

Fitting generalized linear mixed-effects models using **lme4**

Anna Ly
McMaster University

Rune Haubo Bojesen Christensen
Copenhagen Research Centre
for Biological and Precision Psychiatry

Douglas Bates
University of Wisconsin - Madison

Martin Mächler
ETH Zurich

Benjamin M. Bolker
McMaster University

Abstract

The **lme4** R package can be used to fit generalized linear mixed models (GLMMs), which extend the class of linear mixed models (LMMs). The two main extensions provided by GLMMs are (1) allowing for the conditional distribution of the response given the random effects to be non-Gaussian (e.g. binomial, Poisson) and (2) allowing the conditional mean to be a nonlinear function of a linear combination of the fixed and random effect coefficients, via an inverse link function. The conditional mode of the random effects given the observed data, the variance-covariance matrix of the random effects, and the fixed effect parameters are determined using penalized iteratively reweighted least squares. We compute an approximation of the integral over the distributions of the conditional modes to compute the maximum likelihood estimate for a given set of parameters (by default we use the Laplace approximation or, alternatively, the more computationally expensive adaptive Gauss-Hermite quadrature). The package provides all the standard features available for GLMs in base R, including the standard set of accessor functions as well as the possibility of user-specified distributions (within the exponential dispersion family) and link functions.

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

1. Introduction

1 The **lme4** package for R can be used to fit a broad range of mixed-effects models. One major
2 advantage of **lme4** over its predecessor, **nlme**, is that it can be used to fit generalized linear
3 mixed models (GLMMs), which combine the flexibility of linear mixed models (LMMs) and
4 generalized linear models (GLMs). In a companion paper, we have described the facilities
5 in **lme4** for fitting linear mixed models (LMMs). Here we describe the facilities for fitting
6 GLMMs.

2. Generalized Linear Mixed Models

Generalized linear mixed models extend the class of linear models in two ways: allowing for both fixed-effects parameters and random effects in the linear predictor, as in linear mixed models (LMMs), and allowing for non-Gaussian distributions of the response, whose mean is an inverse link applied to the linear predictor, as in generalized linear models (GLMs). We consider several characteristics of each of these model types.

2.1. The Linear Model

A linear model can be written as $\mathcal{Y} \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I})$, where $\mathcal{Y}$ is the n -dimensional random variable representing the response, $\mathbf{X}$ is an $n \times p$ model matrix and $\boldsymbol{\beta}$ is a p -dimensional coefficient vector. The vector $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta}$ (the *linear predictor*) is mapped to the *mean vector*, $\boldsymbol{\mu} = \mathbf{g}^{-1}(\boldsymbol{\eta})$, in this case using an *identity* link and inverse link function. (The *link* function, $\mathbf{g}$, maps $\boldsymbol{\mu}$ to $\boldsymbol{\eta}$ and the *inverse link*, $\mathbf{g}^{-1}$, maps $\boldsymbol{\eta}$ to $\boldsymbol{\mu}$.) The maximum likelihood estimates (MLEs) of the coefficients, $\hat{\boldsymbol{\beta}}$, are the values that minimize the (weighted) sum of squared residuals

$$\hat{\boldsymbol{\beta}} = \arg \min_{\boldsymbol{\beta}} \sum_{i=1}^n w_i (y_i - \mu_i)^2. \quad (1)$$

with weights w_i . The weighted squared residuals, $d(y_i, \mu_i, w_i) = w_i (y_i - \mu_i)^2$, are the *unit deviances* for the Gaussian family.¹

2.2. The Generalized Linear Model

In a *generalized linear model* (GLM) the linear predictor is still $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta}$ but the distribution of the response, $\mathcal{Y}$, with mean $\boldsymbol{\mu} = \mathbf{g}^{-1}(\boldsymbol{\eta})$, can be from distribution families other than the Gaussian, specifically from the class of exponential dispersion families; the most common examples are Bernoulli, binomial, or Poisson. Individual components of the response, $\mathcal{Y}_i, i = 1, \dots, n$, are pairwise independent. The distribution family determines the *unit deviance* function, $d(y_i, \mu_i)$, and the MLEs of the coefficients are those that minimize the sum of the unit deviances

$$\hat{\boldsymbol{\beta}} = \arg \min_{\boldsymbol{\beta}} \sum_{i=1}^n d(y_i, \mu_i, w_i). \quad (2)$$

