

Linear Mixed Effects Modeling using R.

By Dr. Jon Starkweather
Research and Statistical Support consultant

There are a great many ways to do linear (and non-linear) mixed effects modeling in R. The following article discusses the use of the [lme4](#) package, because; it has been developed thoroughly over time and provides reliable, easy to interpret output for mixed effect models. The motivation for this article comes from the growing recognition of the prevalence of nested effects. For those new to R, I would suggest reviewing the Research and Statistical Support (RSS) Do-it-Yourself (DIY) [Introduction to R short course](#). A script file containing all the commands used in this article can be found [here](#).

1. Mixed Effects Models

Mixed effects models refer to a variety of models which have as a key feature both fixed and random effects. The distinction between fixed and random effects is a murky one. As pointed out by [Gelman \(2005\)](#), there are several, often conflicting, definitions of fixed effects as well as definitions of random effects. Gelman offers a fairly intuitive solution in the form of renaming fixed effects and random effects and providing his own clear definitions of each. “We define effects (or coefficients) in a multilevel model as *constant* if they are identical for all groups in a population and *varying* if they are allowed to differ from group to group” (Gelman, p. 21). Other ways of thinking about fixed and random effects, which may be useful but are not always consistent with one another or those given by Gelman above, are discussed in the next paragraph.

Fixed effects are ones in which the possible values of the variable are fixed. Random effects refer to variables in which the set of potential values can change. Stated in terms of populations, fixed effects can be thought of as effects for which the population elements are fixed. Cases or individuals do not move into or out of the population. Random effects can be thought of as effects for which the population elements are changing or can change (i.e. random variable). Cases or individuals can and do move into and out of the population. Another way of thinking about the distinction between fixed and random effects is at the observation level. Fixed effects assume scores or observations are independent while random effects assume *some* type of relationship exists between *some* scores or observations. For instance, it can be said that gender is a fixed effect variable because we know all the values of that variable (male & female) and those values are independent of one another (mutually exclusive); and they (typically) do not change. A variable such as high school class has random effects because we can only sample some of the classes which exist; not to mention, students move into and out of those classes each year.

There are many types of random effects, such as repeated measures of the same individuals; where the scores at each time of measure constitute samples from the same participants among a virtually infinite (and possibly *random*) number of times of measure from those participants. Another example of a random effect can be seen in nested designs, where for example; achievement scores of students are nested within classes and those classes are nested within schools. That would be an example of a hierarchical design structure with a random effect for scores nested within classes and a second random effect for classes nested within schools. The nested data structure assumes a relationship among groups such that members of a class are thought to be similar to others in their class in such a way as to distinguish them from members of other classes and members of a school are thought to be similar to others in their school in such a way as to distinguish them from members of other schools. The example used below deals with a similar design which focuses on multiple fixed effects and a single nested random effect.

2. Linear Mixed Effects Models

Linear mixed effects models simply model the fixed and random effects as having a linear form. Similar to the General Linear Model, an outcome variable is contributed to by additive fixed and random effects (as well as an error term). Using the familiar notation, the linear mixed effect model takes the form:

$$y_{ij} = \beta_1 x_{1ij} + \beta_2 x_{2ij} \dots \beta_n x_{nij} + b_{i1} z_{1ij} + b_{i2} z_{2ij} \dots b_{in} z_{nij} + \varepsilon_{ij}$$

where y_{ij} is the value of the outcome variable for a particular ij case, β_1 through β_n are the fixed effect coefficients (like regression coefficients), x_{1ij} through x_{nij} are the fixed effect variables (predictors) for observation j in group i (usually the first is reserved for the intercept/constant; $x_{1ij} = 1$), b_{i1} through b_{in} are the random effect coefficients which are assumed to be multivariate normally distributed, z_{1ij} through z_{nij} are the random effect variables (predictors), and ε_{ij} is the error for case j in group i where each group's error is assumed to be multivariate normally distributed.

3. Example Data

The example used for this article is fictional data where the interval scaled outcome variable Extroversion (extro) is predicted by fixed effects for the interval scaled predictor Openness to new experiences (open), the interval scaled predictor Agreeableness (agree), the interval scaled predictor Social engagement (social), and the nominal scaled predictor Class (class); as well as the random (nested) effect of Class within School (school). The data contains 1200 cases evenly distributed among 24 nested groups (4 classes within 6 schools). The data set is available [here](#).

