

Iteration History			
Iteration	Evaluations	Objective Function Change	Max Gradient
0	2	225641.67142	3.369E-8

Convergence criterion (ABSGCONV=0.00001) satisfied.			
---	--	--	--

Figure 54.1 continued

Covariance Parameter Estimates					
Cov Parm	Subject	Estimate			
UN(1,1)	Neighborho(SchoolID)	2.0902			
UN(2,1)	Neighborho(SchoolID)	0.000349			
UN(2,2)	Neighborho(SchoolID)	2.0517			
UN(3,1)	Neighborho(SchoolID)	0.01448			
UN(3,2)	Neighborho(SchoolID)	0.01599			
UN(3,3)	Neighborho(SchoolID)	0.08047			
Residual		1.0083			

Fit Statistics		
-2 Res Log Likelihood		225642
AIC (Smaller is Better)		225656
AICC (Smaller is Better)		225656
BIC (Smaller is Better)		225704

Solution for Fixed Effects					
Effect	Estimate	Standard Error	DF	t Value	Pr > t
Intercept	2.5070	0.02828	6E4	88.66	<.0001
Time	2.5124	0.02659	6E4	94.48	<.0001
Time*Time	0.5010	0.005247	6E4	95.48	<.0001

Syntax: HPLMIXED Procedure

The following statements are available in PROC HPLMIXED.

```

PROC HPLMIXED < options > ;
  CLASS variables ;
  ID variables ;
  MODEL dependent = < fixed-effects > < / options > ;
  OUTPUT < OUT=SAS-data-set > < keyword < =name > > . . . < keyword < =name > > < / options > ;
  RANDOM random-effects < / options > ;
  REPEATED repeated-effect < / options > ;
  PARMS < (value-list) . . . > < / options > ;
  PERFORMANCE < options > ;

```

Items within angle brackets (< >) are optional. The **RANDOM** statement can appear multiple times. Other statements can appear only once.

The **PROC HPLMIXED** and **MODEL** statements are required, and the **MODEL** statement must appear after the **CLASS** statement if a **CLASS** statement is included. The **RANDOM** statement must follow the **MODEL** statement.

Table 54.1 summarizes the basic functions and important options of the PROC HPLMIXED statements. The syntax of each statement in Table 54.1 is described in the following sections in alphabetical order after the description of the **PROC HPLMIXED** statement.

Table 54.1 Summary of PROC HPLMIXED Statements

Statement	Description	Important Options
PROC HPLMIXED	Invokes the procedure	DATA= specifies the input data set; METHOD= specifies the estimation method.
CLASS	Declares qualitative variables that create indicator variables in X and Z matrices.	None
ID	Lists additional variables to be included in predicted values tables	None
MODEL	Specifies dependent variable and fixed effects, setting up X	S requests a solution for fixed-effects parameters.
RANDOM	Specifies random effects, setting up Z and G	SUBJECT= creates block-diagonality; TYPE= specifies the covariance structure; S requests a solution for the random effects.
REPEATED	Sets up R	SUBJECT= creates block-diagonality; TYPE= specifies the covariance structure.
OUTPUT	Creates a data set that contains observationwise statistics	ALLSTATS requests that all statistics be computed.
PARMS	Specifies a grid of initial values for the covariance parameters	HOLD= and NOITER hold the covariance parameters or their ratios constant; PARMSDATA= reads the initial values from a SAS data set.
PERFORMANCE	Invokes the distributed computing connection	NODES= specifies the number of nodes to use.

PROC HPLMIXED Statement

PROC HPLMIXED < options > ;

The PROC HPLMIXED statement invokes the procedure. Table 54.2 summarizes important options in the PROC HPLMIXED statement by function. These and other options in the PROC HPLMIXED statement are then described fully in alphabetical order.

Table 54.2 PROC HPLMIXED Statement Options

Option	Description
Basic Options	
DATA=	Specifies the input data set
METHOD=	Specifies the estimation method

Table 54.2 *continued*

Option	Description
NAMELEN=	Limits the length of effect names
BLUP	Computes the best linear unbiased prediction
Options Related to Output	
NOCLPRINT	Suppresses the “Class Level Information” table completely or in parts
MAXCLPRINT=	Specifies the maximum levels of CLASS variables to print
RANKS	Displays the rank of design matrix X
Optimization Options	
ABSCONV=	Tunes an absolute function convergence criterion
ABSFCNV=	Tunes an absolute function difference convergence criterion
ABSGCONV=	Tunes the absolute gradient convergence criterion
FCNV=	Tunes the relative function convergence criterion
GCONV=	Tunes the relative gradient convergence criterion
MAXITER=	Chooses the maximum number of iterations in any optimization
MAXFUNC=	Specifies the maximum number of function evaluations in any optimization
MAXTIME=	Specifies the upper limit on seconds of CPU time for any optimization
MINITER=	Specifies the minimum number of iterations in any optimization
TECHNIQUE=	Selects the optimization technique
XCONV=	Tunes the relative parameter convergence criterion

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

ABSCONV=*r*

specifies an absolute function convergence criterion. For minimization, termination requires $f(\boldsymbol{\psi}^{(k)}) \leq r$, where $\boldsymbol{\psi}$ is the vector of parameters in the optimization and $f(\cdot)$ is the objective function. The default value of r is the negative square root of the largest double-precision value, which serves only as a protection against overflows.

ABSFCNV=*r*

specifies an absolute function difference convergence criterion. For all techniques except Nelder–Mead simplex (NMSIMP), termination requires a small change of the function value in successive iterations:

$$|f(\boldsymbol{\psi}^{(k-1)}) - f(\boldsymbol{\psi}^{(k)})| \leq r$$

Here, $\boldsymbol{\psi}$ denotes the vector of parameters that participate in the optimization and $f(\cdot)$ is the objective function. The same formula is used for the NMSIMP technique, but $\boldsymbol{\psi}^{(k)}$ is defined as the vertex with the lowest function value and $\boldsymbol{\psi}^{(k-1)}$ is defined as the vertex with the highest function value in the simplex. The default value is $r = 0$.

ABSGCONV=*r*

specifies an absolute gradient convergence criterion. Termination requires the maximum absolute gradient element to be small:

$$\max_j |g_j(\boldsymbol{\psi}^{(k)})| \leq r$$

Here, $\boldsymbol{\psi}$ denotes the vector of parameters that participate in the optimization and $g_j(\cdot)$ is the gradient of the objective function with respect to the j parameter. This criterion is not used by the NMSIMP technique. The default value is $r=1\text{E-}5$.

BLUP<(suboptions)>

requests that best linear unbiased predictions (BLUPs) for the random effects be displayed. To use this option, you must also use the PARMS statement to specify fixed values for the covariance parameters, which means that the NOITER option in the PARMS statement will be implied by the BLUP option. Also, the iterations in the ODS Table IterHistory will refer to iterations used to compute the BLUP rather than optimization iterations.

The BLUP solution might be sensitive to the order of observations, and hence to how the data are distributed on the grid. If there are multiple measures of a repeated effect, then the BLUP solution is not unique. If the order of these multiple measures on the grid differs for different runs, then different BLUP solutions will result.

You can specify the following *suboptions*:

ITPRINT=number specifies that the iteration history be displayed after every *number* of iterations. The default value is 10, which means the procedure displays the iteration history for every 10 iterations.

MAXITER=number specifies the maximum number of iterations allowed. The default value is the number of parameters in the BLUP option plus 2.

TOL=number specifies the tolerance value. The default value is the square root of machine precision.

DATA=SAS-data-set

names the SAS data set to be used as the input data set. The default is the most recently created data set.

FCONV=*r*

specifies a relative function convergence criterion. For all techniques except NMSIMP, termination requires a small relative change of the function value in successive iterations,

$$\frac{|f(\boldsymbol{\psi}^{(k)}) - f(\boldsymbol{\psi}^{(k-1)})|}{|f(\boldsymbol{\psi}^{(k-1)})|} \leq r$$

Here, $\boldsymbol{\psi}$ denotes the vector of parameters that participate in the optimization and $f(\cdot)$ is the objective function. The same formula is used for the NMSIMP technique, but $\boldsymbol{\psi}^{(k)}$ is defined as the vertex with the lowest function value and $\boldsymbol{\psi}^{(k-1)}$ is defined as the vertex with the highest function value in the simplex.

The default is $r = 10^{-\text{FDIGITS}}$, where FDIGITS is $-\log_{10}(\epsilon)$ and ϵ is the machine precision.

GCONV=*r*

specifies a relative gradient convergence criterion. For all techniques except CONGRA and NMSIMP, termination requires that the normalized predicted function reduction be small,

$$\frac{\mathbf{g}(\boldsymbol{\psi}^{(k)})' [\mathbf{H}^{(k)}]^{-1} \mathbf{g}(\boldsymbol{\psi}^{(k)})}{|f(\boldsymbol{\psi}^{(k)})|} \leq r$$

Here, $\boldsymbol{\psi}$ denotes the vector of parameters that participate in the optimization, $f(\cdot)$ is the objective function, and $\mathbf{g}(\cdot)$ is the gradient. For the CONGRA technique (where a reliable Hessian estimate $\mathbf{H}$ is not available), the following criterion is used:

$$\frac{\|\mathbf{g}(\boldsymbol{\psi}^{(k)})\|_2^2 \|\mathbf{s}(\boldsymbol{\psi}^{(k)})\|_2}{\|\mathbf{g}(\boldsymbol{\psi}^{(k)}) - \mathbf{g}(\boldsymbol{\psi}^{(k-1)})\|_2 |f(\boldsymbol{\psi}^{(k)})|} \leq r$$

This criterion is not used by the NMSIMP technique. The default value is $r=1\text{E-}8$.

MAXCLPRINT=*number*

specifies the maximum levels of CLASS variables to print in the ODS table “ClassLevels.” The default value is 20. MAXCLPRINT=0 enables you to print all levels of each CLASS variable. However, the option **NOCLPRINT** takes precedence over MAXCLPRINT.

MAXFUNC=*n*

specifies the maximum number n of function calls in the optimization process. The default values are as follows, depending on the optimization technique:

- TRUREG, NRRIDG, NEWRAP: 125
- QUANEW, DBLDOG: 500
- CONGRA: 1,000
- NMSIMP: 3,000

The optimization can terminate only after completing a full iteration. Therefore, the number of function calls that are actually performed can exceed n . You can choose the optimization technique with the **TECHNIQUE=** option.

MAXITER=*n*

specifies the maximum number n of iterations in the optimization process. The default values are as follows, depending on the optimization technique:

- TRUREG, NRRIDG, NEWRAP: 50
- QUANEW, DBLDOG: 200
- CONGRA: 400
- NMSIMP: 1,000

These default values also apply when n is specified as a missing value. You can choose the optimization technique with the **TECHNIQUE=** option.

MAXTIME=*r*

specifies an upper limit of *r* seconds of CPU time for the optimization process. The default value is the largest floating-point double representation of your computer. The time specified by the MAXTIME= option is checked only once at the end of each iteration. Therefore, the actual running time can be longer than *r*.

METHOD=REML**METHOD=ML**

specifies the estimation method for the covariance parameters. METHOD=REML performs residual (restricted) maximum likelihood; it is the default method. METHOD=ML performs maximum likelihood.

MINITER=*n*

specifies the minimum number of iterations. The default value is 0. If you request more iterations than are actually needed for convergence to a stationary point, the optimization algorithms can behave strangely. For example, the effect of rounding errors can prevent the algorithm from continuing for the required number of iterations.

NAMELEN=*number*

specifies the length to which long effect names are shortened. The minimum value is 20, which is also the default.

NOCLPRINT<=*number*>

suppresses the display of the “Class Level Information” table if you do not specify *number*. If you specify *number*, the values of the classification variables are displayed for only those variables whose number of levels is less than *number*. Specifying a *number* helps to reduce the size of the “Class Level Information” table if some classification variables have a large number of levels.

NOPRINT

suppresses the generation of ODS output.

RANKS

displays the rank of design matrix **X**.

SINGCHOL=*number*

tunes the singularity criterion in Cholesky decompositions. The default is 1E4 times the machine epsilon; this product is approximately 1E–12 on most computers.

SINGSWEEP=*number*

tunes the singularity criterion for sweep operations. The default is 1E4 times the machine epsilon; this product is approximately 1E–12 on most computers.

SINGULAR=*number*

tunes the general singularity criterion applied by the HPLMIXED procedure in sweeps and inversions. The default is 1E4 times the machine epsilon; this product is approximately 1E–12 on most computers.

TECHNIQUE=*keyword*

specifies the optimization technique for obtaining maximum likelihood estimates. You can specify any of the following *keywords*:

CONGRA	performs a conjugate-gradient optimization.
DBLDOG	performs a version of double-dogleg optimization.
NEWRAP	performs a Newton-Raphson optimization combining a line-search algorithm with ridging.
NMSIMP	performs a Nelder-Mead simplex optimization.
NONE	performs no optimization.
NRRIDG	performs a Newton-Raphson optimization with ridging.
QUANEW	performs a dual quasi-Newton optimization.
TRUREG	performs a trust-region optimization.

The default value is `TECHNIQUE=NRRIDG`.

XCONV=r

specifies the relative parameter convergence criterion:

- For all techniques except NMSIMP, termination requires a small relative parameter change in subsequent iterations:

$$\frac{\max_j |\psi_j^{(k)} - \psi_j^{(k-1)}|}{\max(|\psi_j^{(k)}|, |\psi_j^{(k-1)}|)} \leq r$$

- For the NMSIMP technique, the same formula is used, but $\psi_j^{(k)}$ is defined as the vertex with the lowest function value and $\psi_j^{(k-1)}$ is defined as the vertex with the highest function value in the simplex.