The unit deviances are the (weighted) difference between negative twice the log-likelihood at (y_i, μ_i) and negative twice the log-likelihood at (y_i, y_i) (which is often zero). For the Poisson distribution the unit deviance function is

$$d(y_i, \mu_i, w_i) = \begin{cases} 2w_i \mu_i & \text{if } y_i = 0 \\ 2w_i (y_i \log(y_i/\mu_i) - (y_i - \mu_i)) & \text{otherwise} \end{cases} \quad (3)$$

where w_i is an optional prior case weight. For simplicity we generally neglect the prior weights and write $d(y_i, \mu_i)$ below.

Minimizing the sum of the squared residuals in a linear model (i.e. finding the *least squares* estimates $\hat{\boldsymbol{\beta}}$ for an LM) can be performed as a direct (i.e. non-iterative) calculation. For a

¹In the R *family* lists of functions that encapsulate the distribution families in arguments to `glm` and `glmer`, the function that returns the unit deviances is named `dev.resids`, a source of confusion in some cases because they are not the values returned by `residuals(m, type="deviance")`.

generalized linear model minimizing the sum of unit deviances usually must be done through iteration. Fortunately, this can be done using a fast, robust algorithm, *iteratively reweighted least squares* (IRLS: McCullagh and Nelder 1989).

The link and inverse link functions are *diagonal mappings*, 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 $\boldsymbol{\mu}_{\mathcal{Y}}$. (The name “diagonal” reflects the fact that the Jacobian matrix, $\frac{d\boldsymbol{\eta}}{d\boldsymbol{\mu}}$, of such a mapping will be diagonal.) If the distribution family is one of the exponential distribution families a *canonical link* can be derived from the form of the distribution. The canonical link is the identity for the Gaussian family, the *logit* or “log-odds” function, $g(\mu) = \log(\mu/(1 - \mu))$, for the Bernoulli or binomial families, and the natural logarithm, $g(\mu) = \log(\mu)$, for the Poisson.

2.3. The Linear Mixed Model

In a linear mixed model (LMM), described in Bates, Mächler, Bolker, and Walker (2015), there are two vector-valued random variables: $\mathcal{Y}$, the n -dimensional response, and $\mathcal{B}$, the q -dimensional random effects vector. The unconditional distribution $\mathcal{B} \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_{\boldsymbol{\theta}})$ depends on a *covariance parameter* $\boldsymbol{\theta}$. The conditional distribution, $(\mathcal{Y}|\mathcal{B} = \mathbf{b}) \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{b}, \sigma^2\mathbf{I})$, has mean, $\boldsymbol{\mu}$, from the identity link applied to the linear predictor $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{b}$, where $\mathbf{Z}$ is an $n \times q$ model matrix for the random effects and $\mathbf{b}$ is a value of $\mathcal{B}$. In practice the covariance parameter, $\boldsymbol{\theta}$, determines a lower triangular *relative covariance factor*, $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}$, with $\boldsymbol{\Sigma}_{\boldsymbol{\theta}} = \sigma^2\boldsymbol{\Lambda}_{\boldsymbol{\theta}}\boldsymbol{\Lambda}_{\boldsymbol{\theta}}^{\top}$ and we work with a “spherical” random effects variable, $\mathcal{U} \sim \mathcal{N}(\mathbf{0}, \sigma^2\mathbf{I})$, such that $\mathcal{B} = \boldsymbol{\Lambda}_{\boldsymbol{\theta}}\mathcal{U}$, and for which the linear predictor is

$$\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{b} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\boldsymbol{\Lambda}_{\boldsymbol{\theta}}\mathbf{u}. \quad (4)$$

($\mathcal{U}$ is said to be “spherical” because contours of constant probability density are spheres centered at the origin.)

Most generally, $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}$ can be any lower triangular matrix, since $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}\boldsymbol{\Lambda}_{\boldsymbol{\theta}}^{\top}$ is then guaranteed to be positive semidefinite. In **lme4** version 2.0 and later, we have also implemented *structured covariance matrices*, which introduce a more elaborate mapping from $\boldsymbol{\theta}$ to $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}$ that lets users specify a range of structured forms such as diagonal, compound symmetric, or first-order autoregressive (AR1) covariance matrices. We introduce the syntax for these structures in the CBPP example (Section 5.1); see the **lme4 covariance structures vignette** for the details of how $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}$ is constructed for each structure.