4. Linear Mixed Effects Modeling with package lme4 in R.

4.1. Preparation.

The lme4 (Linear Mixed Effects version 4; [Bates & Maechler, 2010](#)) is designed to analyze linear mixed effects models. The three primary functions are very similar. Function lmer is used to fit linear mixed models, function glmer is used to fit generalized (non-Gaussian) linear mixed models, and function nlmer is used to fit non-linear mixed models. For the purpose of this article, the example used involves a linear mixed model and thus, the lmer function. First, import the data into R using the read table function.

```
> lmm.data <- read.table("http://www.unt.edu/rss/class/Jon/R_SC/Module9/lmm.data.txt",
+   header=TRUE, sep=" ", na.strings="NA", dec=".", strip.white=TRUE)
> summary(lmm.data)
```

id	extro	open	agree
Min. : 1.0	Min. :30.20	Min. :22.30	Min. :18.48
1st Qu.: 300.8	1st Qu.:54.17	1st Qu.:36.20	1st Qu.:31.90
Median : 600.5	Median :60.15	Median :39.98	Median :35.05
Mean : 600.5	Mean :60.27	Mean :40.06	Mean :35.07
3rd Qu.: 900.2	3rd Qu.:66.50	3rd Qu.:43.93	3rd Qu.:38.42
Max. :1200.0	Max. :90.83	Max. :57.87	Max. :58.44

social	class	school
Min. : 46.31	a:300	I :200
1st Qu.: 89.32	b:300	II :200
Median : 99.20	c:300	III:200
Mean : 99.53	d:300	IV :200
3rd Qu.:109.83		V :200
Max. :151.96		VI :200

Next, we need to load the lme4 package.

```
> library(lme4)
Loading required package: Matrix
```

```
Loading required package: lattice
```

```
Attaching package: 'Matrix'
```

```
The following object(s) are masked from 'package:base':
```

```
det
```

```
Attaching package: 'lme4'
```

```
The following object(s) are masked from 'package:stats':
```

```
AIC
```

4.2. Running the Analysis.

Now we are prepared and can proceed to fit the model, which is named “lmm.2”, using the `lmer` function. Some of the optional arguments are shown here, each with the default value specified. For example, the `family = gaussian` argument can be used to specify other distributions (e.g. binomial, poisson, etc.). The `REML = TRUE` argument is used to specify that the REstricted Maximum Likelihood criterion be used rather than the log-likelihood criterion for optimization of parameter estimates. The `verbose = FALSE` argument suppresses the iteration history which if TRUE would display “the iteration number, the value of the deviance (negative twice the log-likelihood) and the value of the parameter s which is the standard deviation of the random effects relative to the standard deviation of the residuals” (Bates, 2010, p. 4). Also note the form of the formula for specifying the model. The formula (from left to right) begins with the outcome variable then the tilde, followed by all the predictors. The first five predictors represent fixed effects and then, in parentheses each random effect is listed. The random effect specifies the nested effect of class within (or under) school; as class would be considered the level one variable and school the level two variable -- which is why the forward slash is used. By default, the `lmer` function will also model the random effect for the highest level variable (school) of the nesting. A standard interaction term can be specified using the colon, for example `(1|school:class)` would specify a random effect (the parentheses) for the interaction of school and class (the colon). Likewise, a fixed effect interaction could be specified with the colon separating the two variables; for example `... + open:agree + open:agree:social + ...` which would specify the interaction of open and agree, then the interaction of open, agree, and social; no parentheses would identify these interactions as fixed effects.

```
> lmm.2 <- lmer(formula = extro ~ open + agree + social + class + (1|school/class), data = lmm.data, family = gaussian, REML = TRUE, verbose = FALSE)
> summary(lmm.2)
```

```
Linear mixed model fit by REML
```

```
Formula: extro ~ open + agree + social + class + (1 | school/class)
```

```
Data: lmm.data
```

	AIC	BIC	logLik	deviance	REMLdev
	3548	3599	-1764	3509	3528

```
Random effects:
```

