



Fitting linear mixed-effects models using lme4

Douglas Bates
University of Wisconsin - Madison

Martin Mächler
ETH Zurich

Ben Bolker
McMaster University

Abstract

Maximum likelihood or REML estimates of the parameters in linear mixed-effects models can be determined using the `lmer` function in the **lme4** package for R. As in most model-fitting functions, the model is described in an `lmer` call by a formula, in this case including both fixed-effects terms and random-effects terms. The formula and data together determine a numerical representation of the model from which the profiled deviance or the profiled REML criterion can be evaluated as a function of some of the model parameters. The appropriate criterion is optimized, using one of the constrained optimization functions in R, to provide the parameter estimates. We describe the structure of the model, the steps in evaluating the profiled deviance or REML criterion and the structure of the S4 class that represents such a model. Sufficient detail is included to allow specialization of these structures by those who wish to write functions to fit specialized linear mixed models, such as models incorporating pedigrees or smoothing splines, that aren't easily expressible in the formula language used by `lmer`.

Keywords: sparse matrix methods, linear mixed models, penalized least squares, Cholesky decomposition.

1. Introduction

The **lme4** package for R provides functions to fit and analyze linear mixed models (LMMs), generalized linear mixed models (GLMMs) and nonlinear mixed models (NLMMs). In each of these names, the term “mixed” or, more fully, “mixed-effects”, denotes a model that incorporates both fixed-effects terms and random-effects terms in a linear predictor expression from which the conditional mean of the response can be evaluated. In this paper we describe the formulation and representation of linear and generalized linear mixed models. The techniques used for nonlinear mixed models will be described separately.

Just as a linear model can be described in terms of the distribution of $\mathcal{Y}$, the vector-valued random variable whose observed value is $\mathbf{y}_{\text{obs}}$, the observed response vector, a linear mixed

model can be described by the distribution of two vector-valued random variables: $\mathcal{Y}$, the response and $\mathcal{B}$, the vector of random effects. In a linear model the distribution of $\mathcal{Y}$ is multivariate normal,

$$\mathcal{Y} \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}_n), \quad (1)$$

where n is the dimension of the response vector, $\mathbf{I}_n$ is the identity matrix of size n , $\boldsymbol{\beta}$ is a p -dimensional coefficient vector and $\mathbf{X}$ is an $n \times p$ model matrix. The parameters of the model are the coefficients, $\boldsymbol{\beta}$, and the scale parameter, σ .

In a linear mixed model it is the *conditional* distribution of $\mathcal{Y}$ given $\mathcal{B} = \mathbf{b}$ that has such a form,

$$(\mathcal{Y}|\mathcal{B} = \mathbf{b}) \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{b}, \sigma^2 \mathbf{I}_n) \quad (2)$$

where $\mathbf{Z}$ is the $n \times q$ model matrix for the q -dimensional vector-valued random effects variable, $\mathcal{B}$, whose value we are fixing at $\mathbf{b}$. The unconditional distribution of $\mathcal{B}$ is also multivariate normal with mean zero and a parameterized $q \times q$ variance-covariance matrix, $\boldsymbol{\Sigma}$,

$$\mathcal{B} \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}). \quad (3)$$

As a variance-covariance matrix, $\boldsymbol{\Sigma}$ must be positive semidefinite. It is convenient to express the model in terms of a *relative covariance factor*, $\boldsymbol{\Lambda}_\theta$, which is a $q \times q$ matrix, depending on the *variance-component parameter*, $\boldsymbol{\theta}$, and generating the symmetric $q \times q$ variance-covariance matrix, $\boldsymbol{\Sigma}$, according to

$$\boldsymbol{\Sigma}_\theta = \sigma^2 \boldsymbol{\Lambda}_\theta \boldsymbol{\Lambda}_\theta', \quad (4)$$

where σ is the same scale factor as in the conditional distribution (2).

Although q , the number of columns in $\mathbf{Z}$ and the size of $\boldsymbol{\Sigma}_\theta$, can be very large indeed, the dimension of $\boldsymbol{\theta}$ is small, frequently less than 10.