The default value is $r = 1\text{E-}8$ for the NMSIMP technique and $r = 0$ otherwise.

CLASS Statement

CLASS *variables* ;

The CLASS statement names the classification variables to be used as explanatory variables in the analysis. These variables enter the analysis not through their values, but through levels to which the unique values are mapped. For more information about these mappings, see the section “Levelization of Classification Variables” (Chapter 3, *SAS/STAT User’s Guide: High-Performance Procedures*).

If a CLASS statement is specified, it must precede the MODEL statement in high-performance analytical procedures that support a MODEL statement.

Levels of classification variables are ordered by their external formatted values, except for numeric variables with no explicit format, which are ordered by their unformatted (internal) values.

ID Statement

ID *variables* ;

The ID statement specifies which variables from the input data set are to be included in the OUT= data set that is created by the **OUTPUT** statement. If you do not specify an ID statement, then no variables from the input data set are included in the OUT= data set.

MODEL Statement

MODEL *dependent* = < *fixed-effects* > < / *options* > ;

The MODEL statement names a single dependent variable and the fixed effects, which determine the X matrix of the mixed model. (For more information, see the section “Specification and Parameterization of Model Effects” (Chapter 3, *SAS/STAT User’s Guide: High-Performance Procedures*). The MODEL statement is required.

An intercept is included in the fixed-effects model by default. If no fixed effects are specified, only this intercept term is fit. The intercept can be removed by using the NOINT option.

Table 54.3 summarizes options in the MODEL statement. These are subsequently discussed in detail in alphabetical order.

Table 54.3 Summary of Important MODEL Statement Options

Option	Description
Model Building	
NOINT	Excludes the fixed-effect intercept from model
Statistical Computations	
ALPHA= α	Determines the confidence level $(1 - \alpha)$ for fixed effects
DDFM=	Specifies the method for computing denominator degrees of freedom
Statistical Output	
CL	Displays confidence limits for fixed-effects parameter estimates
SOLUTION	Displays fixed-effects parameter estimates

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

ALPHA=*number*

sets the confidence level to be $1 - \text{number}$ for each confidence interval of the fixed-effects parameters. The value of *number* must be between 0 and 1; the default is 0.05.

CL

requests that t -type confidence limits be constructed for each of the fixed-effects parameter estimates. The confidence level is 0.95 by default; this can be changed with the **ALPHA=** option.

DDFM=NONE | RESIDUAL

specifies the method for computing the denominator degrees of freedom for the tests of fixed effects.

The **DDFM=RESIDUAL** option performs all tests by using the residual degrees of freedom, $n - \text{rank}(\mathbf{X})$, where n is the number of observations used. It is the default degrees-of-freedom method.

DDFM=NONE specifies that no denominator degrees of freedom be applied. PROC HPLMIXED then essentially assumes that infinite degrees of freedom are available in the calculation of p -values. The p -values for t tests are then identical to p -values that are derived from the standard normal distribution. In the case of F tests, the p -values equal those of chi-square tests determined as follows: if F_{obs} is the observed value of the F test with l numerator degrees of freedom, then

$$p = \Pr\{F_{l,\infty} > F_{obs}\} = \Pr\{\chi_l^2 > lF_{obs}\}$$

NOINT

requests that no intercept be included in the model. (An intercept is included by default.)

SOLUTION**S**

requests that a solution for the fixed-effects parameters be produced. Using notation from the section “[Linear Mixed Models Theory](#)” on page 4297, the fixed-effects parameter estimates are $\hat{\beta}$ and their approximate standard errors are the square roots of the diagonal elements of $(\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{X})^{-1}$.

Along with the estimates and their approximate standard errors, a t statistic is computed as the estimate divided by its standard error. The **Pr > |t|** column contains the two-tailed p -value that corresponds to the t statistic and associated degrees of freedom. You can use the **CL** option to request confidence intervals for all of the parameters; they are constructed around the estimate by using a radius that is the product of the standard error times a percentage point from the t distribution.

OUTPUT Statement

OUTPUT < **OUT=***SAS-data-set* > < *keyword* <=*name*> >... < *keyword* <=*name*> > < / *options* > ;

The **OUTPUT** statement creates a data set that contains predicted values and residual diagnostics, which are computed after the model is fit. The variables in the input data set are *not* included in the output data set, in order to avoid data duplication for large data sets; however, variables that are specified in the **ID statement** are included. By default, only predicted values are included in the output data set.

If the input data are in distributed form, in which access of data in a particular order cannot be guaranteed, the HPLMIXED procedure copies the distribution or partition key to the output data set so that its contents can be joined with the input data.

For example, suppose that the data set **Scores** contains the variables **Score**, **Machine**, and **Person**. The following statements fit a model that has fixed machine and random person effects and save the predicted and residual values to the data set **lgausOut**:

```
proc hplmixed data = Scores;
  class machine person score;
  model score = machine;
  random person;
  output out=igausout pred=p resid=r;
run;
```

You can specify the following syntax element in the OUTPUT statement:

OUT=SAS-data-set

specifies the name of the output data set. If the OUT= option is omitted, PROC HPLMIXED uses the DATA n convention to name the output data set.

A *keyword* can appear multiple times in the OUTPUT statement. Table 54.4 lists the keywords and the default names that PROC HPLMIXED assigns if you do not specify a *name*. In this table, y denotes the response variable.

Table 54.4 Keywords for Output Statistics

Keyword	Description	Expression	Name
PRED	Linear predictor	$\hat{\eta} = \mathbf{x}'\hat{\boldsymbol{\beta}} + \mathbf{z}'\hat{\boldsymbol{\gamma}}$	Pred
PREDPA	Marginal linear predictor	$\hat{\eta}_m = \mathbf{x}'\hat{\boldsymbol{\beta}}$	PredPA
RESIDUAL	Residual	$r = y - \hat{\eta}$	Resid
RESIDUALPA	Marginal residual	$r_m = y - \hat{\eta}_m$	ResidPA

The marginal linear predictor and marginal residual are also referred to as the predicted population average (PREDPA) and residual population average (RESIDUALPA), respectively. You can use the following shortcuts to request statistics: PRED for PREDICTED and RESID for RESIDUAL.

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

ALLSTATS

requests that all statistics be computed. If you do not use a *keyword* to assign a name, PROC HPLMIXED uses the default name.

PARMS Statement

PARMS < (value-list) . . . > < / options > ;

The PARMS statement specifies initial values for the covariance parameters, or it requests a grid search over several values of these parameters. You must specify the values in the order in which they appear in the “Covariance Parameter Estimates” table.

The *value-list* specification can take any of several forms:

m	a single value
$m_1, m_2, \dots, m_n$	several values
m to n	a sequence in which m equals the starting value, n equals the ending value, and the increment equals 1
m to n by i	a sequence in which m equals the starting value, n equals the ending value, and the increment equals i
m_1, m_2 to m_3	mixed values and sequences

You can use the PARMS statement to input known parameters.

If you specify more than one set of initial values, PROC HPLMIXED performs a grid search of the likelihood surface and uses the best point on the grid for subsequent analysis. Specifying a large number of grid points can result in long computing times.

The results from the PARMS statement are the values of the parameters on the specified grid (denoted by CovP1 through CovPn), the residual variance (possibly estimated) for models with a residual variance parameter, and various functions of the likelihood.

If there are multiple PARMS statements, the first one is used and the rest are ignored.

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

HOLD=*all*

EQCONS=*all*

specifies that all parameter values be held to equal the specified values.

For example, the following statement constrains all covariance parameters to equal 5, 3, 2, and 3:

```
parms (5) (3) (2) (3) / hold=all;
```

LOWERB=*value-list*

enables you to specify lower boundary constraints on the covariance parameters. The *value-list* specification is a list of numbers or missing values (.) separated by commas. You must list the numbers in the order that PROC HPLMIXED uses for the covariance parameters, and each number corresponds to the lower boundary constraint. A missing value instructs PROC HPLMIXED to use its default constraint. If you do not specify numbers for all of the covariance parameters, PROC HPLMIXED assumes the remaining ones are missing.

This option is useful when you want to constrain the G matrix to be positive definite in order to avoid the more computationally intensive algorithms that would be required when G becomes singular. The corresponding statements for a random coefficients model are as follows:

```
proc hplmixed;
  class person;
  model y = time;
  random int time / type=fa0(2) sub=person;
  parms / lowerb=1e-4,.,1e-4;
run;
```

The **TYPE=FA0(2)** structure specifies a Cholesky root parameterization for the 2×2 unstructured blocks in **G**. This parameterization ensures that the **G** matrix is nonnegative definite, and the **PARMS** statement then ensures that it is positive definite by constraining the two diagonal terms to be greater than or equal to $1E-4$.

NOITER

requests that no optimization iterations be performed and that PROC HPLMIXED use the best value from the grid search to perform inferences. By default, iterations begin at the best value from the **PARMS** grid search. The **NOITER** option will be implied by the specification of the **BLUP** option in the HPLMIXED statement.

PARMSDATA=SAS-data-set

PDATA=SAS-data-set

reads in covariance parameter values from a SAS data set. The data set should contain the **Est** or **Covp1** through **Covpn** variables.

UPPERB=value-list

enables you to specify upper boundary constraints on the covariance parameters. The *value-list* specification is a list of numbers or missing values (.) separated by commas. You must list the numbers in the order that PROC HPLMIXED uses for the covariance parameters, and each number corresponds to the upper boundary constraint. A missing value instructs PROC HPLMIXED to use its default constraint. If you do not specify numbers for all of the covariance parameters, PROC HPLMIXED assumes that the remaining ones are missing.

PERFORMANCE Statement

PERFORMANCE < *performance-options* > ;

The **PERFORMANCE** statement defines performance parameters for multithreaded and distributed computing, passes variables about the distributed computing environment, and requests detailed results about the performance characteristics of a SAS high-performance analytical procedure.

You can also use the **PERFORMANCE** statement to control whether a SAS high-performance analytical procedure executes in single-machine mode or distributed mode.

The **PERFORMANCE** statement for SAS high-performance analytical procedures is documented in the section “**PERFORMANCE Statement**” (Chapter 2, *SAS/STAT User’s Guide: High-Performance Procedures*).

RANDOM Statement

RANDOM *random-effects* < / *options* > ;

The **RANDOM** statement defines the random effects that constitute the $\boldsymbol{\gamma}$ vector in the mixed model. You can use this statement to specify traditional variance component models and to specify random coefficients. The random effects can be classification or continuous, and multiple **RANDOM** statements are possible.

Using notation from the section “**Linear Mixed Models Theory**” on page 4297, the purpose of the **RANDOM** statement is to define the **Z** matrix of the mixed model, the random effects in the $\boldsymbol{\gamma}$ vector, and the structure

of **G**. The **Z** matrix is constructed exactly as the **X** matrix for the fixed effects is constructed, and the **G** matrix is constructed to correspond with the effects that constitute **Z**. The structure of **G** is defined by using the **TYPE=** option.

You can specify **INTERCEPT** (or **INT**) as a random effect to indicate the intercept. **PROC HPLMIXED** does not include the intercept in the **RANDOM** statement by default as it does in the **MODEL** statement.

Table 54.5 summarizes important options in the **RANDOM** statement. All options are subsequently discussed in alphabetical order.

Table 54.5 Summary of Important **RANDOM** Statement Options

Option	Description
Construction of Covariance Structure	
SUBJECT=	Identifies the subjects in the model
TYPE=	Specifies the covariance structure
Statistical Output	
ALPHA=α	Determines the confidence level ($1 - \alpha$)
CL	Requests confidence limits for predictors of random effects
SOLUTION	Displays solutions $\hat{\boldsymbol{\gamma}}$ of the random effects

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

ALPHA=number

sets the confidence level to be $1 - \text{number}$ for each confidence interval of the random-effects estimates. The value of *number* must be between 0 and 1; the default is 0.05.

CL

requests that *t*-type confidence limits be constructed for each of the random-effect estimates. The confidence level is 0.95 by default; this can be changed with the **ALPHA=** option.

SOLUTION

S

requests that the solution for the random-effects parameters be produced. Using notation from the section “[Linear Mixed Models Theory](#)” on page 4297, these estimates are the empirical best linear unbiased predictors (EBLUPs), $\hat{\boldsymbol{\gamma}} = \hat{\mathbf{G}}\mathbf{Z}'\hat{\mathbf{V}}^{-1}(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})$. They can be useful for comparing the random effects from different experimental units and can also be treated as residuals in performing diagnostics for your mixed model.

The numbers displayed in the SE Pred column of the “Solution for Random Effects” table are not the standard errors of the $\hat{\boldsymbol{\gamma}}$ displayed in the Estimate column; rather, they are the standard errors of predictions $\hat{\boldsymbol{\gamma}}_i - \boldsymbol{\gamma}_i$, where $\hat{\boldsymbol{\gamma}}_i$ is the *i*th EBLUP and $\boldsymbol{\gamma}_i$ is the *i*th random-effect parameter.

SUBJECT=effect

SUB=effect

identifies the subjects in your mixed model. Complete independence is assumed across subjects; thus, for the **RANDOM** statement, the **SUBJECT=** option produces a block-diagonal structure in **G** with

identical blocks. In fact, specifying a subject effect is equivalent to nesting all other effects in the RANDOM statement within the subject effect.