Groups	Name	Variance	Std.Dev.
class:school	(Intercept)	2.88365	1.69813
school	(Intercept)	95.17339	9.75569
Residual		0.96837	0.98406

```
Number of obs: 1200, groups: class:school, 24; school, 6
```

```
Fixed effects:
```

	Estimate	Std. Error	t value
(Intercept)	57.3838787	4.0559632	14.148
open	0.0061302	0.0049634	1.235
agree	-0.0077361	0.0056985	-1.358
social	0.0005313	0.0018523	0.287
classb	2.0547978	0.9837345	2.089

classc	3.7049300	0.9837165	3.766
classd	5.6657332	0.9837285	5.759

Correlation of Fixed Effects:

	(Intr)	open	agree	social	classb	classc
open	-0.048					
agree	-0.047	-0.012				
social	-0.045	-0.006	-0.009			
classb	-0.121	-0.002	-0.006	0.005		
classc	-0.121	-0.001	-0.005	0.001	0.500	
classd	-0.121	0.000	-0.007	0.002	0.500	0.500

4.3. Interpreting the Default Summary Output.

The output (above) begins by showing what was done; a linear mixed model was fit using REML criterion and the model (formula) and data are listed. Next, two rows of fit statistics are shown; beginning with the Akaike Information Criterion (AIC; [Akaike, 1974](#)) followed by the Bayesian Information Criterion (BIC; [Schwarz, 1978](#)), the log-likelihood, the deviance for the maximum likelihood criterion (smaller deviance indicates better fit), and the deviance for the REML criterion. Generally I tend to use and recommend the BIC for comparing models and assessing fit; the lower the BIC the better the model fits the data (e.g., a BIC of -55.22 indicates a better fitting model than one with a BIC of +23.56). One common way to test the model's fit is to rerun the analysis but include only the intercept terms which is often called the null model and compare the BIC of that model to the hypothesized (full) model BIC.

The next section of the output provides estimates for the random effects in the form of variances and standard deviations. Notice that there are three values shown; the nested effect of class within school, the random effect of the higher level variable, school and the residual term which represents error. The variance estimates are of interest here because we can add them together to find the total variance (of the random effects) and then divide that total by each random effect to see what proportion of the random effect variance is attributable to each random effect (similar to R^2 in traditional regression). So, if we add the variance components:

```
> 2.88365 + 95.17339 + 0.96837
[1] 99.02541
```

Then we can divide this total variance by our nested effect variance to give us the proportion of variance accounted for, which indicates whether or not this effect is meaningful.

```
> 2.88365/99.02541
[1] 0.02912030
```

So, we can see that only 2.9% of the total variance of the random effects is attributed to the nested effect. If all the percentages for each random effect are very small, then the random effects are not present and linear mixed modeling is not appropriate (i.e. remove the random effects from the model and use general linear or generalized linear modeling instead). We can see that the effect of school alone is quite substantial (96%):

```
> 95.17339/99.02541
[1] 0.9611007
```

Another way to think about these variance components is in terms used with standard Analysis of Variance (ANOVA). The residual variance estimate can be thought of as the within groups variance and each random effect variance estimate can be thought of as a between groups estimate (recall the ubiquitous ANOVA summary table).

The next section of the output details the estimates of the fixed effects. These estimates are interpreted the same way as one would interpret estimates from a traditional ordinary least squares linear regression. They are

interpreted as the constant (intercept) and slopes of each fixed effect predictor. The intercept is interpreted as the mean of the outcome (extro) when all the predictors have a value of zero. The predictor estimates (coefficients or slopes) are interpreted the same way as the coefficients from a traditional regression. For instance, a one unit increase in the predictor Openness to new experiences (open) corresponds to a 0.0061302 increase in the outcome Extroversion (extro). Likewise, a one unit increase in the predictor Agreeableness (agree) corresponds to a 0.0077361 *decrease* in the outcome Extroversion (extro). Furthermore, the categorical predictor classb has a coefficient of 2.0547978; which means, the mean Extroversion score of the second group of class (b) is 2.0547978 higher than the mean Extroversion score of the first group of class (a). Class (a) was automatically coded as the reference category by the lmer function because, like most R functions the category with the lower numeric value (or alphabetically first letter) is coded as the reference category. This is very important to note because, both SPSS and SAS use the opposite strategy; they code categorical variables so that the reference category is the category with the highest numerical value (or alphabetically last letter). This difference in strategies means that output from SPSS and SAS will agree but be very different from output produced using the lmer function in R. The key differences will be with the intercept term (which will be substantially different) and the categorical fixed effects coefficients (which will be similar, but not the same). Of course, the really important thing to note is that those differences then produce very different predicted values. If interested in getting the three programs to match, simply reverse code the categorical variable values in SPSS and SAS versions of the data.

