



Linear mixed-effects modeling in SPSS[®]

AN INTRODUCTION TO THE MIXED PROCEDURE

Table of contents

Introduction	3
Data preparation for MIXED	3
Fitting fixed-effects models with iid residual errors	6
Example 1 — Fixed-effects model using MIXED	7
Example 2 — Fixed-effects model using GLM	8
Fitting fixed-effects models with non-iid residual errors	8
Example 3 — Fixed-effects model with correlated residual errors	9
Fitting simple mixed-effects models (balanced design)	10
Example 4 — Simple mixed-effects model with balanced design using MIXED ..	11
Example 5 — Simple mixed-effects model with balanced design using GLM	12
Example 6 — Variance components model with balanced design	13
Fitting simple mixed-effects models (unbalanced design)	14
Example 4a — Mixed-effects model with unbalanced design using MIXED	15
Example 5a — Mixed-effects model with unbalanced design using GLM	16
Example 6a — Variance components model with unbalanced design	17
Fitting mixed-effects models with subjects	17
Example 7 — Fitting random effect*subject interaction using MIXED	18
Example 8 — Fitting random effect*subject interaction using GLM	18
Example 9 — Fitting random effect*subject interaction using	
<i>SUBJECT</i> specification	19
Example 10 — Using <i>COVTYPE</i> in a random-effects model	20
Multilevel analysis	21
Example 11 — Multilevel mixed-effects model	21
Custom hypothesis tests	23
Example 12 — Custom hypothesis testing in mixed-effects model	23
Covariance structure selection	24
Information criteria for Example 9	25
Information criteria for Example 10	25
Reference	26
About SPSS Inc.	26

Introduction

The linear mixed-effects model (MIXED) procedure in SPSS enables you to fit linear mixed-effects models to data sampled from normal distributions. Recent texts, such as those by McCulloch and Searle (2000) and Verbeke and Molenberghs (2000), comprehensively reviewed mixed-effects models. The MIXED procedure fits models more general than those of the general linear model (GLM) procedure and it encompasses all models in the variance components (VARCOMP) procedure. This report illustrates the types of models that MIXED handles. We begin with an explanation of simple models that can be fitted using GLM and VARCOMP, to show how they are translated into MIXED. We then proceed to fit models that are unique to MIXED.

The major capabilities that differentiate MIXED from GLM are that MIXED handles correlated data and unequal variances. Correlated data are very common in such situations as repeated measurements of survey respondents or experimental subjects. MIXED also handles more complex situations in which experimental units are nested in a hierarchy. MIXED can, for example, process data obtained from a sample of students selected from a sample of schools in a district.

In a linear mixed-effects model, responses from a subject are thought to be the sum (linear) of so-called fixed and random effects. If an effect, such as a medical treatment, affects the population mean, it is fixed. If an effect is associated with a sampling procedure (e.g., subject effect), it is random. In a mixed-effects model, random effects contribute only to the covariance structure of the data. The presence of random effects, however, often introduces correlations between cases as well. Though the fixed effect is the primary interest in most studies or experiments, it is necessary to adjust for the covariance structure of the data. The adjustment made in procedures like GLM-Univariate is often not appropriate because it assumes the independence of the data.

The MIXED procedure solves these problems by providing the tools necessary to estimate fixed and random effects in one model. MIXED is based, furthermore, on maximum likelihood (ML) and restricted maximum likelihood (REML) methods, versus the analysis of variance (ANOVA) methods in GLM. ANOVA methods produce only an optimum estimator (minimum variance) for balanced designs, whereas ML and REML yield asymptotically efficient estimators for balanced and unbalanced designs. ML and REML thus present a clear advantage over ANOVA methods in modeling real data, since data are often unbalanced. The asymptotic normality of ML and REML estimators, furthermore, conveniently allows us to make inferences on the covariance parameters of the model, which is difficult to do in GLM.

Data preparation for MIXED

Many datasets store repeated observations on a sample of subjects in “one subject per row” format. MIXED, however, expects that observations from a subject are encoded in separate rows. To illustrate, we select a subset of cases from the data that appear in Potthoff and Roy (1964). The data shown in Figure 1 on the next page encode, in one row, three repeated measurements of a dependent variable (“dist1” to “dist3”) from a subject observed at different ages (“age1” to “age3”).

Figure 1 shows the SPSS Data Editor window for a file named 'Growth_subj.sav'. The data is organized by subject (sub.id) and age (age1, age2, age3). The 'dist' variables (dist1, dist2, dist3) represent measurements at different ages.

	sub.id	gender	age1	age2	age3	dist1	dist2	dist3
1	1	F	10	12	14	20.0	21.5	23.0
2	2	F	10	12	14	21.5	24.0	25.5
3	3	F	10	12	14	24.0	24.5	26.0
4	4	F	10	12	14	24.5	25.0	26.5
5	5	F	10	12	14	23.0	22.5	23.5
6	6	M	10	12	14	25.0	29.0	31.0
7	7	M	10	12	14	22.5	23.0	26.5
8	8	M	10	12	14	22.5	24.0	27.5
9	9	M	10	12	14	27.5	26.5	27.0
10	10	M	10	12	14	23.5	22.5	26.0

Figure 1

MIXED, however, requires that measurements at different ages be collapsed into one variable, so that each subject has three cases. The Data Restructure Wizard in SPSS simplifies the tedious data conversion process. We choose “Data- >Restructure” from the pull-down menu and select the option, “Restructure selected variables into cases.” We then click the “Next” button to reach the following dialog box:

Restructure Data Wizard - Step 2 of 7

Variables to Cases: Number of Variable Groups

You have chosen to restructure selected variables into groups of related cases in the new file.

i A group of related variables, called a variable group, represents measurements on one variable. For example, the variable may be width. If it is recorded in three separate measurements, each one representing a different point in time—w1, w2, and w3, then the data are arranged in a group of variables. If there is more than one variable in the file often it is also recorded in a variable group, for example height, recorded in h1, h2, and h3.

How many variable groups do you want to restructure?

☐ One (for example, w1, w2, and w3)

☒ More than one (for example, w1, w2, w3 and h1, h2, h3, etc.)

How Many?