In calls to the `lm` function for fitting linear models the form of the model matrix $\mathbf{X}$ is determined by the `formula` and `data` arguments. The right-hand side of the formula consists of one or more terms that each generate one or more columns in the model matrix, $\mathbf{X}$. For `lmer` the formula language is extended to allow for random-effects terms that generate the model matrix $\mathbf{Z}$ and the mapping from $\boldsymbol{\theta}$ to $\boldsymbol{\Lambda}_\theta$.

To understand why the formulation in equations 2 and 3 is particularly useful, we first show that the profiled deviance (negative twice the log-likelihood) and the profiled REML criterion can be expressed as a function of $\boldsymbol{\theta}$ only. Furthermore these criteria can be evaluated quickly and accurately.

2. Profiling the deviance and the REML criterion

As stated above, $\boldsymbol{\theta}$ determines the $q \times q$ matrix $\boldsymbol{\Lambda}_\theta$ which, together with a value of σ^2 , determines $\text{Var}(\mathcal{B}) = \boldsymbol{\Sigma}_\theta = \sigma^2 \boldsymbol{\Lambda}_\theta \boldsymbol{\Lambda}_\theta'$. If we define a *spherical*¹ random effects variable, $\mathcal{U}$, with distribution

$$\mathcal{U} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}_q), \quad (5)$$

and set

$$\mathcal{B} = \boldsymbol{\Lambda}_\theta \mathcal{U}, \quad (6)$$

¹ $\mathcal{N}(\boldsymbol{\mu}, \sigma^2 \mathbf{I})$ distributions are called “spherical” because contours of the probability density are spheres.

then $\mathcal{B}$ will have the desired $\mathcal{N}(\mathbf{0}, \Sigma_\theta)$ distribution.

Although it may seem more natural to define $\mathcal{U}$ in terms of $\mathcal{B}$ we must write the relationship as in eqn.~6 because Λ_θ may be singular. In fact, it is important to allow for Λ_θ to be singular because situations where the parameter estimates, $\hat{\theta}$, produce a singular $\Lambda_{\hat{\theta}}$ do occur in practice. And even if the parameter estimates do not correspond to a singular Λ_θ , it may be necessary to evaluate the estimation criterion at such values during the course of the numerical optimization of the criterion.

The model can now be defined in terms of

$$(\mathcal{Y}|\mathcal{U} = \mathbf{u}) \sim \mathcal{N}(\mathbf{Z}\Lambda_\theta\mathbf{u} + \mathbf{X}\beta, \sigma^2\mathbf{I}_n) \quad (7)$$

producing the joint density function

$$\begin{aligned} f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}, \mathbf{u}) &= f_{\mathcal{Y}|\mathcal{U}}(\mathbf{y}|\mathbf{u}) f_{\mathcal{U}}(\mathbf{u}) \\ &= \frac{\exp(-\frac{1}{2\sigma^2}\|\mathbf{y} - \mathbf{X}\beta - \mathbf{Z}\Lambda_\theta\mathbf{u}\|^2)}{(2\pi\sigma^2)^{n/2}} \frac{\exp(-\frac{1}{2\sigma^2}\|\mathbf{u}\|^2)}{(2\pi\sigma^2)^{q/2}} \\ &= \frac{\exp(-[\|\mathbf{y} - \mathbf{X}\beta - \mathbf{Z}\Lambda_\theta\mathbf{u}\|^2 + \|\mathbf{u}\|^2]/(2\sigma^2))}{(2\pi\sigma^2)^{(n+q)/2}}. \end{aligned} \quad (8)$$

The *likelihood* of the parameters, θ , β and σ^2 , given the observed data is the value of the marginal density of $\mathcal{Y}$, evaluated at $\mathbf{y}_{\text{obs}}$. That is

$$L(\theta, \beta, \sigma^2|\mathbf{y}_{\text{obs}}) = \int_{\mathbb{R}^q} f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u}) d\mathbf{u}. \quad (9)$$

The integrand of eqn.~9 is the *unscaled conditional density* of $\mathcal{U}$ given $\mathcal{Y} = \mathbf{y}_{\text{obs}}$. The conditional density of $\mathcal{U}$ given $\mathcal{Y} = \mathbf{y}_{\text{obs}}$ is

$$f_{\mathcal{U}|\mathcal{Y}}(\mathbf{u}|\mathbf{y}_{\text{obs}}) = \frac{f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u})}{\int f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u}) d\mathbf{u}} \quad (10)$$

which is, up to a scale factor, the joint density, $f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u})$. The unscaled conditional density will be, up to a scale factor, a q -dimensional multivariate Gaussian with an integral that is easily evaluated if we know the mean and variance-covariance of the conditional density.