The last section of output simply provides the correlations among the fixed effects variables. This can be used to assess multicollinearity. As we can see in our output (above), the predictors are not related; with the obvious and expected exception of the categories of class. Therefore, multicollinearity is not a concern.

4.4. Extracting Elements of the Output.

The default output shown by the summary function (above) has elements which can be extracted and either viewed or assigned to an object. There are also several other elements of the lmer object which can be extracted and may be useful or meaningful.

To extract the estimates of the fixed effects:

```
> fixef(lmm.2)
(Intercept)          open          agree          social          classb          classc
57.3838786610  0.0061301543 -0.0077360956  0.0005312872  2.0547977919  3.7049300287
classd
5.6657331872
```

To extract the estimates of the random effects:

```
> ranef(lmm.2)
$class:school`
(Intercept)
a:I      -3.4073092
a:II     0.9313800
a:III    1.3514649
a:IV     1.2673700
a:V      1.2019177
a:VI    -1.3448235
b:I      0.3040888
b:II     0.2722975
b:III    0.2902197
b:IV     0.2664209
b:V      0.3434285
b:VI    -1.4764554
c:I      1.3893242
c:II    -0.2505738
```

```

c:III -0.3458363
c:IV -0.2497661
c:V -0.3678312
c:VI -0.1753169
d:I 1.2898957
d:II -1.1384331
d:III -1.3554610
d:IV -1.2252249
d:V -0.9876851
d:VI 3.4169085

```

```

$school
  (Intercept)
I -13.991584
II -6.115677
III -1.967158
IV 1.940334
V 6.264193
VI 13.869891

```

To extract the coefficients for the random effects intercept (2 groups of school) and each group of the random effect factor, which here is a nested set of groups (4 groups of class within 6 groups of school):

```

> coef(lmm.2)
$`class:school`
  (Intercept)      open      agree      social      classb      classc      classd
a:I      53.97657 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
a:II     58.31526 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
a:III    58.73534 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
a:IV     58.65125 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
a:V      58.58580 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
a:VI     56.03906 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:I      57.68797 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:II     57.65618 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:III    57.67410 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:IV     57.65030 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:V      57.72731 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
b:VI     55.90742 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:I      58.77320 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:II     57.13330 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:III    57.03804 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:IV     57.13411 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:V      57.01605 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
c:VI     57.20856 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:I      58.67377 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:II     56.24545 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:III    56.02842 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:IV     56.15865 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:V      56.39619 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
d:VI     60.80079 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733

```

```

$school
  (Intercept)      open      agree      social      classb      classc      classd
I      43.39230 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
II     51.26820 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
III    55.41672 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
IV     59.32421 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
V      63.64807 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733
VI     71.25377 0.006130154 -0.007736096 0.0005312872 2.054798 3.70493 5.665733

```

As you can see above, we can further specify using the “\$” to extract just the coefficients for the random effect of school (or just the coefficients for the nested effect \$class:school):

```
> coef(lmm.2)$'school'
```

	(Intercept)	open	agree	social	classb	classc	classd
I	43.39230	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733
II	51.26820	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733
III	55.41672	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733
IV	59.32421	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733
V	63.64807	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733
VI	71.25377	0.006130154	-0.007736096	0.0005312872	2.054798	3.70493	5.665733

To extract the fitted or predicted values based on the model parameters and data, here the predicted values are assigned the name yhat:

```
> yhat <- fitted(lmm.2)
> summary(yhat)
```