< Back Next > Finish Cancel Help

Figure 2

We need to convert two groups of variables (“age” and “dist”) into cases. We therefore enter “2” and click “Next.” This brings us to the “Select Variables” dialog box on the next page:

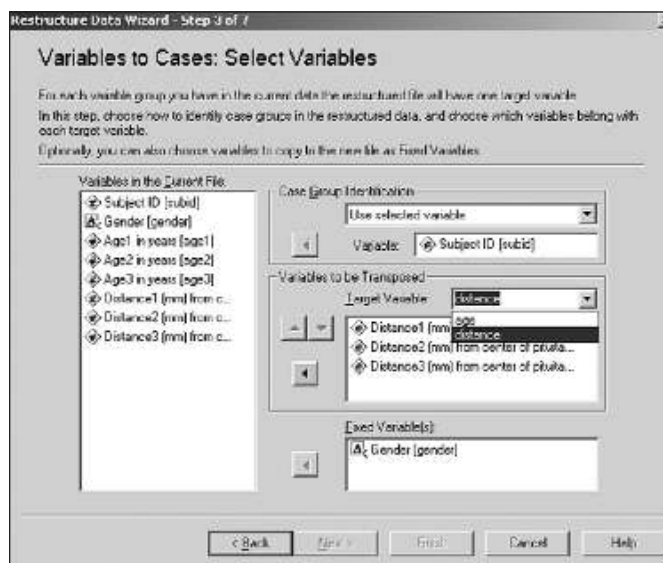


Figure 3

In the “Select Variables” dialog box, we first specify “Subject ID [subid]” as the case group identification. We then enter the names of new variables in the target variable drop-down list. For the target variable “age,” we drag “age1,” “age2” and “age3” to the list box in the “Variables to be Transposed” group. We similarly associate variables “dist1,” “dist2” and “dist3” with the target variable “distance.” We then drag variables that do not vary within a subject to the “Fixed Variable(s)” box. Clicking “Next” brings us to the “Create Index Variables” dialog box. We accept the default of one index variable, then click “Next” to arrive at the final dialog box:

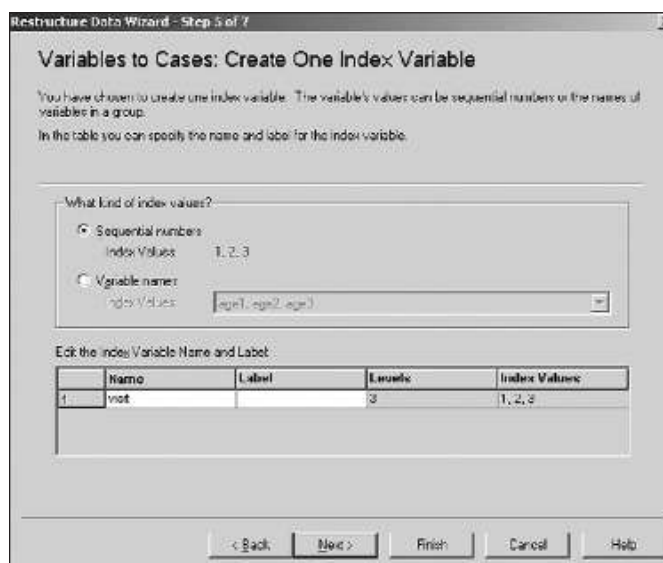
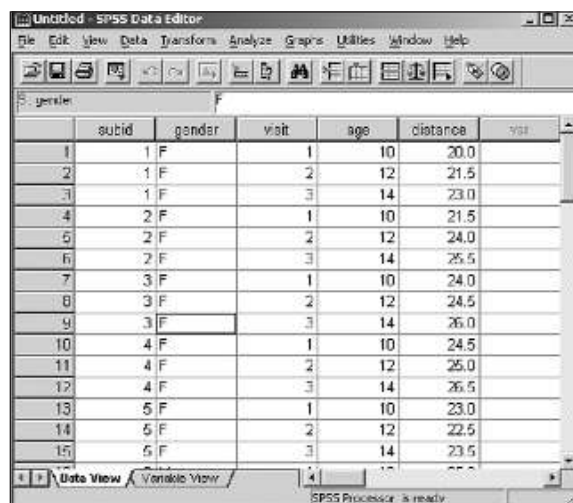


Figure 4

In the “Create One Index Variable” dialog box, we enter “visit” as the name of the indexing variable and click “Finish.” The data file is transformed into a format that MIXED can analyze. We now have three cases for each subject, as shown in Figure 5 on the next page:



	subid	gender	visit	age	distance	visit
1	1	F	1	10	20.0	
2	1	F	2	12	21.5	
3	1	F	3	14	23.0	
4	2	F	1	10	21.5	
5	2	F	2	12	24.0	
6	2	F	3	14	25.5	
7	3	F	1	10	24.0	
8	3	F	2	12	24.5	
9	3	F	3	14	26.0	
10	4	F	1	10	24.5	
11	4	F	2	12	25.0	
12	4	F	3	14	26.5	
13	5	F	1	10	23.0	
14	5	F	2	12	22.5	
15	5	F	3	14	23.5	

Figure 5

We can also perform the conversion using the following syntax:

```

VARSTOCASES
/MAKE age FROM age1 age2 age3
/MAKE distance FROM dist1 dist2 dist3
/INDEX = visit(3)
/KEEP = subid gender.

```

The syntax is easy to interpret — it collapses the three age variables into “age” and the three response variables into “distance.” At the same time, a new variable, “visit,” is created to index the three new cases within each subject. The last subcommand means that all variables that are constant within a subject should be kept.

Fitting fixed-effects models with iid residual errors

A fitted model has the form $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$, where $\mathbf{y}$ is a vector of responses, $\mathbf{X}$ is the fixed-effects design matrix, $\boldsymbol{\beta}$ is a vector of fixed-effects parameters and $\boldsymbol{\epsilon}$ is a vector of residual errors. In this model, we assume that $\boldsymbol{\epsilon}$ is distributed as $\mathbf{N}[\mathbf{0}, \mathbf{R}]$, where $\mathbf{R}$ is an unknown covariance matrix. A common belief is that $\mathbf{R} = \sigma^2 \mathbf{I}$. We can use GLM or MIXED to fit a model with this assumption. Using a subset of the growth study dataset, we illustrate how to use MIXED to fit a fixed-effects model. The following syntax (Example 1) fits a fixed-effects model that investigates the effect of the variables “gender” and “age” on “distance,” which is a measure of the growth rate.

Example 1 — Fixed-effects model using MIXED**Syntax:**

```
MIXED DISTANCE BY GENDER WITH AGE
/FIXED = GENDER AGE | SSTYPE(3)
/PRINT = SOLUTION TESTCOV.
```

Output:**Type III Tests of Fixed Effects^a**

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	27	38.356	.000
GENDER	1	27	7.621	.010
AGE	1	27	11.040	.003

Figure 6

^a. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	17.050	2.620	27	6.507	.000	11.673	22.427
[GENDER=F]	-1.933	.700	27	-2.761	.010	-3.370	-.496
[GENDER=M]	.000 ^b	.000	.	.	.	.	.
AGE	.713	.214	27	3.323	.003	.273	1.152

Figure 7

^a. This parameter is set to zero because it is redundant.