The conditional mean, $\mu_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$, is also the mode of the conditional distribution. Because a constant factor in a function does not affect the location of the optimum, we can determine the conditional mode, and hence the conditional mean, by maximizing the unscaled conditional density. This is in the form of a *penalized linear least squares* problem,

$$\mu_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \arg \min_{\mathbf{u}} \left(\|\mathbf{y}_{\text{obs}} - \mathbf{X}\beta - \mathbf{Z}\Lambda_\theta\mathbf{u}\|^2 + \|\mathbf{u}\|^2 \right). \quad (11)$$

2.1. Solving the penalized least squares problem

In the so-called “pseudo-data” approach to penalized least squares problems we write the objective as a residual sum of squares for an extended response vector and model matrix

$$\|\mathbf{y}_{\text{obs}} - \mathbf{X}\beta - \mathbf{Z}\Lambda_\theta\mathbf{u}\|^2 + \|\mathbf{u}\|^2 = \left\| \begin{bmatrix} \mathbf{y}_{\text{obs}} - \mathbf{X}\beta \\ \mathbf{0} \end{bmatrix} - \begin{bmatrix} \mathbf{Z}\Lambda_\theta \\ \mathbf{I}_q \end{bmatrix} \mathbf{u} \right\|^2. \quad (12)$$

The contribution to the residual sum of squares from the “pseudo” observations appended to $\mathbf{y}_{\text{obs}} - \mathbf{X}\boldsymbol{\beta}$, is exactly the penalty term, $\|\mathbf{u}\|^2$.

From eqn.~12 we can see that the conditional mean satisfies

$$(\boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' \mathbf{Z} \boldsymbol{\Lambda}_{\theta} + \mathbf{I}_q) \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' (\mathbf{y}_{\text{obs}} - \mathbf{X}\boldsymbol{\beta}), \quad (13)$$

which would be interesting, but not terribly useful, were it not for the fact that we can determine the solution to eqn.~13 quickly and accurately, even when q , the size of the system to solve, is very large indeed. (We have done so in cases where q is in the millions.)

The key to solving eqn.~13 is the *sparse Cholesky factor*, $\mathbf{L}_{\theta}$, which is a sparse, lower-triangular matrix such that

$$\mathbf{L}_{\theta} \mathbf{L}'_{\theta} = \mathbf{P} (\boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' \mathbf{Z} \boldsymbol{\Lambda}_{\theta} + \mathbf{I}_q) \mathbf{P}', \quad (14)$$

where $\mathbf{P}$ is a permutation matrix representing a fill-reducing permutation~(Davis 2006, Ch.~7).

As for most sparse matrix methods, the sparse Cholesky factorization can be split into two phases: a symbolic phase in which the positions of the non-zero elements in the result are determined and a numeric phase in which the actual numeric values in these positions are determined. Determining the fill-reducing permutation represented by $\mathbf{P}$ is part of the symbolic phase, which often takes much longer than the numeric phase. During the course of determining the maximum likelihood or REML estimates of the parameters in a linear mixed-effects model we may need to evaluate $\mathbf{L}_{\theta}$ for many different values of $\boldsymbol{\theta}$, but each evaluation after the first requires only the numeric phase. Changing $\boldsymbol{\theta}$ can change the values of the non-zero elements in $\mathbf{L}$ but does not change their positions. Hence, the symbolic phase must be done only once.

The `Cholesky` function in the **Matrix** package for R performs both the symbolic and numeric phases of the factorization to produce $\mathbf{L}_{\theta}$ from $\boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' \mathbf{Z} \boldsymbol{\Lambda}_{\theta}$. The resulting object has S4 class “CHMsuper” or “CHMsimp” depending on whether it is in the supernodal~(Davis 2006, §4.8) or simplicial form. Both these classes inherit from the virtual class “CHMfactor”. Optional arguments to the `Cholesky` function control determination of a fill-reducing permutation and addition of multiple of the identity to the symmetric matrix before factorization. Once the factor has been determined for the initial value, $\boldsymbol{\theta}_0$, it can be updated for new values of $\boldsymbol{\theta}$ in a single call to the `update` method.