When you specify the SUBJECT= option and a classification random effect, computations are usually much quicker if the levels of the random effect are duplicated within each level of the SUBJECT= effect.

TYPE=covariance-structure

specifies the covariance structure of **G**. Valid values for *covariance-structure* and their descriptions are listed in Table 54.6. Although a variety of structures are available, most applications call for either TYPE=VC or TYPE=UN. The TYPE=VC (variance components) option is the default structure, and it models a different variance component for each random effect.

The TYPE=UN (unstructured) option is useful for correlated random coefficient models. For example, the following statement specifies a random intercept-slope model that has different variances for the intercept and slope and a covariance between them:

```
random intercept age / type=un subject=person;
```

You can also use TYPE=FA0(2) here to request a **G** estimate that is constrained to be nonnegative definite.

If you are constructing your own columns of **Z** with continuous variables, you can use the TYPE=TOEP(1) structure to group them together to have a common variance component. If you want to have different covariance structures in different parts of **G**, you must use multiple RANDOM statements with different TYPE= options.

Table 54.6 Covariance Structures

Structure	Description	Parms	(i, j) element
ANTE(1)	Antedependence	$2t - 1$	$\sigma_i \sigma_j \prod_{k=i}^{j-1} \rho_k$
AR(1)	Autoregressive(1)	2	$\sigma^2 \rho^{ i-j }$
ARH(1)	Heterogeneous AR(1)	$t + 1$	$\sigma_i \sigma_j \rho^{ i-j }$
ARMA(1,1)	Autoregressive moving average(1,1)	3	$\sigma^2 [\gamma \rho^{ i-j -1} 1(i \neq j) + 1(i = j)]$
CS	Compound symmetry	2	$\sigma_1 + \sigma^2 1(i = j)$
CSH	Heterogeneous compound symmetry	$t + 1$	$\sigma_i \sigma_j [\rho 1(i \neq j) + 1(i = j)]$
FA(q)	Factor analytic	$\frac{q}{2}(2t - q + 1) + t$	$\sum_{k=1}^{\min(i,j,q)} \lambda_{ik} \lambda_{jk} + \sigma_i^2 1(i = j)$
FA0(q)	No diagonal FA	$\frac{q}{2}(2t - q + 1)$	$\sum_{k=1}^{\min(i,j,q)} \lambda_{ik} \lambda_{jk}$
FA1(q)	Equal diagonal FA	$\frac{q}{2}(2t - q + 1) + 1$	$\sum_{k=1}^{\min(i,j,q)} \lambda_{ik} \lambda_{jk} + \sigma^2 1(i = j)$
HF	Huynh-Feldt	$t + 1$	$(\sigma_i^2 + \sigma_j^2)/2 + \lambda 1(i \neq j)$
SIMPLE	An alias for VC	q	$\sigma_k^2 1(i = j)$ for the kth effect
TOEP	Toeplitz	t	$\sigma_{ i-j +1}$
TOEP(q)	Banded Toeplitz	q	$\sigma_{ i-j +1} 1(i - j < q)$
TOEPH	Heterogeneous TOEP	$2t - 1$	$\sigma_i \sigma_j \rho_{ i-j }$

Table 54.6 continued

Structure	Description	Parms	(i, j) element
TOEPH(<i>q</i>)	Banded heterogeneous TOEP	$t + q - 1$	$\sigma_i \sigma_j \rho_{ i-j } 1(i - j < q)$
UN	Unstructured	$t(t + 1)/2$	σ_{ij}
UN(<i>q</i>)	Banded	$\frac{q}{2}(2t - q + 1)$	$\sigma_{ij} 1(i - j < q)$
UNR	Unstructured correlation	$t(t + 1)/2$	$\sigma_i \sigma_j \rho_{\max(i,j) \min(i,j)}$
UNR(<i>q</i>)	Banded correlations	$\frac{q}{2}(2t - q + 1)$	$\sigma_i \sigma_j \rho_{\max(i,j) \min(i,j)}$
VC	Variance components	q	$\sigma_k^2 1(i = j)$ for the <i>k</i> th effect

In Table 54.6, the Parms column represents the number of covariance parameters in the structure, t is the overall dimension of the covariance matrix, and $1(A)$ equals 1 when A is true and 0 otherwise. For example, $1(i = j)$ equals 1 when $i = j$ and 0 otherwise, and $1(|i - j| < q)$ equals 1 when $|i - j| < q$ and 0 otherwise. For the **TYPE=TOEPH** structures, $\rho_0 = 1$; for the **TYPE=UNR** structures, $\rho_{ii} = 1$ for all i .

Table 54.7 lists some examples of the structures in Table 54.6.

Table 54.7 Covariance Structure Examples

Description	Structure	Example
Variance components	VC (default)	$\begin{bmatrix} \sigma_B^2 & 0 & 0 & 0 \\ 0 & \sigma_B^2 & 0 & 0 \\ 0 & 0 & \sigma_{AB}^2 & 0 \\ 0 & 0 & 0 & \sigma_{AB}^2 \end{bmatrix}$
Compound symmetry	CS	$\begin{bmatrix} \sigma^2 + \sigma_1 & \sigma_1 & \sigma_1 & \sigma_1 \\ \sigma_1 & \sigma^2 + \sigma_1 & \sigma_1 & \sigma_1 \\ \sigma_1 & \sigma_1 & \sigma^2 + \sigma_1 & \sigma_1 \\ \sigma_1 & \sigma_1 & \sigma_1 & \sigma^2 + \sigma_1 \end{bmatrix}$
Unstructured	UN	$\begin{bmatrix} \sigma_1^2 & \sigma_{21} & \sigma_{31} & \sigma_{41} \\ \sigma_{21} & \sigma_2^2 & \sigma_{32} & \sigma_{42} \\ \sigma_{31} & \sigma_{32} & \sigma_3^2 & \sigma_{43} \\ \sigma_{41} & \sigma_{42} & \sigma_{43} & \sigma_4^2 \end{bmatrix}$
Banded main diagonal	UN(1)	$\begin{bmatrix} \sigma_1^2 & 0 & 0 & 0 \\ 0 & \sigma_2^2 & 0 & 0 \\ 0 & 0 & \sigma_3^2 & 0 \\ 0 & 0 & 0 & \sigma_4^2 \end{bmatrix}$
First-order autoregressive	AR(1)	$\sigma^2 \begin{bmatrix} 1 & \rho & \rho^2 & \rho^3 \\ \rho & 1 & \rho & \rho^2 \\ \rho^2 & \rho & 1 & \rho \\ \rho^3 & \rho^2 & \rho & 1 \end{bmatrix}$

Table 54.7 *continued*

Description	Structure	Example
Toeplitz	TOEP	$\begin{bmatrix} \sigma^2 & \sigma_1 & \sigma_2 & \sigma_3 \\ \sigma_1 & \sigma^2 & \sigma_1 & \sigma_2 \\ \sigma_2 & \sigma_1 & \sigma^2 & \sigma_1 \\ \sigma_3 & \sigma_2 & \sigma_1 & \sigma^2 \end{bmatrix}$
Toeplitz with two bands	TOEP(2)	$\begin{bmatrix} \sigma^2 & \sigma_1 & 0 & 0 \\ \sigma_1 & \sigma^2 & \sigma_1 & 0 \\ 0 & \sigma_1 & \sigma^2 & \sigma_1 \\ 0 & 0 & \sigma_1 & \sigma^2 \end{bmatrix}$
Heterogeneous autoregressive(1)	ARH(1)	$\begin{bmatrix} \sigma_1^2 & \sigma_1\sigma_2\rho & \sigma_1\sigma_3\rho^2 & \sigma_1\sigma_4\rho^3 \\ \sigma_2\sigma_1\rho & \sigma_2^2 & \sigma_2\sigma_3\rho & \sigma_2\sigma_4\rho^2 \\ \sigma_3\sigma_1\rho^2 & \sigma_3\sigma_2\rho & \sigma_3^2 & \sigma_3\sigma_4\rho \\ \sigma_4\sigma_1\rho^3 & \sigma_4\sigma_2\rho^2 & \sigma_4\sigma_3\rho & \sigma_4^2 \end{bmatrix}$
First-order autoregressive moving average	ARMA(1,1)	$\sigma^2 \begin{bmatrix} 1 & \gamma & \gamma\rho & \gamma\rho^2 \\ \gamma & 1 & \gamma & \gamma\rho \\ \gamma\rho & \gamma & 1 & \gamma \\ \gamma\rho^2 & \gamma\rho & \gamma & 1 \end{bmatrix}$
Heterogeneous compound symmetry	CSH	$\begin{bmatrix} \sigma_1^2 & \sigma_1\sigma_2\rho & \sigma_1\sigma_3\rho & \sigma_1\sigma_4\rho \\ \sigma_2\sigma_1\rho & \sigma_2^2 & \sigma_2\sigma_3\rho & \sigma_2\sigma_4\rho \\ \sigma_3\sigma_1\rho & \sigma_3\sigma_2\rho & \sigma_3^2 & \sigma_3\sigma_4\rho \\ \sigma_4\sigma_1\rho & \sigma_4\sigma_2\rho & \sigma_4\sigma_3\rho & \sigma_4^2 \end{bmatrix}$
First-order factor analytic	FA(1)	$\begin{bmatrix} \lambda_1^2 + d_1 & \lambda_1\lambda_2 & \lambda_1\lambda_3 & \lambda_1\lambda_4 \\ \lambda_2\lambda_1 & \lambda_2^2 + d_2 & \lambda_2\lambda_3 & \lambda_2\lambda_4 \\ \lambda_3\lambda_1 & \lambda_3\lambda_2 & \lambda_3^2 + d_3 & \lambda_3\lambda_4 \\ \lambda_4\lambda_1 & \lambda_4\lambda_2 & \lambda_4\lambda_3 & \lambda_4^2 + d_4 \end{bmatrix}$
Huynh-Feldt	HF	$\begin{bmatrix} \sigma_1^2 & \frac{\sigma_1^2 + \sigma_2^2}{2} - \lambda & \frac{\sigma_1^2 + \sigma_3^2}{2} - \lambda \\ \frac{\sigma_2^2 + \sigma_1^2}{2} - \lambda & \sigma_2^2 & \frac{\sigma_2^2 + \sigma_3^2}{2} - \lambda \\ \frac{\sigma_3^2 + \sigma_1^2}{2} - \lambda & \frac{\sigma_3^2 + \sigma_2^2}{2} - \lambda & \sigma_3^2 \end{bmatrix}$
First-order antedependence	ANTE(1)	$\begin{bmatrix} \sigma_1^2 & \sigma_1\sigma_2\rho_1 & \sigma_1\sigma_3\rho_1\rho_2 \\ \sigma_2\sigma_1\rho_1 & \sigma_2^2 & \sigma_2\sigma_3\rho_2 \\ \sigma_3\sigma_1\rho_2\rho_1 & \sigma_3\sigma_2\rho_2 & \sigma_3^2 \end{bmatrix}$
Heterogeneous Toeplitz	TOEPH	$\begin{bmatrix} \sigma_1^2 & \sigma_1\sigma_2\rho_1 & \sigma_1\sigma_3\rho_2 & \sigma_1\sigma_4\rho_3 \\ \sigma_2\sigma_1\rho_1 & \sigma_2^2 & \sigma_2\sigma_3\rho_1 & \sigma_2\sigma_4\rho_2 \\ \sigma_3\sigma_1\rho_2 & \sigma_3\sigma_2\rho_1 & \sigma_3^2 & \sigma_3\sigma_4\rho_1 \\ \sigma_4\sigma_1\rho_3 & \sigma_4\sigma_2\rho_2 & \sigma_4\sigma_3\rho_1 & \sigma_4^2 \end{bmatrix}$
Unstructured correlations	UNR	$\begin{bmatrix} \sigma_1^2 & \sigma_1\sigma_2\rho_{21} & \sigma_1\sigma_3\rho_{31} & \sigma_1\sigma_4\rho_{41} \\ \sigma_2\sigma_1\rho_{21} & \sigma_2^2 & \sigma_2\sigma_3\rho_{32} & \sigma_2\sigma_4\rho_{42} \\ \sigma_3\sigma_1\rho_{31} & \sigma_3\sigma_2\rho_{32} & \sigma_3^2 & \sigma_3\sigma_4\rho_{43} \\ \sigma_4\sigma_1\rho_{41} & \sigma_4\sigma_2\rho_{42} & \sigma_4\sigma_3\rho_{43} & \sigma_4^2 \end{bmatrix}$

The following list provides some further information about these covariance structures:

- TYPE=ANTE(1) specifies the first-order antedependence structure (Kenward 1987; Patel 1991; Macchiavelli and Arnold 1994). In Table 54.6, σ_i^2 is the i variance parameter, and ρ_k is the k autocorrelation parameter that satisfies $|\rho_k| < 1$.
- TYPE=AR(1) specifies a first-order autoregressive structure. PROC HPLMIXED imposes the constraint $|\rho| < 1$ for stationarity.
- TYPE=ARH(1) specifies a heterogeneous first-order autoregressive structure. As with TYPE=AR(1), PROC HPLMIXED imposes the constraint $|\rho| < 1$ for stationarity.
- TYPE=ARMA(1,1) specifies the first-order autoregressive moving average structure. In Table 54.6, ρ is the autoregressive parameter, γ models a moving average component, and σ^2 is the residual variance. In the notation of Fuller (1976, p. 68), $\rho = \theta_1$ and

$$\gamma = \frac{(1 + b_1\theta_1)(\theta_1 + b_1)}{1 + b_1^2 + 2b_1\theta_1}$$

The example in Table 54.7 and $|b_1| < 1$ imply that

$$b_1 = \frac{\beta - \sqrt{\beta^2 - 4\alpha^2}}{2\alpha}$$

where $\alpha = \gamma - \rho$ and $\beta = 1 + \rho^2 - 2\gamma\rho$. PROC HPLMIXED imposes the constraints $|\rho| < 1$ and $|\gamma| < 1$ for stationarity, although the resulting covariance matrix is not positive definite for some values of ρ and γ in this region. When the estimated value of ρ becomes negative, the computed covariance is multiplied by $\cos(\pi d_{ij})$ to account for the negativity.