^b. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Residual	3.679	1.001	3.674	.000	2.158	6.271

Figure 8

^a. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

The syntax in Example 1 produces a “Type III Tests of Fixed Effects” table (Figure 6). Both “gender” and “age” are significant at the .05 level. This means that “gender” and “age” are potentially important predictors of the dependent variable. More detailed information on fixed-effects parameters may be obtained by using the subcommand */PRINT SOLUTION*. The “Estimates of Fixed Effects” table (Figure 7) gives estimates of individual parameters, as well as their standard errors and confidence intervals. We can see that the mean distance for males is larger than that for females. Distance, moreover, increases with age. MIXED also produces an estimate of the residual error variance and its standard error. The */PRINT TESTCOV* option gives us the Wald statistic and the confidence interval for the residual error variance estimate.

Example 1 is simple — users familiar with the GLM procedure can fit the same model using GLM.

Example 2 — Fixed-effects model using GLM

Syntax:

```
GLM DISTANCE BY GENDER WITH AGE
/METHOD = SSTYPE(3)
/PRINT = PARAMETER
/DESIGN = GENDER AGE.
```

Output:

Tests of Between-Subjects Effects

Dependent Variable: Distance (mm) from center of pituitary to pteryo-maxillary fissure

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	68.649 ^a	2	34.323	9.331	.001
Intercept	141.095	1	141.095	38.356	.000
GENDER	28.033	1	28.033	7.621	.010
AGE	40.613	1	40.613	11.040	.003
Error	99.321	27	3.679		
Total	18372.000	30			
Corrected Total	167.967	29			

Figure 9

^a. R Squared = .409 (Adjusted R Squared = .365)

Parameter Estimates

Dependent Variable: Distance (mm) from center of pituitary to pteryo-maxillary fissure

Parameter	B	Std. Error	t	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Intercept	17.050	2.620	6.507	.000	11.673	22.427
[GENDER=F]	-1.933	.700	-2.761	.010	-3.370	-.496
[GENDER=M]	0 ^a					
AGE	.713	.214	3.323	.003	.273	1.152

Figure 10

^a. This parameter is set to zero because it is redundant.

We see in Figure 9 that GLM and MIXED produced the same Type III tests and parameter estimates. Note, however, that in the “Parameter Estimates” table (Figure 10), there is no column for the sum of squares. This is because, for some complex models, the test statistics in MIXED may not be expressed as a ratio of two sums of squares. They are thus omitted from the ANOVA table.

Fitting fixed-effects models with non-iid residual errors

The assumption $\mathbf{R} = \sigma^2 \mathbf{I}$ may be violated in some situations. This often happens when repeated measurements are made on each subject. In the growth study dataset, for example, the response variable of each subject is measured at various ages. We may suspect that error terms within a subject are correlated. We may assume that $\mathbf{R}$ is a block diagonal matrix, where each block is a first-order, autoregressive (AR1) covariance matrix.

Example 3 — Fixed-effects model with correlated residual errors**Syntax:**

```

MIXED DISTANCE BY GENDER WITH AGE
/FIXED GENDER AGE
/REPEATED VISIT | SUBJECT(SUBID) COVTYPE(AR1)
/PRINT SOLUTION TESTCOV R.

```

Output:**Type III Tests of Fixed Effects^a**

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	25.723	75.036	.000
GENDER	1	8.701	3.702	.088
AGE	1	23.687	22.772	.000

Figure 11

a. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	17.243	1.947	26.760	8.857	.000	13.246	21.239
[GENDER=F]	-2.072	1.077	8.701	-1.924	.088	-4.522	.377
[GENDER=M]	.000 ^a	.000	.	.	.	.	.
AGE	.712	.149	23.687	4.772	.000	.404	1.021

Figure 12

a. This parameter is set to zero because it is redundant.

b. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Repeated Measures: AR1 diagonal	3.806	1.467	2.597	.009	1.791	8.101
AR1 rho	.729	.120	6.072	.000	.401	.852

Figure 13

a. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Residual Covariance (R) Matrix^a

	[VISIT = 1]	[VISIT = 2]	[VISIT = 3]
[VISIT = 1]	3.809	2.778	2.026
[VISIT = 2]	2.778	3.809	2.778
[VISIT = 3]	2.026	2.778	3.809

Figure 14

First-Order Autoregressive

^a. Dependent Variable: Distance (mm) from center of pituitary to pterygo-maxillary fissure.

Example 3 uses the */REPEATED* subcommand to specify a more general covariance structure for the residual errors. Since there are three observations per subject, we assume that the set of three residual errors for each subject is a sample from a three-dimensional, normal distribution, with a first-order, autoregressive (AR1) covariance matrix. Residual errors within each subject are therefore correlated, but are independent across subjects. The MIXED procedure, by default, uses the restricted maximum likelihood (REML) method to estimate the covariance matrix. An option is to request maximum likelihood (ML) estimates.

The syntax in Example 3 also produces the “Residual Covariance (R) Matrix” (Figure 14), which shows the estimated covariance matrix of the residual error for one subject. We see from the “Estimates of Covariance Parameters” table (Figure 13) that the correlation parameter has a relatively large value (.729) and that the p-value of the Wald test is less than .05. The autoregressive structure may thus fit the data better than the model in Example 1.

We also see that, for the tests of fixed effects, the denominator degrees of freedoms are not integers. This is because these statistics do not have exact F distributions. The denominator degrees of freedoms are obtained by a Satterthwaite approximation. We see in the new model that gender is not significant at the .05 level. This demonstrates that ignoring the possible correlations in your data may lead to incorrect conclusions. MIXED is therefore usually a better alternative to GLM and VARCOMP when you suspect correlation in the data.

Fitting simple mixed-effects models (balanced design)

MIXED, as its name implies, handles complicated models that involve fixed and random effects. Levels of an effect are, in some situations, only a sample of all possible levels. If we want to study the efficiency of workers in different environments, for example, we don't need to include all workers in the study — a sample of workers is usually enough. The worker effect should be considered random, due to the sampling process. A mixed-effects model has, in general, the form $y = X\beta + Z\gamma + \epsilon$ where the extra term $Z\gamma$ models the random effects. Z is the design matrix of random effects and γ is a vector of random-effects parameters. We can use GLM and MIXED to fit mixed-effects models. MIXED, however, fits a much wider class of models. To understand the functionality of MIXED, we first look at several simpler models that can be created in MIXED and GLM. We also look at the similarity between MIXED and VARCOMP in these models.