The likelihood of the parameters, $\boldsymbol{\beta}$, $\boldsymbol{\theta}$ and σ^2 , given $\mathcal{Y} = \mathbf{y}$, is the marginal density of $\mathcal{Y}$ evaluated at the observed $\mathbf{y}$. Conceptually, we determine an expression for the integral of the joint density, $f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}, \mathbf{u})$, with respect to $\mathbf{u}$, then evaluate that expression at the observed $\mathbf{y}$. In practice, it is easier to reverse the order of the evaluation at the observed $\mathbf{y}$ and the integration with respect to $\mathbf{u}$.

The joint density is

$$f_{\mathcal{Y},\mathcal{U}}(\mathbf{y}, \mathbf{u}) = \frac{1}{(2\pi\sigma^2)^{(n+q)/2}} \exp\left(\frac{\|\mathbf{y} - \boldsymbol{\mu}\|^2 + \|\mathbf{u}\|^2}{-2\sigma^2}\right), \quad (5)$$

where the expression in the numerator of the integrand is the *penalized sum of the unit deviances* (PSUD) for the Gaussian family, which is the penalized sum of squared residuals (PSSR). (The sum of the unit deviances is $\|\mathbf{y} - \boldsymbol{\mu}\|^2$ and the penalty is $\|\mathbf{u}\|^2$.)

To evaluate the integral we determine the *conditional mode* of $\mathcal{U}$, given $\mathcal{Y} = \mathbf{y}$, which minimizes the PSSR with respect to $\mathbf{u}$

$$(\tilde{\mathbf{u}}|\mathcal{Y} = \mathbf{y}, \boldsymbol{\theta}, \boldsymbol{\beta}) = \arg \min_{\mathbf{u}} \left[\sum_{i=1}^n (y_i - \mu_i)^2 + \|\mathbf{u}\|^2 \right]. \quad (6)$$

The PSSR can be minimized directly (i.e. without iteration) and the minimum doesn't depend on σ^2 . When the minimum PSSR and the determinant of the Hessian of the PSSR, which can be evaluated from a Cholesky factor created in the process of solving eq. 6, are available the MLE of σ^2 can be evaluated separately, as is commonly done in linear models.

Because $\boldsymbol{\beta}$ and $\mathbf{u}$ both occur as coefficients in the linear predictor, the conditional estimate, $\hat{\boldsymbol{\beta}}|\boldsymbol{\theta}$, can be determined along with $\tilde{\mathbf{u}}$ by extending the penalized least squares problem, (eq. 6), to minimize with respect to both $\mathbf{u}$ and $\boldsymbol{\beta}$. Thus, given a value of $\boldsymbol{\theta}$ and the solution to this extended PSSR problem, the profiled log-likelihood, which is a function of $\boldsymbol{\theta}$ only, can be evaluated. General nonlinear optimizers applied to the profiled log-likelihood as a function of $\boldsymbol{\theta}$ are used to determine the MLEs for all the parameters.

The point of using the profiled log-likelihood instead of the original log-likelihood is that the dimension of the general nonlinear optimization problem is much smaller in the profiled problem.

2.4. The Generalized Linear Mixed Model

As stated earlier, a GLMM combines the distribution family and link function, from a GLM, with random effects in the linear predictor, as in an LMM. The unconditional distribution of the random effects, $\mathcal{B} \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_{\boldsymbol{\theta}})$, and the definitions of the spherical random effects, $\mathcal{U}$, such that $\mathcal{B} = \boldsymbol{\Lambda}_{\boldsymbol{\theta}} \mathcal{U}$ where $\boldsymbol{\Lambda}_{\boldsymbol{\theta}} \boldsymbol{\Lambda}_{\boldsymbol{\theta}}^{\top} = \boldsymbol{\Sigma}_{\boldsymbol{\theta}}$ are as in the LMM except that the scale parameter, σ , is not present in the most common GLMMs (binomial and Poisson families). The mean of the conditional distribution, $\mathcal{Y}|\mathcal{U} = \mathbf{u}$, is the inverse link applied to the linear predictor $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\boldsymbol{\Lambda}_{\boldsymbol{\theta}}\mathbf{u}$.