Although the `solve` method for the “CHMfactor” class has an option to evaluate $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ directly as the solution to

$$\mathbf{P}' \mathbf{L}_{\theta} \mathbf{L}'_{\theta} \mathbf{P} \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}). \quad (15)$$

we will express the solution in two stages:

1. Solve $\mathbf{L} \mathbf{c}_{\mathbf{u}} = \mathbf{P} \boldsymbol{\Lambda}'_{\theta} \mathbf{Z}' (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})$ for $\mathbf{c}_{\mathbf{u}}$.
2. Solve $\mathbf{L}' \mathbf{P} \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \mathbf{c}_{\mathbf{u}}$ for $\mathbf{P} \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ and then for $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \mathbf{P}' (\mathbf{P} \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}})$.

2.2. Evaluating the likelihood

After solving for $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ the exponent in $f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u})$ can be written

$$\|\mathbf{y}_{\text{obs}} - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\boldsymbol{\Lambda}_{\theta}\mathbf{u}\|^2 + \|\mathbf{u}\|^2 = r^2(\boldsymbol{\theta}, \boldsymbol{\beta}) + \|\mathbf{L}' \mathbf{P} (\mathbf{u} - \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}})\|^2. \quad (16)$$

where $r^2(\boldsymbol{\theta}, \boldsymbol{\beta}) = \|\mathbf{y}_{\text{obs}} - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\boldsymbol{\Lambda}_\theta \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}\|^2 + \|\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}\|^2$, is the minimum penalized residual sum of squares for these values of $\boldsymbol{\theta}$ and $\boldsymbol{\beta}$.

With expression (16) and the change of variable $\mathbf{v} = \mathbf{L}'\mathbf{P}(\mathbf{u} - \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}})$, for which $d\mathbf{v} = \text{abs}(|\mathbf{L}||\mathbf{P}|) d\mathbf{u}$, we have

$$\int \frac{\exp\left(\frac{-\|\mathbf{L}'\mathbf{P}(\mathbf{u} - \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}})\|^2}{2\sigma^2}\right)}{(2\pi\sigma^2)^{q/2}} d\mathbf{u} = \int \frac{\exp\left(\frac{-\|\mathbf{v}\|^2}{2\sigma^2}\right)}{(2\pi\sigma^2)^{q/2}} \frac{d\mathbf{v}}{\text{abs}(|\mathbf{L}||\mathbf{P}|)} = \frac{1}{\text{abs}(|\mathbf{L}||\mathbf{P}|)} = \frac{1}{|\mathbf{L}|} \quad (17)$$

because $\text{abs}|\mathbf{P}| = 1$ (one property of a permutation matrix is $|\mathbf{P}| = \pm 1$) and $|\mathbf{L}|$, which, because $\mathbf{L}$ is triangular, is the product of its diagonal elements, all of which are positive, is positive.

Using this expression we can write the deviance (negative twice the log-likelihood) as

$$-2\ell(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2|\mathbf{y}_{\text{obs}}) = -2\log L(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2|\mathbf{y}_{\text{obs}}) = n\log(2\pi\sigma^2) + \frac{r^2(\boldsymbol{\theta}, \boldsymbol{\beta})}{\sigma^2} + \log(|\mathbf{L}_\theta|^2) \quad (18)$$

Because the dependence of eqn.~18 on σ^2 is straightforward, we can form the conditional estimate

$$\widehat{\sigma^2}(\boldsymbol{\theta}, \boldsymbol{\beta}) = \frac{r^2(\boldsymbol{\theta}, \boldsymbol{\beta})}{n} m \quad (19)$$

producing the *profiled deviance*

$$-2\tilde{\ell}(\boldsymbol{\theta}, \boldsymbol{\beta}|\mathbf{y}_{\text{obs}}) = \log(|\mathbf{L}_\theta|^2) + n \left[1 + \log\left(\frac{2\pi r^2(\boldsymbol{\theta}, \boldsymbol{\beta})}{n}\right) \right] \quad (20)$$