In examples 4 through 6, we use a semiconductor dataset that appeared in Pinheiro and Bates (2000) to illustrate the similarity between GLM, MIXED and VARCOMP. The dependent variable in this dataset is “current” and the predictor is “voltage.” The data are collected from a sample of ten silicon wafers. There are eight sites on each wafer and five measurements are taken at each site. We have, therefore, a total of 400 observations and a balanced design.

Example 4 — Simple mixed-effects model with balanced design using MIXED

Syntax:

```
MIXED CURRENT BY WAFER WITH VOLTAGE
/FIXED VOLTAGE | SSTYPE(3)
/RANDOM WAFER
/PRINT SOLUTION TESTCOV.
```

Output:

Type III Tests of Fixed Effects^a

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	16.531	3774.499	.000
VOLTAGE	1	389.000	67958.177	.000

Figure 15

^a. Dependent Variable: CURRENT.

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	-7.06267	.1152858	16.531	-61.437	.000	-7.3266275	-6.8391076
VOLTAGE	9.5486505	.0370123	389.000	260.588	.000	9.5758903	9.7214287

Figure 16

^a. Dependent Variable: CURRENT

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Residual	.175	.013	13.946	.000	.152	.202
WAFER ID diagonal	.093	.046	2.026	.043	.036	.246

Figure 17

^a. Dependent Variable: CURRENT.

Example 5 — Simple mixed-effects model with balanced design using GLM**Syntax:**

```

GLM CURRENT BY WAFER WITH VOLTAGE
/RANDOM = WAFER
/METHOD = SSTYPE(3)
/PRINT = PARAMETER
/DESIGN = WAFER VOLTAGE.

```

Output:

Tests of Between-Subjects Effects

Dependent Variable: CURRENT

Source		Type III Sum of Squares	df	Mean Square	F	Sig.
Intercept	Hypothesis	2229.645	1	2229.645	3774.499	.000
	Error	9.782	16.559	.591 ^a		
WAFER	Hypothesis	35.223	9	3.914	22.319	.000
	Error	68.211	389	.175 ^b		
VOLTAGE	Hypothesis	11916.369	1	11916.369	67958.177	.000
	Error	68.211	389	.175 ^b		

Figure 18

- a. .111 MS(WAFER) + .389 MS(Error)
- b. MS(Error)

Expected Mean Squares^{a,b}

Source	Variance Component		
	Var(WAFER)	Var(Error)	Quadratic Term
Intercept	4.444	1.000	Intercept
WAFER	40.000	1.000	
VOLTAGE	.000	1.000	VOLTAGE
Error	.000	1.000	

Figure 19

- a. For each source, the expected mean square equals the sum of the coefficients in the cells times the variance components, plus a quadratic term involving effects in the Quadratic Term cell.
- b. Expected Mean Squares are based on the Type III Sums of Squares.

Example 6 — Variance components model with balanced design**Syntax:**

```

VARCOMP CURRENT BY WAFER WITH VOLTAGE
/RANDOM = WAFER
/METHOD = REML.

```

Output:**Variance Estimates**

Component	Estimate
Var(WAFER)	.093
Var(Error)	.175

Figure 20

Dependent Variable: CURRENT

Method: Restricted Maximum Likelihood Estimation

In Example 4, “voltage” is entered as a fixed effect and “wafer” is entered as a random effect. This example tries to model the relationship between “current” and “voltage” using a straight line, but the intercept of the regression line will vary from wafer to wafer according to a normal distribution. In the Type III tests for “voltage,” we see a significant relationship between “current” and “voltage.” If we delve deeper into the parameter estimates table, the regression coefficient of “voltage” is 9.65. This indicates a positive relationship between “current” and “voltage.” In the “Estimates of Covariance Parameters” table (Figure 17), we have estimates for the residual error variance and the variance due to the sampling of wafers.

We repeat the same model in Example 5 using GLM. Note that MIXED produces Type III tests for fixed effects only, but GLM includes fixed and random effects. GLM treats all effects as fixed during computation and constructs F statistics by taking the ratio of the appropriate sums of squares. Mean squares of random effects in GLM are estimates of functions of the variance parameters of random and residual effects. These functions can be recovered from “Expected Mean Squares” (Figure 19). In MIXED, the outputs are much simpler because the variance parameters are estimated directly using maximum likelihood (ML) or restricted maximum likelihood (REML). As a result, there is no random-effect sum of squares.

When we have a balanced design, as in examples 4 through 6, the tests of fixed effects are the same for GLM and MIXED. We can also recover the variance parameter estimates of MIXED by using the sum of squares in GLM. In MIXED, for example, the estimate of the residual variance is 0.175, which is the same as the MS(Error) in GLM. The variance estimate of random effect “wafer” is 0.093, which can be recovered in GLM using the “Expected Mean Squares” table (Figure 19) in Example 5:

$$\text{Var(WAFER)} = [\text{MS(WAFER)} - \text{MS(Error)}] / 40 = 0.093$$

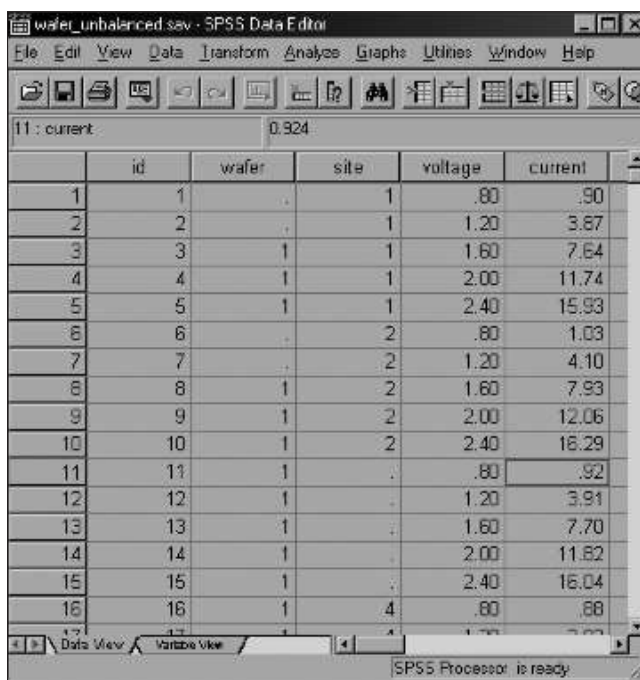
This is equal to MIXED’s estimate. One drawback of GLM, however, is that you cannot compute the standard error of the variance estimates.