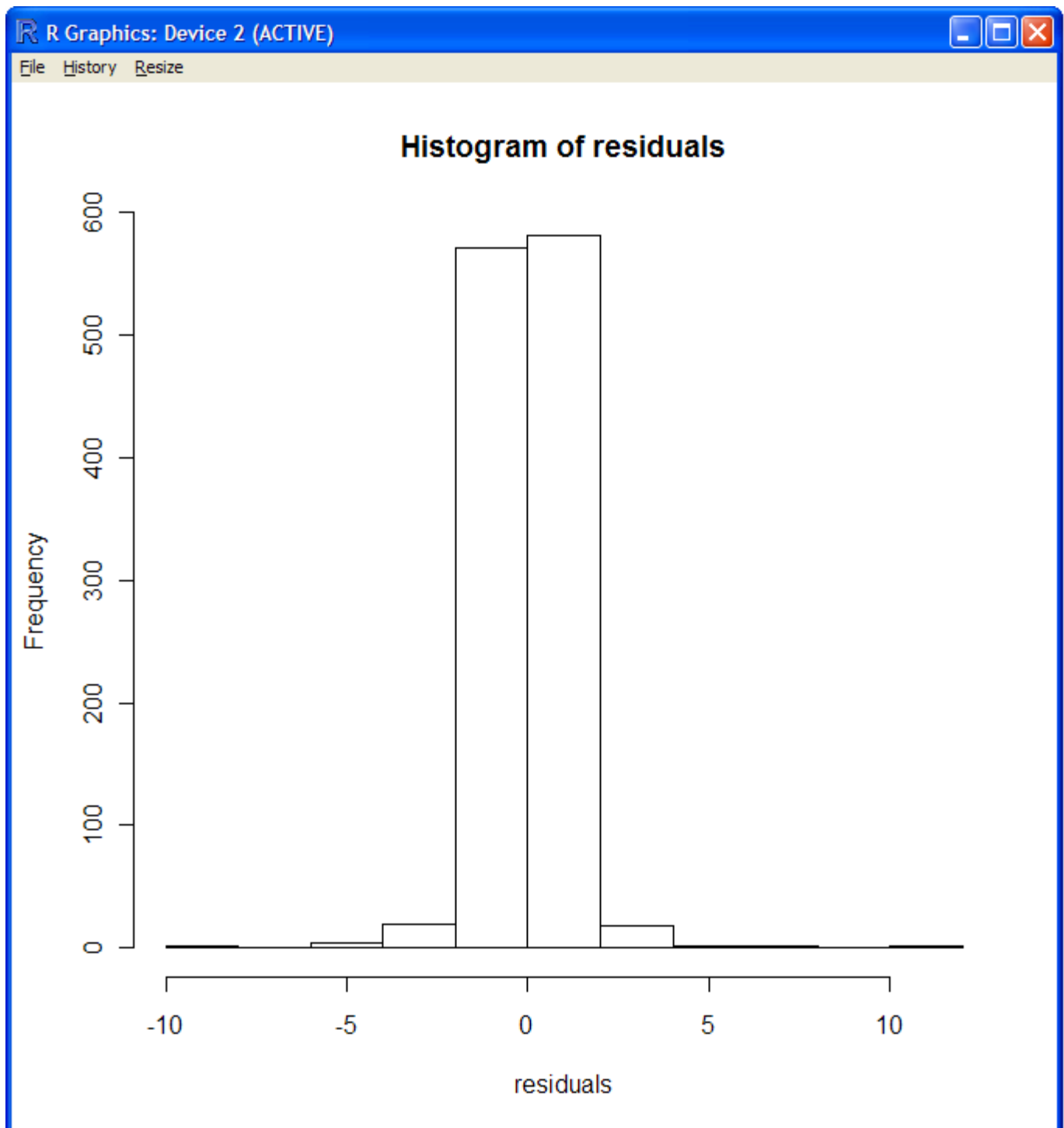
Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
39.91	54.43	60.16	60.27	66.35	80.49

To extract the residuals (errors) and summarize them, as well as plot them (they should be approximately normally distributed around a mean of zero):

```
> residuals <- resid(lmm.2)
> summary(residuals)
```

Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
-9.84100	-0.32980	0.00553	0.00000	0.33520	10.48000

```
> hist(residuals)
```



5. Intra Class Correlation

Intra Class Correlation (ICC) represents a measure of reliability, or dependence among individuals (Kreft & DeLeeuw, 1998). It allows us to assess whether or not the random effect is present in the data. To get the ICC, first create a null model; which for the current example would include just the intercepts (fixed and random) and the random effect for the highest level variable of the nested structure (in this example: school).

```
> lmm.null <- lmer(extro ~ 1 + (1|school), data = lmm.data)
> summary(lmm.null)
Linear mixed model fit by REML
Formula: extro ~ 1 + (1 | school)
```

```

Data: lmm.data
AIC   BIC logLik deviance REMLdev
5812 5827 -2903    5811    5806
Random effects:
Groups   Name             Variance Std.Dev.
school   (Intercept) 95.8720   9.7914
Residual              7.1399   2.6721
Number of obs: 1200, groups: school, 6

Fixed effects:
              Estimate Std. Error t value
(Intercept)    60.267      3.997    15.08

```

Next, add the random effect variance estimates and then divide the random effect of school's variance estimate by the total variance estimate.

```

> 95.8720 + 7.1399
[1] 103.0119
> 95.8720 / 103.0119
[1] 0.9306886

```

So, we see that the ICC is .9306886 (verified below). Another way to get the ICC is with the [multilevel](#) package (Bliese, 2009). First, conduct a standard one way ANOVA using the base 'aov' function.

```

> aov.1 <- aov(extro ~ school, lmm.data)
> summary(aov.1)
          Df Sum Sq Mean Sq F value    Pr(>F)
school      5  95908 19181.5  2686.5 < 2.2e-16 ***
Residuals 1194   8525     7.1
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

```

Then load the 'multilevel' library so that we can use the 'ICC1' and 'ICC2' functions.

```

> library(multilevel)
Loading required package: nlme

Attaching package: 'nlme'

```

The following object(s) are masked from 'package:lme4':

```
BIC, fixef, lmList, ranef, VarCorr
```

```

Loading required package: MASS

```

Next, we can run the ICC1 function to obtain the Intra Class Correlation (which matches the value from above).

```

> ICC1(aov.1)
[1] 0.930689

```

ICC1 indicates that 93.07% of the variance in 'extro' can be "explained" by school group membership.

We can also get the ICC2, which is a measure of reliability.

```

> ICC2(aov.1)
[1] 0.9996278

```

The ICC2 value of .9996 indicates that school groups can be very reliably differentiated in terms of ‘extro’ scores. Remember to detach the multilevel package before continuing with the next section.

```
> detach("package:multilevel")
```

6. Using function mcmcscamp.

Markov Chain Monte Carlo (MCMC) methods represent a type of bridge between traditional frequentist methods and Bayesian methods. MCMC is a type of iterative estimation technique which is used to build an empirical distribution of statistical parameter estimates. The MCMC method can be applied to many types of statistics; such as a *t*-test value, an *F* value, or the multiple parameters of a model such as the linear mixed model used here. MCMC methods use prior information to provide initial parameter estimates in order to evaluate subsequent iteratively re-modeled parameter estimates (e.g. using Bayesian computational methods). The prior information is simply called the *prior* and is generally taken in the form of a distribution (or several distributions if multiple parameters are being estimated). The prior represents a kind of reference point around which iteratively produced parameter estimates are evaluated; thus ensuring convergence on the best set of estimates possible. The empirical distribution of estimates which gets created during a MCMC method is called the *posterior* distribution, because it is created after the data was originally fitted.