However, observing that eqn.~20 depends on $\boldsymbol{\beta}$ only through $r^2(\boldsymbol{\theta}, \boldsymbol{\beta})$ provides a much greater simplification because it allows us to “profile out” the fixed-effects parameter, $\boldsymbol{\beta}$, from the evaluation of the deviance. The conditional estimate, $\widehat{\boldsymbol{\beta}}_\theta$, is the value of $\boldsymbol{\beta}$ at the solution of the joint penalized least squares problem

$$r_\theta^2 = \min_{\mathbf{u}, \boldsymbol{\beta}} \left(\|\mathbf{y} - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\boldsymbol{\Lambda}_\theta \mathbf{u}\|^2 + \|\mathbf{u}\|^2 \right), \quad (21)$$

producing the profiled deviance,

$$-2\tilde{\ell}(\boldsymbol{\theta}) = \log(|\mathbf{L}_\theta|^2) + n \left[1 + \log\left(\frac{2\pi r_\theta^2}{n}\right) \right], \quad (22)$$

which is a function of $\boldsymbol{\theta}$ only. Eqn.~22 is a remarkably compact expression for the deviance.

2.3. Solving the joint penalized least squares problem

The solutions, $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ and $\widehat{\boldsymbol{\beta}}_\theta$, of the joint penalized least squares problem (21) satisfy

$$\begin{bmatrix} \boldsymbol{\Lambda}'_\theta \mathbf{Z}' \mathbf{Z} \boldsymbol{\Lambda}_\theta + \mathbf{I}_q & \boldsymbol{\Lambda}'_\theta \mathbf{Z}' \mathbf{X} \\ \mathbf{X}' \mathbf{Z} \boldsymbol{\Lambda}_\theta & \mathbf{X}' \mathbf{X} \end{bmatrix} \begin{bmatrix} \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} \\ \widehat{\boldsymbol{\beta}}_\theta \end{bmatrix} = \begin{bmatrix} \boldsymbol{\Lambda}'_\theta \mathbf{Z}' \mathbf{y}_{\text{obs}} \\ \mathbf{X}' \mathbf{y}_{\text{obs}} \end{bmatrix} \quad (23)$$

As before we will use the sparse Cholesky decomposition producing, $\mathbf{L}_\theta$, the sparse Cholesky factor, and $\mathbf{P}$, the permutation matrix, satisfying $\mathbf{L}_\theta \mathbf{L}'_\theta = \mathbf{P}(\boldsymbol{\Lambda}'_\theta \mathbf{Z}' \mathbf{Z} \boldsymbol{\Lambda}_\theta + \mathbf{I})\mathbf{P}'$ and $\mathbf{c}_u$, the solution to $\mathbf{L}_\theta \mathbf{c}_u = \mathbf{P} \boldsymbol{\Lambda}'_\theta \mathbf{Z}' \mathbf{y}_{\text{obs}}$.

We extend the decomposition with the $q \times p$ matrix $\mathbf{R}_{ZX}$, the upper triangular $p \times p$ matrix $\mathbf{R}_X$, and the p -vector $\mathbf{c}_\beta$ satisfying

$$\begin{aligned} \mathbf{L}\mathbf{R}_{ZX} &= \mathbf{P}\mathbf{\Lambda}'_\theta\mathbf{Z}'\mathbf{X} \\ \mathbf{R}'_X\mathbf{R}_X &= \mathbf{X}'\mathbf{X} - \mathbf{R}'_{ZX}\mathbf{R}_{ZX} \\ \mathbf{R}'_X\mathbf{c}_\beta &= \mathbf{X}'\mathbf{y}_{\text{obs}} - \mathbf{R}'_{ZX}\mathbf{c}_u \end{aligned}$$