VARCOMP is, in fact, a subset of MIXED. These two procedures, therefore, always provide the same variance estimates, as seen in examples 4 and 6. VARCOMP only fits relatively simple models. It can only handle random effects that are iid normal. No statistics on fixed effects are produced. If, therefore, your primary objective is to make inferences about fixed effects and your data are correlated, MIXED is a better choice.

An important note: due to the different estimation methods that are used, GLM and MIXED do not produce the same results. The next section gives an example of such differences.

Fitting simple mixed-effects models (unbalanced design)

One situation about which MIXED and GLM disagree is an unbalanced design. To illustrate this, we removed some cases in the semiconductor dataset, so that the design is no longer balanced.



	id	wafer	site	voltage	current
1	1	.	1	.80	.90
2	2	.	1	1.20	3.87
3	3	1	1	1.60	7.64
4	4	1	1	2.00	11.74
5	5	1	1	2.40	15.93
6	6	.	2	.80	1.03
7	7	.	2	1.20	4.10
8	8	1	2	1.60	7.93
9	9	1	2	2.00	12.06
10	10	1	2	2.40	16.29
11	11	1	.	.80	.92
12	12	1	.	1.20	3.91
13	13	1	.	1.60	7.70
14	14	1	.	2.00	11.62
15	15	1	.	2.40	16.04
16	16	1	4	.80	.88

Figure 21

We then rerun examples 4 through 6 with this unbalanced dataset. The output is shown in examples 4a through 6a. We want to see whether the three methods — GLM, MIXED and VARCOMP — still agree with each other.

Example 4a — Mixed-effects model with unbalanced design using MIXED**Syntax:**

```

MIXED CURRENT BY WAFER WITH VOLTAGE
/FIXED VOLTAGE | SSTYPE(3)
/RANDOM WAFER
/PRINT SOLUTION TESTCOV.

```

Output:**Type III Tests of Fixed Effects^a**

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	16.467	3709.960	.000
VOLTAGE	1	385.034	67481.118	.000

Figure 22

^a. Dependent Variable: CURRENT.**Estimates of Fixed Effects^a**

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	-7.096	.117	16.467	-60.909	.000	-7.345	-6.852
VOLTAGE	9.555	.037	385.034	259.771	.000	9.553	9.730

Figure 23

^a. Dependent Variable: CURRENT.**Estimates of Covariance Parameters^a**

Parameter	Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Residual	.1744505	.013	13.874	.000	.151	.201
WAFER ID diagonal	.0957247	.047	2.027	.043	.036	.252

Figure 24

^a. Dependent Variable: CURRENT.

Example 5a — Mixed-effects model with unbalanced design using GLM

Syntax:

```
GLM CURRENT BY WAFER WITH VOLTAGE
/RANDOM = WAFER
/METHOD = SSTYPE(3)
/PRINT = PARAMETER /DESIGN = WAFER VOLTAGE.
```

Output:

Tests of Between-Subjects Effects

Dependent Variable: CURRENT

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Intercept	2193.281	1	2193.281	3724.816	.000
WAFER	35.495	9	3.944	22.607	.000
VOLTAGE	11772.307	1	11772.307	67482.629	.000
Error	9.746	15,551	.5886293 ^a		
Error	67.163	385	.1744494 ^b		
Error	67.163	385	.1744494 ^b		

Figure 25

^a. 110 MS(WAFER) + .890 MS(Error)

^b. MS(Error)

Expected Mean Squares^{a,b}

Source	Variance Component		
	Var(WAFER)	Var(Error)	Quadratic Term
Intercept	4.352	1.000	Intercept
WAFER	39.591	1.000	
VOLTAGE	.000	1.000	VOLTAGE
Error	.000	1.000	

Figure 26

- ^a. For each source, the expected mean square equals the sum of the coefficients in the cells times the variance components, plus a quadratic term involving effects in the Quadratic Term cell.
- ^b. Expected Mean Squares are based on the Type III Sums of Squares.

Example 6a — Variance components model with unbalanced design**Syntax:**

```
VARCOMP CURRENT BY WAFER WITH VOLTAGE
/RANDOM = WAFER
/METHOD = REML.
```

Output:

Variance Estimates	
Component	Estimate
Var(WAFER)	.0957247
Var(Error)	.1744505

Figure 27

Dependent Variable: CURRENT
Method: Restricted Maximum Likelihood Estimation

Since the data have changed, we expect examples 4a through 6a to differ from examples 4 through 6. We will focus instead on whether examples 4a, 5a and 6a agree with each other.

In Example 4a, the F statistic for the “voltage” effect is 67481.118, but Example 5a gives an F statistic value of 67482.629. Apart from the test of fixed effects, we also see a difference in covariance parameter estimates.

Examples 4a and 6a, however, show that VARCOMP and MIXED can produce the same variance estimates, even in an unbalanced design. This is because MIXED and VARCOMP use maximum likelihood or restricted maximum likelihood methods in estimation, while GLM estimates are based on the method-of-moments approach.

MIXED is generally preferred because it is asymptotically efficient (minimum variance), whether or not the data are balanced. GLM, however, only achieves its optimum behavior when the data are balanced.

Fitting mixed-effects models with subjects

In the semiconductor dataset, “current” is a dependent variable measured on a batch of wafers. These wafers are therefore considered subjects in a study. An effect of interest (such as “site”) may often vary with subjects (“wafer”). One scenario is that the (population) means of “current” at separate sites are different. When we look at the current measured at these sites on individual wafers, however, they hover below or above the population mean according to some normal distribution. It is therefore common to enter an “effect by subject” interaction term in a GLM or MIXED model to account for the subject variations.

In the dataset there are eight sites and ten wafers. The site*wafer effect, therefore, has 80 parameters, which can be denoted by γ_{ij} , $i=1\ldots 10$ and $j=1\ldots 8$. A common assumption is that γ_{ij} ’s are assumed to be iid normal with zero mean and an unknown variance. The mean is zero because γ_{ij} ’s are used to model only the population variation. The mean of the population is modeled by entering “site” as a fixed effect in GLM and MIXED. The results of this model for MIXED and GLM are shown in examples 7 and 8.