The conditional mode, which minimizes the *penalized sum of unit deviances* (PSUD),

$$(\tilde{\mathbf{u}}|\mathcal{Y} = \mathbf{y}, \boldsymbol{\theta}, \boldsymbol{\beta}) = \arg \min_{\mathbf{u}} \left[\sum_{i=1}^n d(y_i, \mu_i) + \|\mathbf{u}\|^2 \right], \quad (7)$$

is determined by penalized iteratively reweighted least squares (PIRLS).

In the case of an LMM, minimizing the PSSR provided the value of the integral defining the likelihood of the parameters, given the observed response, $\mathcal{Y} = \mathbf{y}$. For a GLMM, minimizing the PSUD provides a local quadratic approximation to the logarithm of the joint density, $f_{\mathcal{Y}, \mathcal{U}}(\mathbf{y}, \mathbf{u})$. Approximating the integral with respect to $\mathbf{u}$ of the joint density (evaluated at the observed $\mathbf{y}$) by the integral of this local quadratic approximation is *Laplace's method*.

For simple models involving scalar random effects from a single random-effects term the multivariate integral with respect to $\mathbf{u}$ can be factored into a product of scalar integrals, which can be evaluated by Gauss-Hermite rules.

To summarize:

1. Like a GLM, the definition of a GLMM incorporates a distribution family and a link function. The distribution family determines the unit deviances, $d(y_i, \mu_i)$, $i = 1, \dots, n$.

- 111 Distributions in the exponential distribution family have a canonical link which is the
 112 default link in `glm` and in `glmer` for those families.
- 113 2. As in an LMM the linear predictor, $\boldsymbol{\eta} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{b}$, incorporates fixed-effects parameters,
 114 $\boldsymbol{\beta}$, and random effects, $\mathbf{b} \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_b)$.
- 115 3. The parameters $\boldsymbol{\theta}$ determine a triangular relative covariance factor, $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}$, such that $\boldsymbol{\Sigma}_{\boldsymbol{\theta}} =$
 116 $\boldsymbol{\Lambda}_{\boldsymbol{\theta}}\boldsymbol{\Lambda}_{\boldsymbol{\theta}}^{\top}$. (A scale parameter, σ , is incorporated for families like the Gaussian that use
 117 one.)
- 118 4. For given values of $\boldsymbol{\beta}$ and $\boldsymbol{\theta}$, the penalized sum of unit deviances (PSUD), $\sum_{i=1}^n d(y_i, \mu_i) +$
 119 $\|\mathbf{u}\|^2$, is minimized with respect to $\mathbf{u}$ using PIRLS.
- 120 5. A local quadratic approximation to the PSUD at the conditional mode provides Laplace's
 121 approximation to the log-likelihood, which is minimized with respect to $\boldsymbol{\beta}$ and $\boldsymbol{\theta}$.
- 122 6. In certain simple models adaptive Gauss-Hermite quadrature (aGHQ) can be used to
 123 approximate the log-likelihood more accurately, at the expense of more evaluations of
 124 the PSUD for each candidate combination of $\boldsymbol{\beta}$ and $\boldsymbol{\theta}$.
- 125 7. If the PSUD is minimized with respect to $\mathbf{u}$ and $\boldsymbol{\beta}$ simultaneously then the approx-
 126 imate profiled log-likelihood becomes a function of $\boldsymbol{\theta}$ only, as discussed below. This
 127 approximation can save computation time for large problems.

128 2.5. Determining the conditional mode

129 The iteratively reweighted least squares (IRLS) algorithm is an efficient method of determining
 130 the maximum likelihood estimates of the coefficients in a GLM. As the name IRLS implies
 131 this is an iterative algorithm where the $(k+1)$ st coefficient vector, $\boldsymbol{\beta}^{(k+1)}$, is determined from
 132 $\boldsymbol{\beta}^{(k)}$, and the process continues until convergence. At the k th iteration, the linear predictor is
 133 $\boldsymbol{\eta}^{(k)} = \mathbf{X}\boldsymbol{\beta}^{(k)}$ and the mean response vector is $\boldsymbol{\mu}^{(k)} = \mathbf{g}^{-1}(\boldsymbol{\eta}^{(k)})$. The residual on the response
 134 scale, $\mathbf{r}^{(k)} = \mathbf{y} - \boldsymbol{\mu}^{(k)}$, is mapped to a *working residual* on the linear predictor scale via a linear
 135 approximation to the link function. Because the link and inverse link are diagonal maps we
 136 can write the working residual component-wise as

$$\tilde{r}_i^{(k)} = (y_i - \mu_i^{(k)})g'(\mu_i^{(k)}) \quad i = 1, \dots, n, \quad (8)$$