- TYPE=CS specifies the compound-symmetry structure, which has constant variance and constant covariance.
- TYPE=CSH specifies the heterogeneous compound-symmetry structure. This structure has a different variance parameter for each diagonal element, and it uses the square roots of these parameters in the off-diagonal entries. In Table 54.6, σ_i^2 is the i variance parameter, and ρ is the correlation parameter that satisfies $|\rho| < 1$.
- TYPE=FA(q) specifies the factor-analytic structure with q factors (Jennrich and Schluchter 1986). This structure is of the form $\mathbf{\Lambda}\mathbf{\Lambda}' + \mathbf{D}$, where $\mathbf{\Lambda}$ is a $t \times q$ rectangular matrix and $\mathbf{D}$ is a $t \times t$ diagonal matrix with t different parameters. When $q > 1$, the elements of $\mathbf{\Lambda}$ in its upper right corner (that is, the elements in the i row and j column for $j > i$) are set to zero to fix the rotation of the structure.
- TYPE=FA0(q) is similar to the FA(q) structure except that no diagonal matrix $\mathbf{D}$ is included. When $q < t$ (that is, when the number of factors is less than the dimension of the matrix), this structure is nonnegative definite but not of full rank. In this situation, you can use this structure for approximating an unstructured $\mathbf{G}$ matrix in the RANDOM statement. When $q = t$, you can use this structure to constrain $\mathbf{G}$ to be nonnegative definite in the RANDOM statement.
- TYPE=FA1(q) is similar to the TYPE=FA(q) structure except that all of the elements in $\mathbf{D}$ are constrained to be equal. This offers a useful and more parsimonious alternative to the full factor-analytic structure.

- TYPE=HF** specifies the Huynh-Feldt covariance structure (Huynh and Feldt 1970). This structure is similar to the **TYPE=CSH** structure in that it has the same number of parameters and heterogeneity along the main diagonal. However, it constructs the off-diagonal elements by taking arithmetic means rather than geometric means.
- You can perform a likelihood ratio test of the Huynh-Feldt conditions by running PROC HPLMIXED twice, once with **TYPE=HF** and once with **TYPE=UN**, and then subtracting their respective values of -2 times the maximized likelihood.
- If PROC HPLMIXED does not converge under your Huynh-Feldt model, you can specify your own starting values with the **PARMS** statement. The default MIVQUE(0) starting values can sometimes be poor for this structure. A good choice for starting values is often the parameter estimates that correspond to an initial fit that uses **TYPE=CS**.
- TYPE=SIMPLE** is an alias for **TYPE=VC**.
- TYPE=TOEP**<(q)> specifies a banded Toeplitz structure. This can be viewed as a moving average structure with order equal to $q - 1$. The **TYPE=TOEP** option is a full Toeplitz matrix, which can be viewed as an autoregressive structure with order equal to the dimension of the matrix. The specification **TYPE=TOEP**(1) is the same as $\sigma^2 I$, where I is an identity matrix, and it can be useful for specifying the same variance component for several effects.
- TYPE=TOEPH**<(q)> specifies a heterogeneous banded Toeplitz structure. In Table 54.6, σ_i^2 is the i variance parameter and ρ_j is the j correlation parameter that satisfies $|\rho_j| < 1$. If you specify the order parameter q , then PROC HPLMIXED estimates only the first q bands of the matrix, setting all higher bands equal to 0. The option **TOEPH**(1) is equivalent to both the **TYPE=UN**(1) and **TYPE=UNR**(1) options.
- TYPE=UN**<(q)> specifies a completely general (unstructured) covariance matrix that is parameterized directly in terms of variances and covariances. The variances are constrained to be nonnegative, and the covariances are unconstrained. This structure is not constrained to be nonnegative definite in order to avoid nonlinear constraints. However, you can use the **TYPE=FA0** structure if you want this constraint to be imposed by a Cholesky factorization. If you specify the order parameter q , then PROC HPLMIXED estimates only the first q bands of the matrix, setting all higher bands equal to 0.
- TYPE=UNR**<(q)> specifies a completely general (unstructured) covariance matrix that is parameterized in terms of variances and correlations. This structure fits the same model as the **TYPE=UN**(q) option but with a different parameterization. The i variance parameter is σ_i^2 . The parameter ρ_{jk} is the correlation between the j and k measurements; it satisfies $|\rho_{jk}| < 1$. If you specify the order parameter r , then PROC HPLMIXED estimates only the first q bands of the matrix, setting all higher bands equal to zero.
- TYPE=VC** specifies standard variance components. This is the default structure for both the **RANDOM** and **REPEATED** statements. In the **RANDOM** statement, a distinct variance component is assigned to each effect.

Jennrich and Schluchter (1986) provide general information about the use of covariance structures, and Wolfinger (1996) presents details about many of the heterogeneous structures.

REPEATED Statement

REPEATED *repeated-effect* < / options > ;

The REPEATED statement specifies the **R** matrix in the mixed model. If no REPEATED statement is specified, **R** is assumed to be equal to $\sigma^2\mathbf{I}$. For this release, you can use the REPEATED statement only with the BLUP option. The statement is ignored when no BLUP option is specified.

The *repeated-effect* is required, because the order of the input data is not necessarily reproducible in a distributed environment. The *repeated-effect* must contain only classification variables. Make sure that the levels of the *repeated-effect* are different for each observation within a subject; otherwise, PROC HPLMIXED constructs identical rows in **R** that correspond to the observations with the same level. This results in a singular **R** matrix and an infinite likelihood.

Table 54.8 summarizes important options in the REPEATED statement. All options are subsequently discussed in alphabetical order.

Table 54.8 Summary of Important REPEATED Statement Options

Option	Description
Construction of Covariance Structure	
SUBJECT=	Identifies the subjects in the R-side model
TYPE=	Specifies the R-side covariance structure

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

SUBJECT=*effect*

SUB=*effect*

identifies the subjects in your mixed model. Complete independence is assumed across subjects; therefore, the SUBJECT= option produces a block-diagonal structure in **R** with identical blocks. When the SUBJECT= effect consists entirely of classification variables, the blocks of **R** correspond to observations that share the same level of that effect. These blocks are sorted according to this effect as well.

If you want to model nonzero covariance among all of the observations in your SAS data set, specify SUBJECT=Dummy_Intercept to treat the data as if they are all from one subject. You need to create this Dummy_Intercept variable in the data set. However, be aware that in this case PROC HPLMIXED manipulates an **R** matrix with dimensions equal to the number of observations.

TYPE=*covariance-structure*

specifies the covariance structure of the **R** matrix. The SUBJECT= option defines the blocks of **R**, and the TYPE= option specifies the structure of these blocks. The default structure is VC. You can specify any of the covariance structures that are described in the TYPE= option in the RANDOM statement.

Details: HPLMIXED Procedure

Linear Mixed Models Theory

This section provides an overview of a likelihood-based approach to linear mixed models. This approach simplifies and unifies many common statistical analyses, including those that involve repeated measures, random effects, and random coefficients. The basic assumption is that the data are linearly related to unobserved multivariate normal random variables. For extensions to nonlinear and nonnormal situations, see the documentation of the GLIMMIX and NLMIXED procedures in the *SAS/STAT User's Guide*. Additional theory and examples are provided in Littell et al. (2006); Verbeke and Molenberghs (1997, 2000); and Burdick and Graybill (1992).

Matrix Notation

Suppose that you observe n data points $y_1, \dots, y_n$ and that you want to explain them by using n values for each of p explanatory variables $x_{11}, \dots, x_{1p}, x_{21}, \dots, x_{2p}, \dots, x_{n1}, \dots, x_{np}$. The x_{ij} values can be either regression-type continuous variables or dummy variables that indicate class membership. The standard linear model for this setup is

$$y_i = \sum_{j=1}^p x_{ij} \beta_j + \epsilon_i \quad i = 1, \dots, n$$

where $\beta_1, \dots, \beta_p$ are unknown fixed-effects parameters to be estimated and $\epsilon_1, \dots, \epsilon_n$ are unknown independent and identically distributed normal (Gaussian) random variables with mean 0 and variance σ^2 .

The preceding equations can be written simultaneously by using vectors and a matrix, as follows:

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{bmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_p \end{bmatrix} + \begin{bmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_n \end{bmatrix}$$

For convenience, simplicity, and extendability, this entire system is written as

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

where $\mathbf{y}$ denotes the vector of observed y_i 's, $\mathbf{X}$ is the known matrix of x_{ij} 's, $\boldsymbol{\beta}$ is the unknown fixed-effects parameter vector, and $\boldsymbol{\epsilon}$ is the unobserved vector of independent and identically distributed Gaussian random errors.

In addition to denoting data, random variables, and explanatory variables in the preceding fashion, the subsequent development makes use of basic matrix operators such as transpose ($'$), inverse ($^{-1}$), generalized inverse ($^{-}$), determinant ($|\cdot|$), and matrix multiplication. See Searle (1982) for details about these and other matrix techniques.

Formulation of the Mixed Model

The previous general linear model is certainly a useful one (Searle 1971), and it is the one fitted by the GLM procedure. However, many times the distributional assumption about ϵ is too restrictive. The mixed model extends the general linear model by allowing a more flexible specification of the covariance matrix of ϵ . In other words, it allows for both correlation and heterogeneous variances, although you still assume normality.

The mixed model is written as

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

where everything is the same as in the general linear model except for the addition of the known design matrix, $\mathbf{Z}$, and the vector of unknown random-effects parameters, $\boldsymbol{\gamma}$. The matrix $\mathbf{Z}$ can contain either continuous or dummy variables, just like $\mathbf{X}$. The name *mixed model* comes from the fact that the model contains both fixed-effects parameters, $\boldsymbol{\beta}$, and random-effects parameters, $\boldsymbol{\gamma}$. See Henderson (1990) and Searle, Casella, and McCulloch (1992) for historical developments of the mixed model.

A key assumption in the foregoing analysis is that $\boldsymbol{\gamma}$ and $\boldsymbol{\epsilon}$ are normally distributed with

$$\begin{aligned} E \begin{bmatrix} \boldsymbol{\gamma} \\ \boldsymbol{\epsilon} \end{bmatrix} &= \begin{bmatrix} \mathbf{0} \\ \mathbf{0} \end{bmatrix} \\ \text{Var} \begin{bmatrix} \boldsymbol{\gamma} \\ \boldsymbol{\epsilon} \end{bmatrix} &= \begin{bmatrix} \mathbf{G} & \mathbf{0} \\ \mathbf{0} & \mathbf{R} \end{bmatrix} \end{aligned}$$

Therefore, the variance of $\mathbf{y}$ is $\mathbf{V} = \mathbf{ZGZ}' + \mathbf{R}$. You can model $\mathbf{V}$ by setting up the random-effects design matrix $\mathbf{Z}$ and by specifying covariance structures for $\mathbf{G}$ and $\mathbf{R}$.

Note that this is a general specification of the mixed model, in contrast to many texts and articles that discuss only simple random effects. Simple random effects are a special case of the general specification with $\mathbf{Z}$ containing dummy variables, $\mathbf{G}$ containing variance components in a diagonal structure, and $\mathbf{R} = \sigma^2 \mathbf{I}_n$, where $\mathbf{I}_n$ denotes the $n \times n$ identity matrix. The general linear model is a further special case with $\mathbf{Z} = \mathbf{0}$ and $\mathbf{R} = \sigma^2 \mathbf{I}_n$.

The following two examples illustrate the most common formulations of the general linear mixed model.

Example: Growth Curve with Compound Symmetry

Suppose that you have three growth curve measurements for s individuals and that you want to fit an overall linear trend in time. Your $\mathbf{X}$ matrix is as follows:

$$\mathbf{X} = \begin{bmatrix} 1 & 1 \\ 1 & 2 \\ 1 & 3 \\ \vdots & \vdots \\ 1 & 1 \\ 1 & 2 \\ 1 & 3 \end{bmatrix}$$

The first column (coded entirely with 1s) fits an intercept, and the second column (coded with series of 1, 2, 3) fits a slope. Here, $n = 3s$ and $p = 2$.