Example 7 — Fitting random effect*subject interaction using MIXED**Syntax:**

```
MIXED CURRENT BY WAFER SITE WITH VOLTAGE
/FIXED SITE VOLTAGE |SSTYPE(3)
/RANDOM SITE*WAFER | COVTYPE(ID).
```

Output:

Type III Tests of Fixed Effects^a

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	264.928	10467.974	.000
SITE	7	68.691	1.140	.349
VOLTAGE	1	319.000	76639.444	.000

Figure 28

^a Dependent Variable: CURRENT.

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Residual	.155	.012	12.629	.000	.133	.182
SITE * WAFER ID diagonal	.104	.023	4.586	.000	.068	.159

Figure 29

^a Dependent Variable: CURRENT.**Example 8 — Fitting random effect*subject interaction using GLM****Syntax:**

```
GLM CURRENT BY WAFER SITE WITH VOLTAGE
/RANDOM = WAFER
/METHOD = SSTYPE(3)
/DESIGN = SITE SITE*WAFER VOLTAGE.
```

Output:

Tests of Between-Subjects Effects

Dependent Variable: CURRENT

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Intercept	2229.645	1	2229.645	10467.974	.000
Error	70.248	329.798	.213 ^a		
SITE	5.371	7	.767	1.140	.348
Error	48.462	72	.673 ^a		
WAFER *	48.462	72	.673	4.328	.000
SITE	49.600	319	.155 ^a		
VOLTAGE	11918.369	1	11918.369	76639.444	.000
Error	49.600	319	.155 ^a		

Figure 30

Expected Mean Squares^{a,b}

Source	Variance Component		
	Var(WAFER * SITE)	Var(Error)	Quadratic Term
Intercept	.556	1.000	Intercept, SITE
SITE	5.000	1.000	SITE
WAFER * SITE	5.000	1.000	
VOLTAGE	.000	1.000	VOLTAGE
Error	.000	1.000	

Figure 31

- a. For each source, the expected mean square equals the sum of the coefficients in the cells times the variance components, plus a quadratic term involving effects in the Quadratic Term cell.
- b. Expected Mean Squares are based on the Type III Sums of Squares.

Since the design is balanced, the results of GLM and MIXED in examples 7 and 8 match. This is similar to examples 4 and 5. We see from the results of Type III tests that “voltage” is still an important predictor of “current,” while “site” is not. The mean currents at different sites are thus not significantly different from each other, so we can use a simpler model without the fixed effect “site.” We should still, however, consider a random-effects model, because ignoring the subject variations may lead to incorrect standard error estimates of fixed effects or false significant tests.

Up to this point, we examined primarily the similarities between GLM and MIXED. MIXED, in fact, has a much more flexible way of modeling random effects. Using the SUBJECT and COVTYPE options, Example 9 presents an equivalent form of Example 7.

Example 9 — Fitting random effect*subject interaction using *SUBJECT* specification

Syntax:

```
MIXED CURRENT BY SITE WITH VOLTAGE
/FIXED SITE VOLTAGE |SSTYPE(3)
/RANDOM SITE | SUBJECT(WAFER) COVTYPE(ID).
```

The *SUBJECT* option tells MIXED that each subject will have its own set of random parameters for the random effect “site.” The *COVTYPE* option will specify the form of the variance covariance matrix of the random parameters within one subject. The syntax attempts to specify the distributional assumption in a multivariate form, which can be written as:

$$\begin{pmatrix} \gamma_{1,1} \\ \vdots \\ \gamma_{1,8} \end{pmatrix}, \begin{pmatrix} \gamma_{2,1} \\ \vdots \\ \gamma_{2,8} \end{pmatrix}, \dots, \begin{pmatrix} \gamma_{10,1} \\ \vdots \\ \gamma_{10,8} \end{pmatrix} \text{ are iid } N \left[\begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_{\epsilon}^2 & 0 & \dots & 0 \\ 0 & \sigma_{\epsilon}^2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_{\epsilon}^2 \end{pmatrix} \right]$$

Figure 32

This assumption is equivalent to that in Example 7 under normality.

One advantage of the multivariate form is that you can easily specify other covariance structures by using the *COVTYPE* option. The flexibility in specifying covariance structures helps us to fit a model that better describes the data. If, for example, we believe that the variances of different sites are different, we can specify a diagonal matrix as covariance type and the assumption becomes:

$$\begin{pmatrix} \mathbf{Y}_{1,1} \\ \vdots \\ \mathbf{Y}_{1,8} \end{pmatrix} \begin{pmatrix} \mathbf{Y}_{2,1} \\ \vdots \\ \mathbf{Y}_{2,8} \end{pmatrix} \dots \begin{pmatrix} \mathbf{Y}_{10,1} \\ \vdots \\ \mathbf{Y}_{10,8} \end{pmatrix} \text{ are iid } \mathbf{N} \left[\begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_{\epsilon}^2 & 0 & \dots & 0 \\ 0 & \sigma_{\epsilon}^2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_{\epsilon}^2 \end{pmatrix} \right] \quad \text{Figure 33}$$

The result of fitting the same model using this assumption is given in Example 10.

Example 10 — Using *COVTYPE* in a random-effects model

Syntax:

```
MIXED CURRENT BY SITE WITH VOLTAGE
/FIXED SITE VOLTAGE |SSTYPE(3)
/RANDOM SITE | SUBJECT(WAFER) COVTYPE(VC)
/PRINT G TESTCOV.
```

Output:

Type III Tests of Fixed Effects^a

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	252.867	10467.974	.000
SITE	7	16.071	1.267	.326
VOLTAGE	1	319.000	76639.444	.000

Figure 34

a. Dependent Variable: CURRENT.

Estimates of Covariance Parameters^a

Parameter		Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Residual		.155	.012	12.629	.000	.133	.182
SITE [subject = WAFER]	VC diagonal 1	.136	.079	1.726	.084	.044	.424
	VC diagonal 2	.096	.060	1.599	.110	.028	.326
	VC diagonal 3	.183	.101	1.812	.070	.062	.539
	VC diagonal 4	.119	.071	1.681	.093	.037	.382
	VC diagonal 5	.071	.048	1.475	.140	.019	.269
	VC diagonal 6	.073	.049	1.484	.138	.019	.272
	VC diagonal 7	.030	.029	1.046	.296	.005	.198
	VC diagonal 8	.120	.071	1.685	.092	.038	.385

Figure 35

a. Dependent Variable: CURRENT.