137 where $g'(\mu_i^{(k)})$ denotes the derivative of the link function at $\mu_i^{(k)}$.

138 Because the variances of the working residuals are not constant we must evaluate these vari-
 139 ances and use weighted least squares, with the weights inversely proportional to the variances,
 140 to determine an increment $\boldsymbol{\delta}^{(k)}$ such that $\boldsymbol{\beta}^{(k+1)} = \boldsymbol{\beta}^{(k)} + \boldsymbol{\delta}^{(k)}$. (The *working weights* based
 141 on the variances are combined with the prior weights w_i , if any have been specified.) An
 142 alternative approach is to define a *working response*

$$\tilde{y}_i^{(k)} = (\eta_i^{(k)} - o_i) + \tilde{r}_i^{(k)}, \quad i = 1, \dots, n, \quad (9)$$

143 where o_i is an optional *offset* to be applied on the linear predictor scale.

144 IRLS is usually implemented by solving a weighted least squares problem on the working
 145 residual for the increment, $\boldsymbol{\delta}^{(k)}$, at each iteration. For the *penalized iteratively reweighted*

least squares (PIRLS) algorithm to determine $\mathbf{u}^{(k+1)}$ from $\mathbf{u}^{(k)}$ it is easiest to use the working response because the penalty is defined as $\|\mathbf{u}^{(k+1)}\|^2$ so the penalized weighted least squares calculation should be in terms of $\mathbf{u}$.

In more detail, the PIRLS algorithm has the form

1. Given parameter values, $\boldsymbol{\beta}$ and $\boldsymbol{\theta}$, incorporate the fixed-effects contribution to the linear predictor, $\mathbf{X}\boldsymbol{\beta}$, as an *offset*, $\mathbf{o}$, to be used in evaluating the working response (eq. 9). Set the initial values of $\mathbf{u}^{(0)} = \mathbf{0}$ so that the PIRLS iterations always start from the same value of $\mathbf{u}$, providing for a stable result that depends only on $\boldsymbol{\beta}$ and $\boldsymbol{\theta}$. Evaluate the lower triangular $\mathbf{\Lambda}_{\boldsymbol{\theta}}$, the linear predictor

$$\boldsymbol{\eta}^{(0)} = \mathbf{Z}\mathbf{\Lambda}_{\boldsymbol{\theta}}\mathbf{u}^{(0)} + \mathbf{o}, \quad (10)$$

the conditional mean of $\mathcal{Y}|\mathcal{U} = \mathbf{u}^{(0)}$,

$$\boldsymbol{\mu}^{(0)} = \mathbf{g}^{-1}(\boldsymbol{\eta}^{(0)}), \quad (11)$$

the (diagonal) Jacobian matrix,

$$\mathbf{J}^{(0)} = \frac{d\boldsymbol{\mu}}{d\boldsymbol{\eta}^\top} = \mathbf{g}^{-1'} \quad (12)$$

and the conditional variance

$$\mathbf{V}^{(0)} = \text{Var}(\mathcal{Y}|\boldsymbol{\mu} = \boldsymbol{\mu}^{(0)}) \quad (13)$$

which is also a diagonal matrix.

2. Establish the weights as the inverse of the combined variance contributions from $\boldsymbol{\eta}$:

$$\mathbf{W}^{(0)} = \left(\mathbf{J}^{(0)\top} \mathbf{V}^{(0)} \mathbf{J}^{(0)} \right)^{-1} \quad (14)$$

(we would also multiply this quantity by a diagonal matrix of prior weights, if any were specified). Eq. 14 is written as a product and an inverse of matrices but all the matrices involved are diagonal and the evaluation of eq. 14 is performed as vector operations.