so that

$$\begin{bmatrix} \mathbf{P}'\mathbf{L} & \mathbf{0} \\ \mathbf{R}'_{ZX} & \mathbf{R}'_X \end{bmatrix} \begin{bmatrix} \mathbf{L}'\mathbf{P} & \mathbf{R}_{ZX} \\ \mathbf{0} & \mathbf{R}_X \end{bmatrix} = \begin{bmatrix} \mathbf{\Lambda}'_\theta\mathbf{Z}'\mathbf{Z}\mathbf{\Lambda}_\theta + \mathbf{I} & \mathbf{\Lambda}'_\theta\mathbf{Z}'\mathbf{X} \\ \mathbf{X}'\mathbf{Z}\mathbf{\Lambda}_\theta & \mathbf{X}'\mathbf{X} \end{bmatrix}, \quad (24)$$

and the solutions, $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ and $\hat{\boldsymbol{\beta}}_\theta$, satisfy

$$\mathbf{R}_X\hat{\boldsymbol{\beta}}_\theta = \mathbf{c}_\beta \quad (25)$$

$$\mathbf{L}'\mathbf{P}\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} = \mathbf{c}_u - \mathbf{R}_{ZX}\hat{\boldsymbol{\beta}}_\theta. \quad (26)$$

2.4. The profiled REML criterion

Laird and Ware (1982) show that the criterion to be optimized by the REML estimates can be expressed as

$$L_R(\boldsymbol{\theta}, \sigma^2 | \mathbf{y}_{\text{obs}}) = \int L(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 | \mathbf{y}_{\text{obs}}) d\boldsymbol{\beta}. \quad (27)$$

Because the joint solutions, $\boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}}$ and $\hat{\boldsymbol{\beta}}_\theta$, to the penalized least squares problem allow us to express

$$\begin{aligned} \|\mathbf{y}_{\text{obs}} - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{\Lambda}_\theta\mathbf{u}\|^2 + \|\mathbf{u}\|^2 = \\ r_\theta^2 + \left\| \mathbf{L}'\mathbf{P} \left[\mathbf{u} - \boldsymbol{\mu}_{\mathcal{U}|\mathcal{Y}=\mathbf{y}_{\text{obs}}} - \mathbf{R}_{ZX}(\boldsymbol{\beta} - \hat{\boldsymbol{\beta}}_\theta) \right] \right\|^2 + \left\| \mathbf{R}_X(\boldsymbol{\beta} - \hat{\boldsymbol{\beta}}_\theta) \right\|^2 \end{aligned} \quad (28)$$

we can use a change of variable, similar to that in eqn.~17, to evaluate the profiled REML criterion. On the deviance scale the criterion can be evaluated as

$$-2\tilde{\ell}_R(\boldsymbol{\theta}) = \log(|\mathbf{L}|^2) + \log(|\mathbf{R}_X|^2) + (n - p) \left[1 + \log \left(\frac{2\pi r_\theta^2}{n - p} \right) \right]. \quad (29)$$

The structures in **lme4** for representing mixed-models are somewhat more general than is required for linear mixed models. In the remainder of this section we briefly describe some of the computational requirements for generalized linear mixed models, to show why these more general structures are employed.

2.5. Definition of GLMMs

The generalized linear mixed models (GLMMs) that can be fit by the **lme4** package preserve the multivariate Gaussian unconditional distribution of the random effects, $\mathcal{B}$ (eqn.~3). Because most families used for the conditional distribution, $\mathcal{Y}|\mathcal{B} = \mathbf{b}$, do not incorporate a separate scale factor, σ , we remove it from the definition of $\boldsymbol{\Sigma}$ and from the distribution of the spherical random effects, $\mathcal{U}$. That is

$$\mathcal{U} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_q) \quad (30)$$

and

$$\Sigma_\theta = \Lambda_\theta \Lambda_\theta'.$$
 (31)

The conditional distributions, $\mathcal{Y}|\mathcal{B} = \mathbf{b}$ and $\mathcal{Y}|\mathcal{U} = \mathbf{u}$, preserve the properties that the components of $\mathcal{Y}$ are conditionally independent and that the mean, $\mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}$, depends on $\mathbf{u}$ only through the linear predictor,

$$\boldsymbol{\eta} = \mathbf{Z}\Lambda_\theta\mathbf{u} + \mathbf{X}\boldsymbol{\beta}.$$
 (32)

The mapping from $\mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}$ to $\boldsymbol{\eta}$, which is called the *link function* and written

$$\mathbf{Z}\Lambda_\theta\mathbf{u} + \mathbf{X}\boldsymbol{\beta} = \boldsymbol{\eta} = \mathbf{g}(\mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}),$$
 (33)

is a *diagonal mapping* in the sense that there is a scalar function, g , such that the i th component of $\boldsymbol{\eta}$ is g applied to the i th component of $\mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}$. (The name “diagonal” reflects the fact that the Jacobian matrix, $\frac{d\boldsymbol{\eta}}{d\boldsymbol{\mu}'}$, of such a mapping will be diagonal.)