Suppose further that you want to introduce a common correlation among the observations from a single individual, with correlation being the same for all individuals. One way of setting this up in the general mixed

model is to eliminate the **Z** and **G** matrices and let the **R** matrix be block-diagonal with blocks corresponding to the individuals and with each block having the *compound-symmetry* structure. This structure has two unknown parameters, one modeling a common covariance and the other modeling a residual variance. The form for **R** would then be

$$\mathbf{R} = \begin{bmatrix} \sigma_1^2 + \sigma^2 & \sigma_1^2 & \sigma_1^2 & & & \\ \sigma_1^2 & \sigma_1^2 + \sigma^2 & \sigma_1^2 & & & \\ \sigma_1^2 & \sigma_1^2 & \sigma_1^2 + \sigma^2 & & & \\ & & & \ddots & & \\ & & & & \sigma_1^2 + \sigma^2 & \sigma_1^2 & \sigma_1^2 \\ & & & & \sigma_1^2 & \sigma_1^2 + \sigma^2 & \sigma_1^2 \\ & & & & \sigma_1^2 & \sigma_1^2 & \sigma_1^2 + \sigma^2 \end{bmatrix}$$

where blanks denote zeros. There are $3s$ rows and columns altogether, and the common correlation is $\sigma_1^2/(\sigma_1^2 + \sigma^2)$.

The following PROC HPLMIXED statements fit this model:

```
proc hplmixed;
  class indiv;
  model y = time;
  repeated morder/ type=cs subject=indiv;
run;
```

Here, INDIV is a classification variable that indexes individuals. The **MODEL** statement fits a straight line for TIME ; the intercept is fit by default just as in PROC GLM. The **REPEATED** statement models the **R** matrix: **TYPE=CS** specifies the compound symmetry structure, and **SUBJECT=INDIV** specifies the blocks of **R**, and **MORDER** is the repeated effect that records the order of the measurements for each individual.

An alternative way of specifying the common intra-individual correlation is to let

$$\mathbf{Z} = \begin{bmatrix} 1 & & & & & \\ 1 & & & & & \\ 1 & & & & & \\ & 1 & & & & \\ & 1 & & & & \\ & 1 & & & & \\ & & \ddots & & & \\ & & & 1 & & \\ & & & 1 & & \\ & & & 1 & & \end{bmatrix}$$

$$\mathbf{G} = \begin{bmatrix} \sigma_1^2 & & & & \\ & \sigma_1^2 & & & \\ & & \ddots & & \\ & & & \sigma_1^2 & \end{bmatrix}$$

and $\mathbf{R} = \sigma^2 \mathbf{I}_n$. The **Z** matrix has $3s$ rows and s columns, and **G** is $s \times s$.

You can set up this model in PROC HPLMIXED in two different but equivalent ways:

```

proc hplmixed;
  class indiv;
  model y = time;
  random indiv;
run;

proc hplmixed;
  class indiv;
  model y = time;
  random intercept / subject=indiv;
run;

```

Both of these specifications fit the same model as the previous one that used the **REPEATED** statement. However, the **RANDOM** specifications constrain the correlation to be positive, whereas the **REPEATED** specification leaves the correlation unconstrained.

Example: Split-Plot Design

The split-plot design involves two experimental treatment factors, A and B, and two different sizes of experimental units to which they are applied (Winer 1971; Snedecor and Cochran 1980; Milliken and Johnson 1992; Steel, Torrie, and Dickey 1997). The levels of A are randomly assigned to the larger-sized experimental units, called *whole plots*, whereas the levels of B are assigned to the smaller-sized experimental units, the *subplots*. The subplots are assumed to be nested within the whole plots, so that a whole plot consists of a cluster of subplots and a level of A is applied to the entire cluster.

Such an arrangement is often necessary by nature of the experiment; the classical example is the application of fertilizer to large plots of land and different crop varieties planted in subdivisions of the large plots. For this example, fertilizer is the whole-plot factor A and variety is the subplot factor B.

The first example is a split-plot design for which the whole plots are arranged in a randomized block design. The appropriate PROC HPLMIXED statements are as follows:

```

proc hplmixed;
  class a b block;
  model y = a b a*b;
  random block a*block;
run;

```

Here

$$\mathbf{R} = \sigma^2 \mathbf{I}_{24}$$

However, GLS requires knowledge of $\mathbf{V}$ and therefore knowledge of $\mathbf{G}$ and $\mathbf{R}$. Lacking such information, one approach is to use an *estimated* GLS, in which you insert some reasonable estimate for $\mathbf{V}$ into the minimization problem. The goal thus becomes finding a reasonable estimate of $\mathbf{G}$ and $\mathbf{R}$.

In many situations, the best approach is to use likelihood-based methods, exploiting the assumption that $\boldsymbol{\gamma}$ and $\boldsymbol{\epsilon}$ are normally distributed (Hartley and Rao 1967; Patterson and Thompson 1971; Harville 1977; Laird and Ware 1982; Jennrich and Schluchter 1986). PROC HPLMIXED implements two likelihood-based methods: maximum likelihood (ML) and restricted (residual) maximum likelihood (REML). A favorable theoretical property of ML and REML is that they accommodate data that are missing at random (Rubin 1976; Little 1995).

PROC HPLMIXED constructs an objective function associated with ML or REML and maximizes it over all unknown parameters. Using calculus, it is possible to reduce this maximization problem to one over only the parameters in $\mathbf{G}$ and $\mathbf{R}$. The corresponding log-likelihood functions are as follows:

$$\begin{aligned}\text{ML : } l(\mathbf{G}, \mathbf{R}) &= -\frac{1}{2} \log |\mathbf{V}| - \frac{1}{2} \mathbf{r}' \mathbf{V}^{-1} \mathbf{r} - \frac{n}{2} \log(2\pi) \\ \text{REML : } l_R(\mathbf{G}, \mathbf{R}) &= -\frac{1}{2} \log |\mathbf{V}| - \frac{1}{2} \log |\mathbf{X}' \mathbf{V}^{-1} \mathbf{X}| - \frac{1}{2} \mathbf{r}' \mathbf{V}^{-1} \mathbf{r} - \frac{n-p}{2} \log(2\pi)\end{aligned}$$

where $\mathbf{r} = \mathbf{y} - \mathbf{X}(\mathbf{X}' \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}' \mathbf{V}^{-1} \mathbf{y}$ and p is the rank of $\mathbf{X}$. By default, PROC HPLMIXED actually minimizes a normalized form of -2 times these functions by using a ridge-stabilized Newton-Raphson algorithm by default. Lindstrom and Bates (1988) provide reasons for preferring Newton-Raphson to the expectation-maximum (EM) algorithm described in Dempster, Laird, and Rubin (1977) and Laird, Lange, and Stram (1987), in addition to analytical details for implementing a QR-decomposition approach to the problem. Wolfinger, Tobias, and Sall (1994) present the sweep-based algorithms that are implemented in PROC HPLMIXED. You can change the optimization technique with the `TECHNIQUE=` option in the `PROC HPLMIXED` statement.

One advantage of using the Newton-Raphson algorithm is that the second derivative matrix of the objective function evaluated at the optima is available upon completion. Denoting this matrix $\mathbf{H}$, the asymptotic theory of maximum likelihood (Serfling 1980) shows that $2\mathbf{H}^{-1}$ is an asymptotic variance-covariance matrix of the estimated parameters of $\mathbf{G}$ and $\mathbf{R}$. Thus, tests and confidence intervals based on asymptotic normality can be obtained. However, these can be unreliable in small samples, especially for parameters such as variance components that have sampling distributions that tend to be skewed to the right.

If a residual variance σ^2 is a part of your mixed model, it can usually be *profiled* out of the likelihood. This means solving analytically for the optimal σ^2 and plugging this expression back into the likelihood formula (Wolfinger, Tobias, and Sall 1994). This reduces the number of optimization parameters by 1 and can improve convergence properties. PROC HPLMIXED profiles the residual variance out of the log likelihood.

Estimating Fixed and Random Effects in the Mixed Model

ML and REML methods provide estimates of $\mathbf{G}$ and $\mathbf{R}$, which are denoted $\hat{\mathbf{G}}$ and $\hat{\mathbf{R}}$, respectively. To obtain estimates of $\boldsymbol{\beta}$ and predicted values of $\boldsymbol{\gamma}$, the standard method is to solve the *mixed model equations* (Henderson 1984):

$$\begin{bmatrix} \mathbf{X}' \hat{\mathbf{R}}^{-1} \mathbf{X} & \mathbf{X}' \hat{\mathbf{R}}^{-1} \mathbf{Z} \\ \mathbf{Z}' \hat{\mathbf{R}}^{-1} \mathbf{X} & \mathbf{Z}' \hat{\mathbf{R}}^{-1} \mathbf{Z} + \hat{\mathbf{G}}^{-1} \end{bmatrix} \begin{bmatrix} \hat{\boldsymbol{\beta}} \\ \hat{\boldsymbol{\gamma}} \end{bmatrix} = \begin{bmatrix} \mathbf{X}' \hat{\mathbf{R}}^{-1} \mathbf{y} \\ \mathbf{Z}' \hat{\mathbf{R}}^{-1} \mathbf{y} \end{bmatrix}$$

The solutions can also be written as

$$\begin{aligned}\hat{\beta} &= (\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{X})^{-1}\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{y} \\ \hat{\gamma} &= \hat{\mathbf{G}}\mathbf{Z}'\hat{\mathbf{V}}^{-1}(\mathbf{y} - \mathbf{X}\hat{\beta})\end{aligned}$$

and have connections with empirical Bayes estimators (Laird and Ware 1982; Carlin and Louis 1996). Note that the γ are random variables and not parameters (unknown constants) in the model. Technically, determining values for γ from the data is thus a prediction task, whereas determining values for β is an estimation task.

The mixed model equations are extended normal equations. The preceding expression assumes that $\hat{\mathbf{G}}$ is nonsingular. For the extreme case where the eigenvalues of $\hat{\mathbf{G}}$ are very large, $\hat{\mathbf{G}}^{-1}$ contributes very little to the equations and $\hat{\gamma}$ is close to what it would be if γ actually contained fixed-effects parameters. On the other hand, when the eigenvalues of $\hat{\mathbf{G}}$ are very small, $\hat{\mathbf{G}}^{-1}$ dominates the equations and $\hat{\gamma}$ is close to 0. For intermediate cases, $\hat{\mathbf{G}}^{-1}$ can be viewed as shrinking the fixed-effects estimates of γ toward 0 (Robinson 1991).

If $\hat{\mathbf{G}}$ is singular, then the mixed model equations are modified (Henderson 1984) as follows:

$$\begin{bmatrix} \mathbf{X}'\hat{\mathbf{R}}^{-1}\mathbf{X} & \mathbf{X}'\hat{\mathbf{R}}^{-1}\mathbf{Z}\hat{\mathbf{G}} \\ \hat{\mathbf{G}}'\mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{X} & \hat{\mathbf{G}}'\mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{Z}\hat{\mathbf{G}} + \mathbf{G} \end{bmatrix} \begin{bmatrix} \hat{\beta} \\ \hat{\tau} \end{bmatrix} = \begin{bmatrix} \mathbf{X}'\hat{\mathbf{R}}^{-1}\mathbf{y} \\ \hat{\mathbf{G}}'\mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{y} \end{bmatrix}$$

Denote the generalized inverses of the nonsingular $\hat{\mathbf{G}}$ and singular $\hat{\mathbf{G}}$ forms of the mixed model equations by $\mathbf{C}$ and $\mathbf{M}$, respectively. In the nonsingular case, the solution $\hat{\gamma}$ estimates the random effects directly. But in the singular case, the estimates of random effects are achieved through a back-transformation $\hat{\gamma} = \hat{\mathbf{G}}\hat{\tau}$ where $\hat{\tau}$ is the solution to the modified mixed model equations. Similarly, while in the nonsingular case $\mathbf{C}$ itself is the estimated covariance matrix for $(\hat{\beta}, \hat{\gamma})$, in the singular case the covariance estimate for $(\hat{\beta}, \hat{\mathbf{G}}\hat{\tau})$ is given by **PMP** where

$$\mathbf{P} = \begin{bmatrix} \mathbf{I} & \\ & \hat{\mathbf{G}} \end{bmatrix}$$

An example of when the singular form of the equations is necessary is when a variance component estimate falls on the boundary constraint of 0.

Statistical Properties

If $\mathbf{G}$ and $\mathbf{R}$ are known, $\hat{\beta}$ is the best linear unbiased estimator (BLUE) of β , and $\hat{\gamma}$ is the best linear unbiased predictor (BLUP) of γ (Searle 1971; Harville 1988, 1990; Robinson 1991; McLean, Sanders, and Stroup 1991). Here, “best” means minimum mean squared error. The covariance matrix of $(\hat{\beta} - \beta, \hat{\gamma} - \gamma)$ is

$$\mathbf{C} = \begin{bmatrix} \mathbf{X}'\mathbf{R}^{-1}\mathbf{X} & \mathbf{X}'\mathbf{R}^{-1}\mathbf{Z} \\ \mathbf{Z}'\mathbf{R}^{-1}\mathbf{X} & \mathbf{Z}'\mathbf{R}^{-1}\mathbf{Z} + \mathbf{G}^{-1} \end{bmatrix}^{-}$$

where $^{-}$ denotes a generalized inverse (Searle 1971).

However, $\mathbf{G}$ and $\mathbf{R}$ are usually unknown and are estimated by using one of the aforementioned methods. These estimates, $\hat{\mathbf{G}}$ and $\hat{\mathbf{R}}$, are therefore simply substituted into the preceding expression to obtain

$$\hat{\mathbf{C}} = \begin{bmatrix} \mathbf{X}'\hat{\mathbf{R}}^{-1}\mathbf{X} & \mathbf{X}'\hat{\mathbf{R}}^{-1}\mathbf{Z} \\ \mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{X} & \mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{Z} + \hat{\mathbf{G}}^{-1} \end{bmatrix}^{-}$$

as the approximate variance-covariance matrix of $(\hat{\beta} - \beta, \hat{\gamma} - \gamma)$. In this case, the BLUE and BLUP acronyms no longer apply, but the word *empirical* is often added to indicate such an approximation. The appropriate acronyms thus become EBLUE and EBLUP.