To obtain a simulated empirical distribution or posterior distribution (here with $n = 5000$) of estimates based on the specified lmer model using MCMC methods:

```
> mcmc.5000 <- mcmcscamp(lmm.2, saveb = TRUE, n = 5000)
```

To then show the structure of elements of the MCMC object; this simply shows how you can then extract elements of the MCMC object using the MCMC object name and “@” (examples are further below):

```
> str(mcmc.5000)
Formal class 'merMCMC' [package "lme4"] with 9 slots
 ..@ Gp      : int [1:3] 0 24 30
 ..@ ST      : num [1:2, 1:5000] 1.73 9.91 1.68 9.96 1.67 ...
 ..@ call    : language lmer(formula = extro ~ open + agree + social + class + (1 |
school/class),
data = lmm.data)
 ..@ deviance: num [1:5000] 3509 3509 3509 3509 3509 ...
 ..@ dims    : Named int [1:18] 2 1200 7 30 1 2 0 1 2 5 ...
 .. ..- attr(*, "names")= chr [1:18] "nt" "n" "p" "q" ...
 ..@ fixef   : num [1:7, 1:5000] 57.383879 0.00613 -0.007736 0.000531 2.054798 ...
 .. ..- attr(*, "dimnames")=List of 2
 .. .. ..$ : chr [1:7] "(Intercept)" "open" "agree" "social" ...
 .. .. ..$ : NULL
 ..@ nc      : int [1:2] 1 1
 ..@ ranef   : num [1:30, 1:5000] -3.407 0.931 1.351 1.267 1.202 ...
 ..@ sigma   : num [1, 1:5000] 0.984 0.968 1.003 0.996 0.98 ...
```

To extract the fixed effect parameter estimates from the MCMC object (output of the matrix of 5000 by 7 parameter estimates not shown):

```
> mcmc.5000@fixef
```

To extract the random effect parameter estimates from the MCMC object (output of the matrix of parameter estimates not shown):

```
> mcmc.5000@ranef
```

Deviance is a measure of fit; the smaller the deviance statistic, the better the model fits the data. To extract and summarize the Maximum Likelihood Deviance:

```
> dev <- as.vector(mcmc.5000@deviance)
> summary(dev)
   Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
  3509   3808   3828   3807   3846   3918
```

To show the Highest Posterior Density (HPD) intervals for the parameters of an MCMC distribution (which essentially provides confidence intervals for the posterior parameters):

```
> HPDinterval(mcmc.5000, prob = 0.95)
$fixef
      lower      upper
(Intercept) 55.648253310 58.975371325
open        -0.005528170  0.016827333
agree       -0.020263349  0.005456398
social      -0.003820699  0.004483850
classb       0.843105909  3.165402291
classc       2.509578421  4.816869867
classd       4.503296720  6.840049292
attr(,"Probability")
[1] 0.95
```

```
$ST
      lower      upper
[1,] 0.7267248 1.018011
[2,] 0.8629657 3.426454
attr(,"Probability")
[1] 0.95
```

```
$sigma
      lower      upper
[1,] 1.014880 1.21526
attr(,"Probability")
[1] 0.95
```

```
$ranef
      lower      upper
[1,] -6.8912029 -3.1898835
[2,] -1.2099349  1.4023157
[3,] -0.1084854  2.2773701
[4,]  0.2550491  2.6539355
[5,]  0.7257380  3.3318154
[6,] -1.4673513  2.0908638
[7,] -3.2369864  0.3176006
[8,] -1.8458023  0.7151407
[9,] -1.1854551  1.1665241
[10,] -0.6519829  1.6672193
[11,] -0.1414238  2.5049994
[12,] -1.4623917  2.0688023
[13,] -2.2266973  1.3912863
[14,] -2.3333123  0.2616849
[15,] -1.7605533  0.6277444
[16,] -1.1287503  1.2296900
[17,] -0.8618493  1.7422243
[18,] -0.2751656  3.3046945
[19,] -2.3236627  1.2985081
[20,] -3.3204482 -0.7280814
[21,] -2.7984293 -0.3895353
[22,] -2.0954523  0.2725904
[23,] -1.4421025  1.1768267
```

```
[24,] 3.2136308 6.7736435
[25,] -13.9539192 -10.1167238
[26,] -6.6835666 -3.5847653
[27,] -3.2718712 -0.2672930
[28,] 0.1375870 3.1094643
[29,] 3.8363158 6.9669211
[30,] 9.9230090 13.8334339
attr(,"Probability")
[1] 0.95
```