The scalar link function must be invertible over its range. The vector-valued *inverse link* function, $\mathbf{g}^{-1}$, will be the scalar inverse link, g^{-1} , applied component-wise to $\boldsymbol{\eta}$.

Common forms of the conditional distribution are Bernoulli, for binary responses, binomial for binary responses that are recorded as the number of trials and the number of successes, and Poisson, for count data. The combination of a distributional form and a link function is called a *family*. For distributional forms in the exponential family there is a *canonical link*. For Bernoulli or binomial forms the canonical link is the *logit* link function

$$\eta_i = \log\left(\frac{\mu_i}{1 - \mu_i}\right);$$
 (34)

for the Poisson distribution the canonical link is the natural logarithm.

The form of the distribution determines the conditional variance, $\text{Var}(\mathcal{Y}|\mathcal{U} = \mathbf{u})$, as a function of the conditional mean and, possibly, a separate scale factor. (In most cases the conditional variance is completely determined by the conditional mean.)

The likelihood of the parameters, given the observed data, is now

$$L(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}_{\text{obs}}) = \int_{\mathbb{R}^q} f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u}) d\mathbf{u}$$
 (35)

where, as in the case of linear mixed models, $f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}_{\text{obs}}, \mathbf{u})$ is the unscaled conditional density of $\mathcal{U}$ given $\mathcal{Y} = \mathbf{y}_{\text{obs}}$. The notation here is a bit blurred because, although the joint distribution of $\mathcal{Y}$ and $\mathcal{U}$ is always continuous with respect to $\mathcal{U}$, it can be (and often is) discrete with respect to $\mathcal{Y}$. However, when we condition on the observed value $\mathcal{Y} = \mathbf{y}_{\text{obs}}$, the resulting function is continuous with respect to $\mathbf{u}$ so the unscaled conditional density is indeed well-defined as a density, up to a scale factor.

2.6. Determining the conditional mode

As for linear mixed models, we simplify evaluation of the integral (35) by determining the value, $\tilde{\mathbf{u}}_{\boldsymbol{\beta},\boldsymbol{\theta}}$, that maximizes the unscaled conditional density. When the conditional density, $\mathcal{U}|\mathcal{Y} = \mathbf{y}_{\text{obs}}$, is multivariate Gaussian, this conditional mode will also be the conditional mean. However, for most families used in GLMMs, the mode and the mean need not coincide so use the more general term and call $\tilde{\mathbf{u}}_{\boldsymbol{\beta},\boldsymbol{\theta}}$ the *conditional mode*.

The iteratively reweighted least squares (IRLS) algorithm is an incredibly efficient method of determining the maximum likelihood estimates of the coefficients in a generalized linear model. We extend it to a *penalized iteratively reweighted least squares* (PIRLS) algorithm for determining the conditional mode, $\tilde{\mathbf{u}}_{\beta, \theta}$. This algorithm has the form

1. Given parameter values, β and θ , and starting estimates, $\mathbf{u}_0$, evaluate the linear predictor, η , the corresponding conditional mean, $\mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}$, and the conditional variance. Establish the weights as the inverse of the variance. We write these weights in the form of a diagonal weight matrix, $\mathbf{W}$, although they are stored and manipulated as a vector.
2. Solve the penalized, weighted, nonlinear least squares problem

$$\arg \min_{\mathbf{u}} \left(\left\| \mathbf{W}^{1/2} (\mathbf{y}_{\text{obs}} - \mu_{\mathcal{Y}|\mathcal{U}=\mathbf{u}}) \right\|^2 + \|\mathbf{u}\|^2 \right) \quad (36)$$