SITE [subject = WAFER]

	[SITE=1] WAFER	[SITE=2] WAFER	[SITE=3] WAFER	[SITE=4] WAFER	[SITE=5] WAFER	[SITE=6] WAFER	[SITE=7] WAFER	[SITE=8] WAFER
[SITE=1] WAFER	.136	.000	.000	.000	.000	.000	.000	.000
[SITE=2] WAFER	.000	.096	.000	.000	.000	.000	.000	.000
[SITE=3] WAFER	.000	.000	.183	.000	.000	.000	.000	.000
[SITE=4] WAFER	.000	.000	.000	.119	.000	.000	.000	.000
[SITE=5] WAFER	.000	.000	.000	.000	.071	.000	.000	.000
[SITE=6] WAFER	.000	.000	.000	.000	.000	.073	.000	.000
[SITE=7] WAFER	.000	.000	.000	.000	.000	.000	.030	.000
[SITE=8] WAFER	.000	.000	.000	.000	.000	.000	.000	.120

Variance Components

a. Dependent Variable: CURRENT.

Figure 36

In Example 10, we request one extra table, the estimated covariance matrix of the random effect “site.” It is an eight-by-eight diagonal matrix in this case. Note that changing the covariance structure of a random effect also changes the estimates and tests of a fixed effect. We want, in practice, an objective method to select a suitable covariance structure for our random effects. In the section “Covariance Structure Selection,” we revisit examples 9 and 10 to show how to select covariance structure for random effects.

Multilevel analysis

The use of the *SUBJECT* and *COVTYPE* options in */RANDOM* and */REPEATED* brings many options for modeling the covariance structures of random effects and residual errors. It is particularly useful when modeling data obtained from a hierarchy. Example 11 illustrates the simultaneous use of these options in a multilevel model. We selected data from six schools from the Junior School Project of Mortimore et al. (1988). We investigate below how the socioeconomic status (SES) of a student affects his or her math scores over a three-year period.

Example 11 — Multilevel mixed-effects model

Syntax:

```
MIXED MATHTEST BY SCHOOL CLASS STUDENT GENDER SES SCHLYEAR
/FIXED GENDER SES SCHLYEAR SCHOOL
/RANDOM SES |SUBJECT(SCHOOL*CLASS) COVTYPE(ID)
/RANDOM SES |SUBJECT(SCHOOL*CLASS*STUDENT) COVTYPE(ID)
/REPEATED SCHLYEAR | SUBJECT(SCHOOL*CLASS*STUDENT) COVTYPE(AR1)
/PRINT SOLUTION TESTCOV.
```

Output:

Type III Tests of Fixed Effects^a

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	15.100	1076.489	.000
GENDER	1	96.609	.979	.325
SES	2	16.513	3.888	.041
SCHLYEAR	2	201.195	55.376	.000
SCHOOL	5	12.969	.872	.526

a. Dependent Variable: Math test.

Figure 37

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	29.097	2.184	19.564	13.324	.000	24.535	33.659
[GENDER=0]	-1.026	1.037	96.609	-.989	.325	-3.084	1.032
[GENDER=1]	.000 ^b	.000					
[SES=1.00]	5.803	2.331	20.899	2.490	.021	.955	10.652
[SES=2.00]	.304	1.782	13.705	.170	.867	-3.527	4.134
[SES=3.00]	.000 ^b	.000					
[SCHLYEAR=0]	-4.377	.457	116.837	-9.575	.000	-5.282	-3.471
[SCHLYEAR=1]	-4.126	.468	219.911	-8.825	.000	-5.047	-3.204
[SCHLYEAR=2]	.000 ^b	.000					
[SCHOOL=1]	-2.751	2.405	12.727	-1.144	.274	-7.958	2.456
[SCHOOL=2]	-.784	2.865	18.151	-.274	.787	-6.801	5.232
[SCHOOL=3]	2.289	2.645	14.332	.858	.405	-3.392	7.929
[SCHOOL=4]	-1.911	2.811	9.275	-.680	.513	-8.241	4.420
[SCHOOL=5]	-.686	2.545	15.323	-.270	.791	-6.100	4.728
[SCHOOL=6]	.000 ^b	.000					

Figure 38

^a This parameter is set to zero because it is redundant.

^b Dependent Variable: Math test.

Estimates of Covariance Parameters^a

Parameter		Estimate	Std. Error	Wald Z	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Repeated Measures	AR1 diagonal	12.686	1.667	7.609	.000	9.805	15.413
	AR1 rho	-.027	.142	-.190	.850	-.296	.246
SES [subject = SCHOOL * CLASS]	ID diagonal	6.450	4.991	1.292	.196	1.415	29.391
SES [subject = SCHOOL * CLASS * STUDENT]	ID diagonal	30.409	4.782	6.358	.000	22.342	41.387

Figure 39

^a Dependent Variable: Math test.

In Example 11, the goal is to discover whether socioeconomic status (“ses”) is an important predictor for mathematics achievement (“mathtest”). To do so, we use the factor “ses” as a fixed effect. We also want to adjust for the possible sampling variation due to different classes and students. “Ses” is therefore also used twice as a random effect. The first random effect tries to adjust for the variation of the “ses” effect owing to class variation. In order to identify all classes in the dataset, school*class is specified in the *SUBJECT* option. The second random effect also tries to adjust for the variation of the “ses” effect owing to student variation. The subject specification is thus school*class*student. All of the students are followed for three years; the school year (“schlyear”) is therefore used as a fixed effect to adjust for possible trends in this period. The */REPEATED* subcommand is also used to model the possible correlation of the residual errors with each student.

We have a relatively small dataset, since there are only six schools, so we can only use it as a fixed effect while adjusting for possible differences between schools. In this example, there is only one random effect in each level. More than one effect can, in general, be specified, but MIXED will assume that these effects are independent. Check to see whether such an assumption is reasonable for your applications. In SPSS 11.5, users are able to specify correlated random effects, which provide additional flexibility in modeling data.

In the Type III tests of fixed effects, in Example 11, we see that socioeconomic status does have an impact on student performance. The parameter estimates of “ses” for students with “ses=1” (fathers have managerial or professional occupations) indicates that they perform better than other groups. The effect “schlyear” is also significant in the model and the students’ performances increase with “schlyear.”

From “Estimates of Covariance Parameters” (Figure 39), we notice that the estimate of the “AR1 rho” parameter is not significant, which means that a simple, scaled-identity structure may be used. For the variation of “ses” due to school*class, the estimate is very small compared to other sources of variance and the Wald test indicates that it is not significant. We can therefore consider removing the random effect from the model.

We see from this example that the major advantages of MIXED are that it is able to look at different aspects of a dataset simultaneously and that all of the statistics are already adjusted for all effects in the model. Without MIXED, we must use different tools to study different aspects of the models. An example of this is using GLM to study the fixed effect and using VARCOMP to study the covariance structure. This is not only time consuming, but the assumptions behind the statistics are usually violated.