3. Evaluate $\mathbf{u}^{(1)}$ as the solution to the penalized, weighted least squares problem

$$\left(\mathbf{\Lambda}^\top \mathbf{Z}^\top \mathbf{W}^{(0)} \mathbf{Z} \mathbf{\Lambda} + \mathbf{I} \right) \mathbf{u}^{(1)} = \mathbf{\Lambda}^\top \mathbf{Z}^\top \mathbf{W}^{(0)} \tilde{\mathbf{y}}^{(0)} \quad (15)$$

using the sparse Cholesky factor, $\mathbf{L}_{\boldsymbol{\beta}, \boldsymbol{\theta}}^{(0)}$, defined by

$$\mathbf{L}_{\boldsymbol{\beta}, \boldsymbol{\theta}}^{(0)} \mathbf{L}_{\boldsymbol{\beta}, \boldsymbol{\theta}}^{(0)\top} = \mathbf{P} \left(\mathbf{\Lambda}^\top \mathbf{Z}^\top \mathbf{W}^{(0)} \mathbf{Z} \mathbf{\Lambda} + \mathbf{I} \right) \mathbf{P}^\top, \quad (16)$$

where $\mathbf{P}$ is the matrix representation of a *fill-reducing permutation* determined from the structure of $\mathbf{Z}^\top \mathbf{Z}$. (In simple GLMM problems $\mathbf{P}$ can be an identity matrix but for general sparse Cholesky evaluations determining and employing a fill-reducing permutation can be important in reducing memory usage and compute time.)

4. Evaluate PSUD, the penalized sum of unit deviances, at $\mathbf{u}^{(1)}$. If the PSUD at $\mathbf{u}^{(1)}$ exceeds that at $\mathbf{u}^{(0)}$ use *step-halving* which, in this case, means replacing $\mathbf{u}^{(1)}$ by the average of $\mathbf{u}^{(0)}$ and $\mathbf{u}^{(1)}$, and trying again. We set a maximum, usually 10 or fewer, on the number of step-halvings that are allowed before declaring the algorithm to have failed.
5. If the iteration succeeds in reducing the PSUD, check for convergence by comparing the weights, $\mathbf{W}^{(i+1)}$, to $\mathbf{W}^{(i)}$. If the relative differences in the weights are small the algorithm is declared to have converged. Otherwise repeat the previous steps using evaluations at $\mathbf{u}^{(i)}$ to produce $\mathbf{u}^{(i+1)}$.

2.6. Evaluating the likelihood for GLMMs using the Laplace approximation

Evaluating the likelihood for generalized linear mixed models requires approximating an intractable integral over the random effects distribution. The `glmer` function offers several approximations, controlled by the `nAGQ` argument. The default value of `nAGQ=1` specifies the *Laplace approximation* (Madsen and Thyregod 2011).

A second-order Taylor series approximation to $-2 \log[f_{\mathcal{Y}, \mathcal{U}}(\mathbf{y}, \mathbf{u})]$ based at $\tilde{\mathbf{u}}$ provides an approximation of the unscaled conditional density as a multiple of the density for the multivariate Gaussian $\mathcal{N}(\tilde{\mathbf{u}}, \mathbf{L}\mathbf{L}^\top)$. The change of variable

$$\mathbf{u} = \tilde{\mathbf{u}} + \mathbf{L}\mathbf{z} \quad (17)$$

provides

$$\begin{aligned} L(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}) &= \int_{\mathbb{R}^q} f_{\mathcal{Y}, \mathcal{U}}(\mathbf{y}, \mathbf{u}) d\mathbf{u} \\ &\approx \tilde{f} |\mathbf{L}| \int_{\mathbb{R}^q} e^{-\|\mathbf{z}\|^2/2} (2\pi)^{-q/2} d\mathbf{z} \\ &= \tilde{f} |\mathbf{L}| \end{aligned} \quad (18)$$

(where $\tilde{f}$ is the penalized deviance evaluated at the conditional modes) or, on the deviance scale,

$$-2\ell(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}) \approx \sum_{i=1}^n d(y_i, \mu_i) + \|\tilde{\mathbf{u}}\|^2 + \log(|\mathbf{L}|^2) + \frac{q}{2} \log(2\pi) \quad (19)$$