McLean and Sanders (1988) show that $\hat{\mathbf{C}}$ can also be written as

$$\hat{\mathbf{C}} = \begin{bmatrix} \hat{\mathbf{C}}_{11} & \hat{\mathbf{C}}'_{21} \\ \hat{\mathbf{C}}_{21} & \hat{\mathbf{C}}_{22} \end{bmatrix}$$

where

$$\hat{\mathbf{C}}_{11} = (\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{X})^{-1}$$

$$\hat{\mathbf{C}}_{21} = -\hat{\mathbf{G}}\mathbf{Z}'\hat{\mathbf{V}}^{-1}\mathbf{X}\hat{\mathbf{C}}_{11}$$

$$\hat{\mathbf{C}}_{22} = (\mathbf{Z}'\hat{\mathbf{R}}^{-1}\mathbf{Z} + \hat{\mathbf{G}}^{-1})^{-1} - \hat{\mathbf{C}}_{21}\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{Z}\hat{\mathbf{G}}$$

Note that $\hat{\mathbf{C}}_{11}$ is the familiar estimated generalized least squares formula for the variance-covariance matrix of $\hat{\beta}$.

Computational Method

Distributed Computing

Distributed computing refers to the use of multiple autonomous computers that communicate through a secure network. Distributed computing solves computational problems by dividing them into many tasks, each of which is solved by one or more computers. Each computer in this distributed environment is referred to as a node.

You can specify the number of nodes to use with the `NODES=` option in the **PERFORMANCE** statement. Specify `NODES=0` to force the execution to be done locally (often referred to as single-machine mode).

Multithreading

Threading refers to the organization of computational work into multiple tasks (processing units that can be scheduled by the operating system). A task is associated with a thread. Multithreading refers to the concurrent execution of threads. When multithreading is possible, substantial performance gains can be realized compared to sequential (single-threaded) execution.

The number of threads spawned by the HPLMIXED procedure is determined by the number of CPUs on a machine and can be controlled in the following ways:

You can specify the `NTHREADS=` option in the **PERFORMANCE** statement to determine the number of threads. This specification overrides the system option. Specify `NTHREADS=1` to force single-threaded execution.

The number of threads per machine is displayed in the “Performance Information” table, which is part of the default output. The HPLMIXED procedure allocates two threads per CPU.

The tasks multithreaded by the HPLMIXED procedure are primarily defined by dividing the data processed on a single machine among the threads—that is, the HPLMIXED procedure implements multithreading through a data-parallel model. For example, if the input data set has 1,000 observations and you are running

with four threads, then 250 observations are associated with each thread. All operations that require access to the data are then multithreaded. These operations include the following:

- variable levelization
- effect levelization
- formation of the crossproducts matrix
- the log-likelihood computation

In addition, operations on matrices such as sweeps might be multithreaded if the matrices are of sufficient size to realize performance benefits from managing multiple threads for the particular matrix operation.

Displayed Output

The following sections describe the output produced by PROC HPLMIXED. The output is organized into various tables, which are discussed in the order of their appearance.

Performance Information

The “Performance Information” table is produced by default. It displays information about the execution mode. For single-machine mode, the table displays the number of threads used. For distributed mode, the table displays the grid mode (symmetric or asymmetric), the number of compute nodes, and the number of threads per node.

If you specify the DETAILS option in the PERFORMANCE statement, PROC HPLMIXED also produces a “Timing” table that displays elapsed times (absolute and relative) for the main tasks of the procedure.

Model Information

The “Model Information” table describes the model, some of the variables it involves, and the method used in fitting it. The “Model Information” table also has a row labeled Fixed Effects SE Method. This row describes the method used to compute the approximate standard errors for the fixed-effects parameter estimates and related functions of them.

Class Level Information

The “Class Level Information” table lists the levels of every variable specified in the CLASS statement.

Dimensions

The “Dimensions” table lists the sizes of relevant matrices. This table can be useful in determining the requirements for CPU time and memory.

Number of Observations

The “Number of Observations” table shows the number of observations read from the data set and the number of observations used in fitting the model.

Optimization Information

The “Optimization Information” table displays important details about the optimization process.

The number of parameters that are updated in the optimization equals the number of parameters in this table minus the number of equality constraints. The number of constraints is displayed if you fix covariance parameters with the **HOLD=** option in the **PARMS** statement. The HPLMIXED procedure also lists the number of upper and lower boundary constraints. PROC HPLMIXED might impose boundary constraints for certain parameters, such as variance components and correlation parameters. If you specify the **HOLD=** option in the **PARMS** statement, covariance parameters have an upper and lower boundary equal to the parameter value.

Iteration History

The “Iteration History” table describes the optimization of the [restricted log likelihood or log likelihood](#). The function to be minimized (the *objective function*) is $-2l$ for ML and $-2l_R$ for REML; the column name of the objective function in the “Iteration History” table is “-2 Log Like” for ML and “-2 Res Log Like” for REML. The minimization is performed by using a ridge-stabilized Newton-Raphson algorithm, and the rows of this table describe the iterations that this algorithm takes in order to minimize the objective function.

The Evaluations column of the “Iteration History” table tells how many times the objective function is evaluated during each iteration.

The Criterion column of the “Iteration History” table is, by default, a relative Hessian convergence quantity given by

$$\frac{\mathbf{g}'_k \mathbf{H}_k^{-1} \mathbf{g}_k}{|f_k|}$$

where f_k is the value of the objective function at iteration k , $\mathbf{g}_k$ is the gradient (first derivative) of f_k , and $\mathbf{H}_k$ is the Hessian (second derivative) of f_k . If $\mathbf{H}_k$ is singular, then PROC HPLMIXED uses the following relative quantity:

$$\frac{\mathbf{g}'_k \mathbf{g}_k}{|f_k|}$$

To prevent division by $|f_k|$, specify the **ABSGCONV** option in the **PROC HPLMIXED** statement. To use a relative function or gradient criterion, specify the **FCONV** or **GCONV** option, respectively.

The Hessian criterion is considered superior to function and gradient criteria because it measures orthogonality rather than lack of progress (Bates et al. 1987). Provided that the initial estimate is feasible and the maximum number of iterations is not exceeded, the Newton-Raphson algorithm is considered to have converged when the criterion is less than the tolerance specified with the **FCONV** or **GCONV** option in the **PROC HPLMIXED** statement. The default tolerance is 1E-8. If convergence is not achieved, PROC HPLMIXED displays the estimates of the parameters at the last iteration.

A convergence criterion that is missing indicates that a boundary constraint has been dropped; it is usually not a cause for concern.

Convergence Status

The “Convergence Status” table displays the status of the iterative estimation process at the end of the optimization. The status appears as a message in the listing, and this message is repeated in the log. The ODS object “ConvergenceStatus” also contains several nonprinting columns that can be helpful in checking the success of the iterative process, in particular during batch processing. The Status variable takes on the value 0 for a successful convergence (even if the Hessian matrix might not be positive definite). The values 1 and 2 of the Status variable indicate lack of convergence and infeasible initial parameter values, respectively. The variable pdG can be used to check whether the **G** matrix is positive definite.

For models that are not fit iteratively, such as models without random effects or when the **NOITER** option is in effect, the “Convergence Status” is not produced.

Covariance Parameter Estimates

The “Covariance Parameter Estimates” table contains the estimates of the parameters in **G** and **R**. (See the section “[Estimating Covariance Parameters in the Mixed Model](#)” on page 4302.) Their values are labeled in the table along with Subject information if applicable. The estimates are displayed in the Estimate column and are the results of either the REML or the ML estimation method.

Fit Statistics

The “Fit Statistics” table provides some statistics about the estimated mixed model. Expressions for -2 times the log likelihood are provided in the section “[Estimating Covariance Parameters in the Mixed Model](#)” on page 4302. If the log likelihood is an extremely large number, then PROC HPLMIXED has deemed the estimated **V** matrix to be singular. In this case, all subsequent results should be viewed with caution.

In addition, the “Fit Statistics” table lists three information criteria: AIC, AICC, and BIC. All these criteria are in smaller-is-better form and are described in [Table 54.9](#).

Table 54.9 Information Criteria

Criterion	Formula	Reference
AIC	$-2\ell + 2d$	Akaike (1974)
AICC	$-2\ell + 2dn^*/(n^* - d - 1)$	Hurvich and Tsai (1989) Burnham and Anderson (1998)
BIC	$-2\ell + d \log n$ for $n > 0$	Schwarz (1978)

Here ℓ denotes the maximum value of the (possibly restricted) log likelihood; d is the dimension of the model; and n equals the number of effective subjects as displayed in the “Dimensions” table, unless this value equals 1, in which case n equals the number of levels of the first random effect specified in the first **RANDOM** statement or the number of levels of the interaction of the first random effect with noncommon subject effect specified in the first **RANDOM** statement. If the number of effective subjects equals 1 and you have no **RANDOM** statements, then n equals the number of valid observations for maximum likelihood estimation and $n - p$ for restricted maximum likelihood estimation, where p equals the rank of $\mathbf{X}$. For AICC (a finite-sample corrected version of AIC), n^* equals the number of valid observations for maximum likelihood estimation and $n - p$ equals the number of valid observations for restricted maximum likelihood estimation, unless this number is less than $d + 2$, in which case it equals $d + 2$. When $n = 0$, the value of the BIC is -2ℓ . For restricted likelihood estimation, d equals q , the effective number of estimated covariance parameters. For maximum likelihood estimation, d equals $q + p$.

Timing Information

If you specify the **DETAILS** option in the **PERFORMANCE** statement, the procedure also produces a “Timing” table in which the elapsed time for each main task of the procedure is displayed.

ODS Table Names

Each table created by PROC HPLMIXED has a name associated with it, and you must use this name to refer to the table when you use ODS statements. These names are listed in Table 54.10.

Table 54.10 ODS Tables Produced by PROC HPLMIXED

Table Name	Description	Required Statement / Option
ClassLevels	Level information from the CLASS statement	Default output
ConvergenceStatus	Convergence status	Default output
CovParms	Estimated covariance parameters	Default output
Dimensions	Dimensions of the model	Default output
FitStatistics	Fit statistics	Default output
IterHistory	Iteration history	Default output
ModelInfo	Model information	Default output
NObs	Number of observations read and used	Default output
OptInfo	Optimization information	Default output
ParmSearch	Parameter search values	PARMS
PerformanceInfo	Information about high-performance computing environment	Default output
Ranks	Rank of designed matrix $\mathbf{X}$	PROC HPLMIXED RANKS
SolutionF	Fixed-effects solution vector	MODEL / S
SolutionR	Random-effects solution vector	RANDOM / S
Timing	Timing breakdown by task	DETAILS option in the PERFORMANCE statement

Examples: HPLMIXED Procedure

Example 54.1: Computing BLUPs for a Large Number of Subjects

Suppose you are using health measurements on patients treated by each medical center to monitor the performance of those centers. Different measurements within each patient are correlated, and there is enough data to fit the parameters of an unstructured covariance model for this correlation. In fact, long experience with historical data provides you with values for the covariance model that are essentially known, and the task is to apply these known values in order to compute best linear unbiased predictors (BLUPs) of the random effect of medical center. You can use these BLUPs to determine the best and worst performing medical centers, adjusting for other factors, on a weekly basis. Another reason why you want to do this with fixed values for the covariance parameters is to make the week-to-week BLUPs more comparable.

Although you cannot use the REPEATED statement in PROC HPLMIXED to fit models in this release, you can use it to compute BLUPs for such models with known values of the variance parameters. For illustration, the following statements create a simulated data set of a given week's worth of patient health measurements across 100 different medical centers. Measurements at three different times are simulated for each patient, and each center has about 50 patients. The simulated model includes a fixed gender effect, a random effect due to center, and covariance between different measurements on the same patient.

```
%let NCenter = 100;
%let NPatient = %eval(&NCenter*50);
%let NTime = 3;
%let SigmaC = 2.0;
%let SigmaP = 4.0;
%let SigmaE = 8.0;
%let Seed = 12345;

data WeekSim;
  keep Gender Center Patient Time Measurement;
  array PGender{&NPatient};
  array PCenter{&NPatient};
  array PEffect{&NPatient};
  array CEffect{&NCenter};
  array GEeffect{2};

  do Center = 1 to &NCenter;
    CEffect{Center} = sqrt(&SigmaC)*rannor(&Seed);
  end;

  GEeffect{1} = 10*ranuni(&Seed);
  GEeffect{2} = 10*ranuni(&Seed);

  do Patient = 1 to &NPatient;
    PGender{Patient} = 1 + int(2 * ranuni(&Seed));
    PCenter{Patient} = 1 + int(&NCenter*ranuni(&Seed));
    PEffect{Patient} = sqrt(&SigmaP)*rannor(&Seed);
  end;
```

```

do Patient = 1 to &NPatient;
  Gender = PGender{Patient};
  Center = PCenter{Patient};
  Mean = 1 + GEffect{Gender} + CEffect{Center} + PEffect{Patient};
  do Time = 1 to &nTime;
    Measurement = Mean + sqrt(&SigmaE)*rannor(&Seed);
    output;
  end;
end;
run;

```

Suppose that the known values for the covariance parameters are

$$\begin{aligned}
 \text{Var}(\text{Center}) &= 1.7564 \\
 \text{Cov}(\text{Patient}) &= \begin{bmatrix} 11.4555 & 3.6883 & 4.5951 \\ 3.6883 & 11.2071 & 3.6311 \\ 4.5951 & 3.6311 & 12.1050 \end{bmatrix}
 \end{aligned}$$

Incidentally, these are not precisely the same estimates you would get if you estimated these parameters based on the preceding data (for example, with the HPLMIXED procedure).