3. Update the weights, $\mathbf{W}$, and check for convergence. If not converged, go to step 2.

We use a Gauss-Newton algorithm with an orthogonality convergence criterion (Bates and Watts 1988, §2.2.3) to solve the penalized, weighted, nonlinear least squares problem in step 2. At the i th iteration we determine an increment, δ_i , as the solution to the penalized, weighted, linear least squares problem

$$\delta_i = \arg \min_{\delta} \left\| \begin{bmatrix} \mathbf{W}^{1/2} (\mathbf{y}_{\text{obs}} - \mu_i) \\ \mathbf{u}_i \end{bmatrix} - \begin{bmatrix} \mathbf{W}^{1/2} \mathbf{M}_i \mathbf{Z} \Lambda_{\theta} \\ \mathbf{I}_q \end{bmatrix} \delta \right\|^2 \quad (37)$$

where $\mathbf{u}_i$ is current value of $\mathbf{u}$, μ_i is the corresponding conditional mean of $\mathcal{Y}|\mathcal{U} = \mathbf{u}_i$ and $\mathbf{M}_i$ is the Jacobian matrix of the vector-valued inverse link, evaluated at μ_i . That is

$$\mathbf{M}_i = \left. \frac{d\mu}{d\eta'} \right|_{\eta_i}, \quad (38)$$

which will be a diagonal matrix so, as for the weights, we store and manipulate the Jacobian as a vector.

The minimizer, δ_i , of (37) satisfies

$$\mathbf{P} (\Lambda'_{\theta} \mathbf{Z}' \mathbf{M}_i \mathbf{W} \mathbf{M}_i \mathbf{Z} \Lambda_{\theta} + \mathbf{I}_q) \mathbf{P}' \delta_i = \Lambda'_{\theta} \mathbf{Z}' \mathbf{M}_i \mathbf{W} (\mathbf{y}_{\text{obs}} - \mu_i) - \mathbf{u}_i \quad (39)$$

which we solve using the sparse Cholesky factor. At convergence, the factor, $\mathbf{L}_{\beta, \theta}$, satisfies

$$\mathbf{L}_{\beta, \theta} \mathbf{L}'_{\beta, \theta} = \mathbf{P} (\Lambda'_{\theta} \mathbf{Z}' \mathbf{M} \mathbf{W} \mathbf{M} \mathbf{Z} \Lambda_{\theta} + \mathbf{I}_q) \mathbf{P}' \quad (40)$$

The integrand in the likelihood (35) is approximately a constant times the density of the $\mathcal{N}(\tilde{\mathbf{u}}, \mathbf{L} \mathbf{L}')$ distribution. The *Laplace approximation* to the deviance is

$$d(\beta, \theta | \mathbf{y}) = d_g(\mathbf{y}_{\text{obs}}, \mu(\tilde{\mathbf{u}})) + \|\tilde{\mathbf{u}}\|^2 + \log(|\mathbf{L}|^2) \quad (41)$$

where $d_g(\mathbf{y}_{\text{obs}}, \mu(\tilde{\mathbf{u}}))$ is the GLM deviance for $\mathbf{y}_{\text{obs}}$ and $\mu(\tilde{\mathbf{u}})$.

References

Bates DM, Watts DG (1988). *Nonlinear Regression Analysis and Its Applications*. Wiley, Hoboken, NJ. ISBN 0-471-81643-4.

Davis T (2006). *Direct Methods for Sparse Linear Systems*. SIAM, Philadelphia, PA.

Laird NM, Ware JH (1982). "Random-Effects Models for Longitudinal Data." *Biometrics*, **38**, 963–974.

Affiliation:

Douglas Bates
Department of Statistics, University of Wisconsin
1300 University Ave.
Madison, WI 53706, U.S.A.
E-mail: bates@stat.wisc.edu

Martin Mächler
Seminar für Statistik, HG G~16
ETH Zurich
8092 Zurich, Switzerland
E-mail: maechler@stat.math.ethz.ch

Benjamin M. Bolker
Departments of Mathematics & Statistics and Biology
McMaster University
1280 Main Street W
Hamilton, ON L8S 4K1, Canada
E-mail: bolker@mcmaster.ca