As with some of the objects above, we can use the “\$” to extract elements of the HPD interval output; here extracting just the intervals for the fixed effects (\$fixef):

```
> HPDinterval(mcmc.5000, prob = 0.95)$fixef
      lower      upper
(Intercept) 55.648253310 58.975371325
open        -0.005528170  0.016827333
agree       -0.020263349  0.005456398
social      -0.003820699  0.004483850
classb       0.843105909  3.165402291
classc       2.509578421  4.816869867
classd       4.503296720  6.840049292
attr(,"Probability")
[1] 0.95
```

6. Alternatives

As mentioned at the beginning of this article, there are other R packages available and in development for doing mixed effect modeling (linear and otherwise) or multilevel modeling; some of these alternatives are also capable of doing MCMC methods on the fitted model. A few of those alternatives are [HGLMMM](#), [MCMCglmm](#), and [multilevel](#). However, the [lme4](#) package represents the long term (and continued) development of the [nlme](#) package for mixed effects modeling which has been developed and used for more than 10 years.

References

- Akaike, H. (1974). A new look at the statistical model identification. *I.E.E.E. Transactions on Automatic Control*, AC 19, 716 – 723. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Akaike_1974.pdf
- Bates, D., & Maechler, M. (2010). Package ‘lme4’. Reference manual for the package, available at:
<http://cran.r-project.org/web/packages/lme4/lme4.pdf>
- Bates, D. (2010). Linear mixed model implementation in lme4. Package lme4 vignette, available at:
<http://cran.r-project.org/web/packages/lme4/vignettes/Implementation.pdf>
- Bliese, P. (2009). Package ‘multilevel’. Reference manual for the package, available at:
<http://cran.r-project.org/web/packages/multilevel/index.html>
- Gelman, A. (2005). Analysis of variance -- why it is more important than ever. *The Annals of Statistics*, 33(1), 1 -- 53. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Gelman_2005.pdf
- Kreft, I., & DeLeeuw, J. (1998). *Introducing multilevel modeling*. London: Sage Publications Ltd.

Schwarz, G. (1978). Estimating the dimension of a model. *Annals of Statistics*, 6, 461 – 464. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Schwarz_1978.pdf

Additional Resources

Bates, D. (2010). Computational methods for mixed models. Package lme4 vignette, available at:
<http://cran.r-project.org/web/packages/lme4/vignettes/Theory.pdf>

Bates, D. (2010). Penalized least squares versus generalized least squares representations of linear mixed models. Package lme4 vignette, available at:
<http://cran.r-project.org/web/packages/lme4/vignettes/PLSvGLS.pdf>

Bliese, P. (2009). Multilevel modeling in R: A brief introduction to R, the multilevel package and the nlme package. Available at:
http://cran.r-project.org/doc/contrib/Bliese_Multilevel.pdf

Draper, D. (1995). Inference and hierarchical modeling in the social sciences. *Journal of Educational and Behavioral Statistics*, 20(2), 115 - 147. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Draper_1995.pdf

Fox, J. (2002). Linear mixed models: An appendix to “An R and S-PLUS companion to applied regression”. Available at:
<http://cran.r-project.org/doc/contrib/Fox-Companion/appendix-mixed-models.pdf>

Hofmann, D. A., Griffin, M. A., & Gavin, M. B. (2000). The application of hierarchical linear modeling to organizational research. In K. J. Klein (Ed.), *Multilevel theory, research, and methods in organizations: Foundations, extensions, and new directions* (p. 467 - 511). San Francisco, CA: Jossey-Bass. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Hofmann_2000.pdf

Raudenbush, S. W. (1995). Reexamining, reaffirming, and improving application of hierarchical models. *Journal of Educational and Behavioral Statistics*, 20(2), 210 - 220. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Raudenbush_1995.pdf

Raudenbush, S. W. (1993). Hierarchical linear models and experimental design. In L. Edwards (Ed.), *Applied analysis of variance in behavioral science* (p. 459 - 496). New York: Marcel Dekker. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Raudenbush_1993.pdf

Rogosa, D., & Saner, H. (1995). Longitudinal data analysis examples with random coefficient models. *Journal of Educational and Behavioral Statistics*, 20(2), 149 - 170. Available at:
http://www.unt.edu/rss/class/Jon/MiscDocs/Rogosa_1995.pdf

Until next time; *Freedom is just another word for nothing left to lose...*