The following statements use PROC HPLMIXED to compute the BLUPs for the random medical center effect. Instead of simply displaying them (as PROC HPMIXED does), PROC HPLMIXED sorts them and displays the five highest and lowest values. To run these statements successfully, you need to set the macro variables GRIDHOST and GRIDINSTALLLOC to resolve to appropriate values, or you can replace the references to macro variables with appropriate values.

```

ods listing close;
proc hplmixed data=WeekSim blup;
  performance host="&GRIDHOST" install="&GRIDINSTALLLOC" nodes=20;
  class Gender Center Patient Time;
  model Measurement = Gender;
  random Center / s;
  repeated Time / sub=Patient type=un;
  parms 1.7564
        11.4555
        3.6883 11.2071
        4.5951 3.6311 12.1050;
  ods output SolutionR=BLUPs;
run;
ods listing;

proc sort data=BLUPs;
  by Estimate;
run;

data BLUPs; set BLUPs;
  Rank = _N_;
run;

```

```
proc print data=BLUPs;
  where ((Rank <= 5) | (Rank >= 96));
  var Center Estimate;
run;
```

Three parts of the PROC HPLMIXED syntax are required in order to compute BLUPs for this model: the BLUP option in the HPLMIXED statement, the REPEATED statement, and the PARMS statement with fixed values for all parameters. The resulting values of the best and worst performing medical centers for this week are shown in [Output 54.1.1](#). Apparently, for this week's data, medical center 54 had the most decreasing effect, and medical center 48 the most increasing effect, on patient measurements overall.

Output 54.1.1 Highest and Lowest Medical Center BLUPs

Obs	Center	Estimate
1	54	-2.9369
2	7	-2.4614
3	50	-2.2467
4	51	-2.2281
5	93	-2.1644
96	26	2.1603
97	99	2.2718
98	44	2.4222
99	60	2.6089
100	48	2.6443

References

- Akaike, H. (1974). "A New Look at the Statistical Model Identification." *IEEE Transactions on Automatic Control* AC-19:716–723.
- Bates, D. M., Lindstrom, M. J., Wahba, G., and Yandell, B. S. (1987). "GCVPACK—Routines for Generalized Cross Validation." *Communications in Statistics—Simulation and Computation* 16:263–297.
- Burdick, R. K., and Graybill, F. A. (1992). *Confidence Intervals on Variance Components*. New York: Marcel Dekker.
- Burnham, K. P., and Anderson, D. R. (1998). *Model Selection and Inference: A Practical Information-Theoretic Approach*. New York: Springer-Verlag.
- Carlin, B. P., and Louis, T. A. (1996). *Bayes and Empirical Bayes Methods for Data Analysis*. London: Chapman & Hall.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). "Maximum Likelihood from Incomplete Data via the EM Algorithm." *Journal of the Royal Statistical Society, Series B* 39:1–38.
- Fuller, W. A. (1976). *Introduction to Statistical Time Series*. New York: John Wiley & Sons.
- Hartley, H. O., and Rao, J. N. K. (1967). "Maximum-Likelihood Estimation for the Mixed Analysis of Variance Model." *Biometrika* 54:93–108.

- Harville, D. A. (1977). "Maximum Likelihood Approaches to Variance Component Estimation and to Related Problems." *Journal of the American Statistical Association* 72:320–338.
- Harville, D. A. (1988). "Mixed-Model Methodology: Theoretical Justifications and Future Directions." In *Proceedings of the Statistical Computing Section*, 41–49. Alexandria, VA: American Statistical Association.
- Harville, D. A. (1990). "BLUP (Best Linear Unbiased Prediction), and Beyond." In *Advances in Statistical Methods for Genetic Improvement of Livestock*, edited by D. Gianola and K. Hammond, 239–276. Vol. 18 of Advanced Series in Agricultural Sciences. Berlin: Springer-Verlag.
- Henderson, C. R. (1984). *Applications of Linear Models in Animal Breeding*. Guelph, ON: University of Guelph.
- Henderson, C. R. (1990). "Statistical Method in Animal Improvement: Historical Overview." In *Advances in Statistical Methods for Genetic Improvement of Livestock*, 1–14. New York: Springer-Verlag.
- Hurvich, C. M., and Tsai, C.-L. (1989). "Regression and Time Series Model Selection in Small Samples." *Biometrika* 76:297–307.
- 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.
- Jennrich, R. I., and Schluchter, M. D. (1986). "Unbalanced Repeated-Measures Models with Structured Covariance Matrices." *Biometrics* 42:805–820.
- Kenward, M. G. (1987). "A Method for Comparing Profiles of Repeated Measurements." *Journal of the Royal Statistical Society, Series C* 36:296–308.
- Laird, N. M., Lange, N. T., and Stram, D. O. (1987). "Maximum Likelihood Computations with Repeated Measures: Application of the EM Algorithm." *Journal of the American Statistical Association* 82:97–105.
- Laird, N. M., and Ware, J. H. (1982). "Random-Effects Models for Longitudinal Data." *Biometrics* 38:963–974.
- Lindstrom, M. J., and Bates, D. M. (1988). "Newton-Raphson and EM Algorithms for Linear Mixed-Effects Models for Repeated-Measures Data." *Journal of the American Statistical Association* 83:1014–1022.
- Littell, R. C., Milliken, G. A., Stroup, W. W., Wolfinger, R. D., and Schabenberger, O. (2006). *SAS for Mixed Models*. 2nd ed. Cary, NC: SAS Institute Inc.
- Little, R. J. A. (1995). "Modeling the Drop-Out Mechanism in Repeated-Measures Studies." *Journal of the American Statistical Association* 90:1112–1121.
- Macchiavelli, R. E., and Arnold, S. F. (1994). "Variable Order Ante-dependence Models." *Communications in Statistics—Theory and Methods* 23:2683–2699.
- McLean, R. A., and Sanders, W. L. (1988). "Approximating Degrees of Freedom for Standard Errors in Mixed Linear Models." In *Proceedings of the Statistical Computing Section*, 50–59. Alexandria, VA: American Statistical Association.
- McLean, R. A., Sanders, W. L., and Stroup, W. W. (1991). "A Unified Approach to Mixed Linear Models." *American Statistician* 45:54–64.

- Milliken, G. A., and Johnson, D. E. (1992). *Designed Experiments*. Vol. 1 of Analysis of Messy Data. Reprint edition. New York: Chapman & Hall.
- Patel, H. I. (1991). "Analysis of Incomplete Data from a Clinical Trial with Repeated Measurements." *Biometrika* 78:609–619.
- Patterson, H. D., and Thompson, R. (1971). "Recovery of Inter-block Information When Block Sizes Are Unequal." *Biometrika* 58:545–554.
- Robinson, G. K. (1991). "That BLUP Is a Good Thing: The Estimation of Random Effects." *Statistical Science* 6:15–51.
- Rubin, D. B. (1976). "Inference and Missing Data." *Biometrika* 63:581–592.
- Schwarz, G. (1978). "Estimating the Dimension of a Model." *Annals of Statistics* 6:461–464.
- Searle, S. R. (1971). *Linear Models*. New York: John Wiley & Sons.
- Searle, S. R. (1982). *Matrix Algebra Useful for Statisticians*. New York: John Wiley & Sons.
- Searle, S. R., Casella, G., and McCulloch, C. E. (1992). *Variance Components*. New York: John Wiley & Sons.
- Serfling, R. J. (1980). *Approximation Theorems of Mathematical Statistics*. New York: John Wiley & Sons.
- Snedecor, G. W., and Cochran, W. G. (1980). *Statistical Methods*. 7th ed. Ames: Iowa State University Press.
- Steel, R. G. D., Torrie, J. H., and Dickey, D. A. (1997). *Principles and Procedures of Statistics: A Biometrical Approach*. 3rd ed. New York: McGraw-Hill.
- Verbeke, G., and Molenberghs, G., eds. (1997). *Linear Mixed Models in Practice: A SAS-Oriented Approach*. New York: Springer.
- Verbeke, G., and Molenberghs, G. (2000). *Linear Mixed Models for Longitudinal Data*. New York: Springer.
- Winer, B. J. (1971). *Statistical Principles in Experimental Design*. 2nd ed. New York: McGraw-Hill.
- Wolfinger, R. D. (1996). "Heterogeneous Variance-Covariance Structures for Repeated Measures." *Journal of Agricultural, Biological, and Environmental Statistics* 1:205–230.
- Wolfinger, R. D., Tobias, R. D., and Sall, J. (1994). "Computing Gaussian Likelihoods and Their Derivatives for General Linear Mixed Models." *SIAM Journal on Scientific Computing* 15:1294–1310.

Subject Index

- 2D geometric anisotropic structure
 - HPLMIXED procedure, [4291](#)
- Akaike's information criterion
 - HPLMIXED procedure, [4308](#)
- Akaike's information criterion (finite sample corrected version)
 - HPLMIXED procedure, [4308](#)
- alpha level
 - HPLMIXED procedure, [4285](#), [4290](#)
- ANTE(1) structure
 - HPLMIXED procedure, [4291](#)
- antedependence structure
 - HPLMIXED procedure, [4291](#)
- AR(1) structure
 - HPLMIXED procedure, [4291](#)
- autoregressive moving-average structure
 - HPLMIXED procedure, [4291](#)
- autoregressive structure
 - HPLMIXED procedure, [4291](#)
- banded Toeplitz structure
 - HPLMIXED procedure, [4291](#)
- BLUE
 - HPLMIXED procedure, [4304](#)
- BLUP
 - HPLMIXED procedure, [4304](#)
- boundary constraints
 - HPLMIXED procedure, [4288](#), [4289](#)
- class level
 - HPLMIXED procedure, [4283](#), [4306](#)
- compound symmetry structure
 - example (HPLMIXED), [4298](#)
 - HPLMIXED procedure, [4291](#)
- computational method
 - HPLMIXED procedure, [4305](#)
- constraints
 - boundary (HPLMIXED), [4288](#), [4289](#)
- convergence criterion
 - HPLMIXED procedure, [4280–4282](#), [4307](#)
- convergence status
 - HPLMIXED procedure, [4308](#)
- covariance parameter estimates
 - HPLMIXED procedure, [4308](#)
- covariance structure
 - antedependence (HPLMIXED), [4294](#)
 - autoregressive (HPLMIXED), [4294](#)
 - autoregressive moving average (HPLMIXED), [4294](#)
 - banded (HPLMIXED), [4295](#)
 - compound symmetry (HPLMIXED), [4294](#)
 - equi-correlation (HPLMIXED), [4294](#)
 - examples (HPLMIXED), [4292](#)
 - factor-analytic (HPLMIXED), [4294](#)
 - heterogeneous autoregressive (HPLMIXED), [4294](#)
 - heterogeneous compound symmetry (HPLMIXED), [4294](#)
 - heterogeneous Toeplitz (HPLMIXED), [4295](#)
 - HPLMIXED procedure, [4274](#), [4291](#)
 - Huynh-Feldt (HPLMIXED), [4295](#)
 - simple (HPLMIXED), [4295](#)
 - Toeplitz (HPLMIXED), [4295](#)
 - unstructured (HPLMIXED), [4295](#)
 - unstructured, correlation (HPLMIXED), [4295](#)
 - variance components (HPLMIXED), [4295](#)
- degrees of freedom
 - HPLMIXED procedure, [4286](#)
 - infinite (HPLMIXED), [4286](#)
 - method (HPLMIXED), [4286](#)
 - residual method (HPLMIXED), [4286](#)
- dimension information
 - HPLMIXED procedure, [4306](#)
- direct product structure
 - HPLMIXED procedure, [4291](#)
- displayed output
 - HPLMIXED procedure, [4306](#)
- distributed computing
 - HPLMIXED procedure, [4305](#)
- effect
 - name length (HPLMIXED), [4283](#)
- estimation
 - mixed model (HPLMIXED), [4302](#)
- estimation methods
 - HPLMIXED procedure, [4283](#)
- examples, HPLMIXED
 - compound symmetry, G-side setup, [4299](#)
 - compound symmetry, R-side setup, [4299](#)
 - holding covariance parameters fixed, [4288](#)
 - specifying lower bounds, [4288](#)
 - split-plot design, [4300](#)
 - subject and no-subject formulation, [4299](#)
 - unstructured covariance, G-side, [4291](#)