Custom hypothesis tests

Apart from predefined statistics, MIXED allows users to construct custom hypotheses on fixed- and random-effects parameters through the use of the */TEST* subcommand. To illustrate, we use a dataset in Pinheiro and Bates (2000). The data consist of a CT scan on a sample of ten dogs. The dogs’ left and right lymph nodes were scanned and the intensity of each scan was recorded in the variable pixel. The following mixed-model syntax tests whether there is a difference between the left and right lymph nodes.

Example 12 — Custom hypothesis testing in mixed-effects model

Syntax:

```
MIXED PIXEL BY SIDE
/FIXED SIDE
/RANDOM SIDE | SUBJECT(DOG) COVTYPE(UN)
/TEST(0) 'Side (fixed)' SIDE 1 -1
/TEST(0) 'Side (random)' SIDE 1 -1 | SIDE 1 -1
/PRINT LMATRIX.
```

Output:

Contrast Coefficients ^a		
		L1
Fixed Effects	Intercept	0
	[SIDE=L]	1
	[SIDE=R]	-1
Random Effects	[SIDE=L] DOG	0
	[SIDE=R] DOG	0

Figure 40

Contrast Estimates^{a,b}

Contrast	Estimate	Std. Error	df	Test Value	t	Sig.	95% Confidence Interval	
							Lower Bound	Upper Bound
L1	8.502	7.337	7.723	.000	1.159	.281	-8.523	25.526

^a. Side (fixed)

^b. Dependent Variable: PIXEL

Figure 41

Contrast Estimates^{a,b}

Contrast	Estimate	Std. Error	df	Test Value	t	Sig.	95% Confidence Interval	
							Lower Bound	Upper Bound
L1	8.502	3.205	82.839	.000	2.653	.010	2.127	14.877

^a. Side (random)

^b. Dependent Variable: PIXEL

Figure 42

Contrast Estimates^{a,b}

Contrast	Estimate	Std. Error	df	Test Value	t	Sig.	95% Confidence Interval	
							Lower Bound	Upper Bound
L1	8.5018356	3.2050572	86.681	0	2.653	.009	2.1311069	14.8725643

^a. Side (random)

^b. Dependent Variable: PIXEL

Figure 43

The output of the two *TEST* subcommands is shown above. The first test looks at differences in the left and right sides in the general population (broad inference space). We should use the second test to test the differences between the left and right sides for the sample of dogs used in this particular study (narrow inference space). In the second test, the average differences of the random effects over the ten dogs are added to the statistics. MIXED automatically calculates the average over subjects. Note that the contrast coefficients for random effects are scaled by one/(number of subjects). Though the average difference for the random effect is zero, it affects the standard error of the statistic. We see that statistics of the two tests are the same, but the second has a smaller standard error. This means that if we make an inference on a larger population, there will be more uncertainty. This is reflected in the larger standard error of the test. The hypothesis in this example is not significant in the general population, but it is significant for the narrow inference. A larger sample size is therefore often needed to test a hypothesis about the general population.

Covariance structure selection

In examples 3 and 11, we see the use of Wald statistics in covariance structure selection. Another approach to testing hypotheses on covariance parameters uses likelihood ratio tests. The statistics are constructed by taking the differences of the -2 Log likelihood of two nested models. It follows a chi-squared distribution with degrees of freedom equal to the difference in the number of parameters of the models.

To illustrate the use of the likelihood ratio test, we again look at the model in examples 9 and 10. In Example 9, we use a scaled identity as the covariance matrix of the random effect "site." In Example 10, however, we use a diagonal matrix with unequal diagonal elements. Our goal is to discover which model better fits the data. We obtain the -2 Log likelihood and other criteria about the two models from the information criteria tables shown on the next page.

Information criteria for Example 9

Information Criteria ^a	
-2 Restricted Log Likelihood	523.532
Akaike's Information Criterion (AIC)	527.532
Hurvich and Tsai's Criterion (AICC)	527.563
Bozdogan's Criterion (CAIC)	537.469
Schwarz's Bayesian Criterion (BIC)	535.469

Figure 44

The information criteria are displayed in smaller-is-better forms.

^a. Dependent Variable: CURRENT.

Information criteria for Example 10

Information Criteria ^a	
-2 Restricted Log Likelihood	519.290
Akaike's Information Criterion (AIC)	537.290
Hurvich and Tsai's Criterion (AICC)	537.763
Bozdogan's Criterion (CAIC)	582.009
Schwarz's Bayesian Criterion (BIC)	573.009

Figure 45

The information criteria are displayed in smaller-is-better forms.

^a. Dependent Variable: CURRENT

The likelihood ratio test statistic for testing Example 9 (null hypothesis) versus Example 10 is $523.532 - 519.290 = 4.242$. This statistic has a chi-squared distribution and the degree of freedom is determined by the difference (seven) in the number of parameters in the two models. The p-value of this statistic is 0.752, which is not significant at level 0.05. The likelihood ratio test indicates, therefore, that we may use the simpler model in Example 9. Apart from Wald statistics and likelihood ratio tests, we can also use such information criteria as Akaike's Information Criterion (AIC) and Schwarz's Bayesian Criterion (BIC) to search for the best model.

Reference

McCulloch, C.E., and Searle, S.R. (2000). *Generalized, Linear, and Mixed Models*. John Wiley and Sons.

Mortimore, P., Sammons, P., Stoll, L., Lewis, D. and Ecob, R. (1988). *School Matters: the Junior Years*. Wells, Open Books.

Pinheiro J.C., and Bates, D.M. (2000). *Mixed-Effects Models in S and S-PLUS*. Springer.

Potthoff, R.F., and Roy, S.N. (1964). "A generalized multivariate analysis of variance model useful especially for growth curve problems." *Biometrika*, 51:313-326.

Verbeke, G., and Molenberghs, G. (2000). *Linear Mixed Models for Longitudinal Data*. Springer.

About SPSS Inc.

SPSS Inc. (Nasdaq: SPSS), headquartered in Chicago, IL, USA, is a multinational computer software company providing technology that transforms data into insight through the use of predictive analytics and other data mining techniques. The company's solutions and products enable organizations to manage the future by learning from the past, understanding the present and predicting potential problems and opportunities. For more information, visit www.spss.com.

SPSS is a registered trademark and the other SPSS products named are trademarks of SPSS Inc. All other products are trademarks of their respective owners. © Copyright 2002 SPSS Inc.