- factor analytic structures
 - HPLMIXED procedure, 4291
- fixed effects
 - HPLMIXED procedure, 4274
- fixed-effects parameters
 - HPLMIXED procedure, 4298
- G matrix
 - HPLMIXED procedure, 4274, 4289, 4298
- general linear covariance structure
 - HPLMIXED procedure, 4291
- generalized inverse, 4304
- gradient
 - HPLMIXED procedure, 4307
- growth curve analysis
 - example (HPLMIXED), 4298
- Hessian matrix
 - HPLMIXED procedure, 4307
- heterogeneous
 - AR(1) structure (HPLMIXED), 4291
 - compound-symmetry structure (HPLMIXED), 4291
 - covariance structure (HPLMIXED), 4295
 - Toeplitz structure (HPLMIXED), 4291
- HPLMIXED procedure, 4272
 - 2D geometric anisotropic structure, 4291
 - Akaike's information criterion, 4308
 - Akaike's information criterion (finite sample corrected version), 4308
 - alpha level, 4285, 4290
 - ANTE(1) structure, 4291
 - antedependence structure, 4291
 - AR(1) structure, 4291
 - ARMA structure, 4291
 - autoregressive moving-average structure, 4291
 - autoregressive structure, 4291
 - banded Toeplitz structure, 4291
 - BLUE, 4281, 4304
 - BLUP, 4281, 4304
 - boundary constraints, 4288, 4289
 - chi-square test, 4286
 - class level, 4283, 4306
 - compound symmetry structure, 4291, 4298
 - computational method, 4305
 - confidence interval, 4290
 - confidence limits, 4286, 4290
 - convergence criterion, 4280–4282, 4307
 - convergence status, 4308
 - covariance parameter estimates, 4308
 - covariance structure, 4274, 4291, 4292
 - degrees of freedom, 4286
 - dimension information, 4306
 - direct product structure, 4291
 - displayed output, 4306
 - distributed computing, 4305
 - EBLUPs, 4290, 4305
 - effect name length, 4283
 - estimation methods, 4283
 - factor analytic structures, 4291
 - fitting information, 4308
 - fixed effects, 4274
 - fixed-effects parameters, 4286, 4298
 - function-based convergence criteria, 4280, 4281
 - G matrix, 4274, 4289, 4298
 - general linear covariance structure, 4291
 - generalized inverse, 4304
 - gradient, 4307
 - gradient-based convergence criteria, 4281, 4282
 - grid search, 4287
 - growth curve analysis, 4298
 - Hessian matrix, 4307
 - heterogeneous AR(1) structure, 4291
 - heterogeneous compound-symmetry structure, 4291
 - heterogeneous covariance structures, 4295
 - heterogeneous Toeplitz structure, 4291
 - Huynh-Feldt structure, 4291
 - infinite degrees of freedom, 4286
 - infinite likelihood, 4296
 - initial values, 4287
 - input data sets, 4281
 - intercept effect, 4286, 4290
 - iteration history, 4307
 - iterations, 4307
 - Kronecker product structure, 4291
 - linear covariance structure, 4291
 - Matérn covariance structure, 4291
 - matrix notation, 4297
 - mixed model, 4298
 - mixed model equations, 4303
 - mixed model theory, 4297
 - model information, 4306
 - multi-threading, 4289
 - multithreading, 4305
 - Newton-Raphson algorithm, 4303
 - number of observations, 4306
 - ODS table names, 4309
 - optimization information, 4307
 - parameter constraints, 4288
 - parameter-based convergence criteria, 4284
 - performance information, 4306
 - R matrix, 4274, 4296, 4298
 - random effects, 4274, 4289
 - random-effects parameters, 4290, 4298
 - repeated measures, 4296
 - residual method, 4286
 - restricted maximum likelihood (REML), 4273

- ridging, 4303
- Schwarz's Bayesian information criterion, 4308
- spatial anisotropic exponential structure, 4291
- split-plot design, 4300
- standard linear model, 4274
- statement positions, 4278
- subject effect, 4290, 4296
- summary of commands, 4278
- table names, 4309
- timing, 4309
- Toeplitz structure, 4291
- unstructured correlations, 4291
- unstructured covariance matrix, 4291
- variance components, 4291
- variance ratios, 4288
- Huynh-Feldt
 - structure (HPLMIXED), 4291
- infinite likelihood
 - HPLMIXED procedure, 4296
- initial values
 - HPLMIXED procedure, 4287
- iteration history
 - HPLMIXED procedure, 4307
- iterations
 - history (HPLMIXED), 4307
- Kronecker product structure
 - HPLMIXED procedure, 4291
- linear covariance structure
 - HPLMIXED procedure, 4291
- Matérn covariance structure
 - HPLMIXED procedure, 4291
- matrix
 - notation, theory (HPLMIXED), 4297
- maximum likelihood estimation
 - mixed model (HPLMIXED), 4303
- mixed model (HPLMIXED)
 - estimation, 4302
 - formulation, 4298
 - maximum likelihood estimation, 4303
 - notation, 4274
 - objective function, 4307
 - theory, 4297
- mixed model equations
 - HPLMIXED procedure, 4303
- model information
 - HPLMIXED procedure, 4306
- multi-threading
 - HPLMIXED procedure, 4289
- multithreading
 - HPLMIXED procedure, 4305

- Newton-Raphson algorithm
 - HPLMIXED procedure, 4303
- number of observations
 - HPLMIXED procedure, 4306
- objective function
 - mixed model (HPLMIXED), 4307
- optimization information
 - HPLMIXED procedure, 4307
- options summary
 - MODEL statement (HPLMIXED), 4285
 - PROC HPLMIXED statement, 4279
 - RANDOM statement (HPLMIXED), 4290
 - REPEATED statement (HPLMIXED), 4296
- parameter constraints
 - HPLMIXED procedure, 4288
- performance information
 - HPLMIXED procedure, 4306
- R matrix
 - HPLMIXED procedure, 4274, 4296, 4298
- random effects
 - HPLMIXED procedure, 4274, 4289
- random-effects parameters
 - HPLMIXED procedure, 4298
- repeated measures
 - HPLMIXED procedure, 4296
- residual maximum likelihood (REML)
 - HPLMIXED procedure, 4303
- restricted maximum likelihood
 - HPLMIXED procedure, 4273
- restricted maximum likelihood (REML)
 - HPLMIXED procedure, 4303
- ridging
 - HPLMIXED procedure, 4303
- Schwarz's Bayesian information criterion
 - HPLMIXED procedure, 4308
- spatial anisotropic exponential structure
 - HPLMIXED procedure, 4291
- split-plot design
 - HPLMIXED procedure, 4300
- standard linear model
 - HPLMIXED procedure, 4274
- subject effect
 - HPLMIXED procedure, 4290, 4296
- summary of commands
 - HPLMIXED procedure, 4278
- table names
 - HPLMIXED procedure, 4309
- timing
 - HPLMIXED procedure, 4309
- Toeplitz structure

HPLMIXED procedure, [4291](#)

unstructured correlations

HPLMIXED procedure, [4291](#)

unstructured covariance matrix

HPLMIXED procedure, [4291](#)

variance components

HPLMIXED procedure, [4291](#)

variance ratios

HPLMIXED procedure, [4288](#)

Syntax Index

- ABSCONV option
 - PROC HPLMIXED statement, [4280](#)
- ABSFCNV option
 - PROC HPLMIXED statement, [4280](#)
- ABSGCONV option
 - PROC HPLMIXED statement, [4281](#)
- ABSGTOL option
 - PROC HPLMIXED statement, [4281](#)
- ABSOLUTE option
 - PROC HPLMIXED statement, [4307](#)
- ABSTOL option
 - PROC HPLMIXED statement, [4280](#)
- ALLSTATS option
 - OUTPUT statement (HPLMIXED), [4287](#)
- ALPHA= option
 - RANDOM statement (HPLMIXED), [4290](#)
- BLUP option
 - PROC HPLMIXED statement, [4281](#)
- CL option
 - MODEL statement (HPLMIXED), [4286](#)
 - RANDOM statement (HPLMIXED), [4290](#)
- CLASS statement
 - HPLMIXED procedure, [4284](#), [4306](#)
- CONVF option
 - PROC HPLMIXED statement, [4307](#)
- CONVG option
 - PROC HPLMIXED statement, [4307](#)
- CONVH option
 - PROC HPLMIXED statement, [4307](#)
- DATA= option
 - PROC HPLMIXED statement, [4281](#)
- DDFM= option
 - MODEL statement (HPLMIXED), [4286](#)
- EQCONS= option
 - PARMS statement (HPLMIXED), [4288](#)
- FCONV option
 - PROC HPLMIXED statement, [4281](#)
- FTOL option
 - PROC HPLMIXED statement, [4281](#)
- GCONV option
 - PROC HPLMIXED statement, [4282](#)
- GTOL option
 - PROC HPLMIXED statement, [4282](#)
- HOLD= option
 - PARMS statement (HPLMIXED), [4288](#)
- HPLMIXED procedure, [4278](#)
 - ID statement, [4285](#)
 - OUTPUT statement, [4286](#)
 - PERFORMANCE statement, [4289](#)
 - PROC HPLMIXED statement, [4279](#)
 - syntax, [4278](#)
- HPLMIXED procedure, CLASS statement, [4284](#), [4306](#)
- HPLMIXED procedure, ID statement, [4285](#)
- HPLMIXED procedure, MODEL statement, [4285](#)
 - CL option, [4286](#)
 - DDFM= option, [4286](#)
 - NOINT option, [4286](#)
 - SOLUTION option, [4286](#)
- HPLMIXED procedure, OUTPUT statement, [4286](#)
 - ALLSTATS option, [4287](#)
 - OUT= option, [4287](#)
- HPLMIXED procedure, PARMS statement, [4287](#)
 - EQCONS= option, [4288](#)
 - HOLD= option, [4288](#)
 - LOWERB= option, [4288](#)
 - NOITER option, [4289](#)
 - PARMSDATA= option, [4289](#)
 - PDATA= option, [4289](#)
 - UPPERB= option, [4289](#)
- HPLMIXED procedure, PERFORMANCE statement, [4289](#)
- HPLMIXED procedure, PROC HPLMIXED statement
 - ABSCONV option, [4280](#)
 - ABSFCNV option, [4280](#)
 - ABSFTOL option, [4280](#)
 - ABSGCONV option, [4281](#)
 - ABSGTOL option, [4281](#)
 - ABSOLUTE option, [4307](#)
 - ABSTOL option, [4280](#)
 - BLUP option, [4281](#)
 - CONVF option, [4307](#)
 - CONVG option, [4307](#)
 - CONVH option, [4307](#)
 - DATA= option, [4281](#)
 - FCONV option, [4281](#)
 - FTOL option, [4281](#)
 - GCONV option, [4282](#)
 - GTOL option, [4282](#)
 - MAXCLPRINT= option, [4282](#)
 - MAXFUNC= option, [4282](#)
 - MAXITER= option, [4282](#), [4283](#)

- MAXTIME= option, 4283
- METHOD= option, 4283
- NAMELEN= option, 4283
- NOCLPRINT option, 4283
- NOPRINT option, 4283
- RANKS option, 4283
- SINGCHOL= option, 4283
- SINGSWEEP= option, 4283
- SINGULAR= option, 4283
- TECHNIQUE= option, 4283
- XCONV option, 4284
- XTOL option, 4284
- HPLMIXED procedure, RANDOM statement, 4289
 - ALPHA= option, 4290
 - CL option, 4290
 - SOLUTION option, 4290
 - SUBJECT= option, 4290
 - TYPE= option, 4291
- HPLMIXED procedure, REPEATED statement, 4296
 - SUBJECT= option, 4296
 - TYPE= option, 4296
- HPLMIXEDC procedure, PROC HPLMIXED statement, 4279
- ID statement
 - HPLMIXED procedure, 4285
- LOWERB= option
 - PARMS statement (HPLMIXED), 4288
- MAXCLPRINT= option
 - PROC HPLMIXED statement, 4282
- MAXFUNC= option
 - PROC HPLMIXED statement, 4282
- MAXITER= option
 - PROC HPLMIXED statement, 4282, 4283
- MAXTIME= option
 - PROC HPLMIXED statement, 4283
- METHOD= option
 - PROC HPLMIXED statement, 4283
- MODEL statement
 - HPLMIXED procedure, 4285
- NAMELEN= option
 - PROC HPLMIXED statement, 4283
- NOCLPRINT option
 - PROC HPLMIXED statement, 4283
- NOINT option
 - MODEL statement (HPLMIXED), 4286
- NOITER option
 - PARMS statement (HPLMIXED), 4289
- NOPRINT option
 - PROC HPLMIXED statement, 4283
- OUT= option
 - OUTPUT statement (HPLMIXED), 4287
- OUTPUT statement
 - HPLMIXED procedure, 4286
- PARMS statement
 - HPLMIXED procedure, 4287
- PARMSDATA= option
 - PARMS statement (HPLMIXED), 4289
- PDATA= option
 - PARMS statement (HPLMIXED), 4289
- PERFORMANCE statement
 - HPLMIXED procedure, 4289
- PROC HPLMIXED statement, *see* HPLMIXED procedure
 - HPLMIXED procedure, 4279
- RANDOM statement
 - HPLMIXED procedure, 4289
- RANKS option
 - PROC HPLMIXED statement, 4283
- REPEATED statement
 - HPLMIXED procedure, 4296
- SINGCHOL= option
 - PROC HPLMIXED statement, 4283
- SINGSWEEP= option
 - PROC HPLMIXED statement, 4283
- SINGULAR= option
 - PROC HPLMIXED statement, 4283
- SOLUTION option
 - MODEL statement (HPLMIXED), 4286
 - RANDOM statement (HPLMIXED), 4290
- SUBJECT= option
 - RANDOM statement (HPLMIXED), 4290
 - REPEATED statement (HPLMIXED), 4296
- TECHNIQUE= option
 - PROC HPLMIXED statement, 4283
- TYPE= option
 - RANDOM statement (HPLMIXED), 4291
 - REPEATED statement (HPLMIXED), 4296
- UPPERB= option
 - PARMS statement (HPLMIXED), 4289
- XCONV option
 - PROC HPLMIXED statement, 4284
- XTOL option
 - PROC HPLMIXED statement, 4284