The Laplace approximation normally conditions on both the fixed effects $\boldsymbol{\beta}$ and the variance-covariance parameters $\boldsymbol{\theta}$. A further approximation, which is denoted in `glmer` by `nAGQ=0`, profiles out the fixed effects by minimizing with respect to $\boldsymbol{\beta}$ and $\mathbf{u}$ simultaneously in the PIRLS algorithm (i.e., replacing $\mathbf{u}$ with $(\boldsymbol{\beta} | \mathbf{u})$ and $\mathbf{Z}\boldsymbol{\Lambda}$ with $(\mathbf{X} | \mathbf{Z}\boldsymbol{\Lambda})$ throughout, except with the penalty term still applying only to $\mathbf{u}$). This approximation is exact when (1) $\partial(\log L)/\partial\boldsymbol{\beta}$ is a linear function of the conditional modes $\mathbf{u}$ and (2) when the conditional mode is equal to the conditional mean (typically, although not necessarily, implying a symmetric conditional distribution). Both assumptions hold for linear mixed models (although a Laplace approximation is not necessary there), consistent with Bates *et al.* (2015) showing that the fixed effects can be profiled out of the log-likelihood for LMMs. The Julia `MixedModels.jl` package offers the same approximation as the `fast` argument to the `pirls!` function (<https://>

juliastats.org/MixedModels.jl/stable/optimization/); Template Model Builder (Kristensen, Nielsen, Berg, Skaug, and Bell 2016), and downstream packages such as glmmTMB, provide this function via a `profile` argument.

By default, `glmer` uses a two-stage optimization procedure (described below) with `nAGQ=0` in the first stage; users can also specify `nAGQ=0` for faster, approximate model fits.

Decomposing the deviance for simple models

A common special case of mixed models is those where scalar (typically intercept) random effects are associated with levels of a single grouping factor, $\mathbf{h}$. In this case the dimension, q , of the random effects is the number of levels of $\mathbf{h}$ — i.e. there is exactly one random effect associated with each level of $\mathbf{h}$. We will write the vector of variance-covariance parameters, which is one-dimensional, as a scalar, θ . The matrix $\mathbf{\Lambda}_\theta$ is a multiple of the identity, $\theta \mathbf{I}_q$, and $\mathbf{Z}$ is the $n \times q$ matrix of indicators of the levels of $\mathbf{f}$. The permutation matrix, $\mathbf{P}$, can be set to the identity and $\mathbf{L}$ is diagonal, although not necessarily homogeneous (i.e., a scalar multiple of the identity matrix).

Because each element of $\boldsymbol{\mu}$ depends on only one element of $\mathbf{u}$ and the elements of $\mathcal{Y}$ are conditionally independent, given $\mathcal{U} = \mathbf{u}$, the conditional densities of the $u_j, j = 1, \dots, q$ given $\mathcal{Y} = \mathbf{y}$ are independent. We partition the indices $1, \dots, n$ as $\mathbb{I}_j, j = 1, \dots, q$ according to the levels of $\mathbf{h}$. That is, the index i is in $\mathbb{I}_j$ if $h_i = j$. This partitioning also applies to the deviance residuals in that the i th deviance residual depends only on u_j when $i \in \mathbb{I}_j$.

Writing the univariate conditional densities as

$$f_j(\mathbf{y}, u_j) = \exp \left(-\frac{\sum_{i \in \mathbb{I}_j} d(y_i, u_j) + u_j^2}{2} \right) (2\pi)^{-1/2} \quad (20)$$

we have

$$f_{\mathcal{Y}, \mathcal{U}}(\mathbf{y}, \mathbf{u}) = \prod_{j=1}^q f_j(\mathbf{y}, u_j) \quad (21)$$

and

$$L(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}) = \prod_{j=1}^q \int_{\mathbb{R}} f_j(\mathbf{y}, u) du \quad (22)$$

We consider this special case both because it occurs frequently and because, for some software, it is the only type of GLMM that can be fit. Also, in this particular case we can graphically assess the quality of the Laplace approximation by comparing the actual integrand to its approximation.

Consider the `cbpp` data on contagious bovine pleuropneumonia (CBPP) incidence according to season and herd, available in the `lme4` package (see 5.1 for more details), and the model

```
> print(m1 <- glmer(cbind(incidence, size-incidence) ~ period + (1|herd),
```

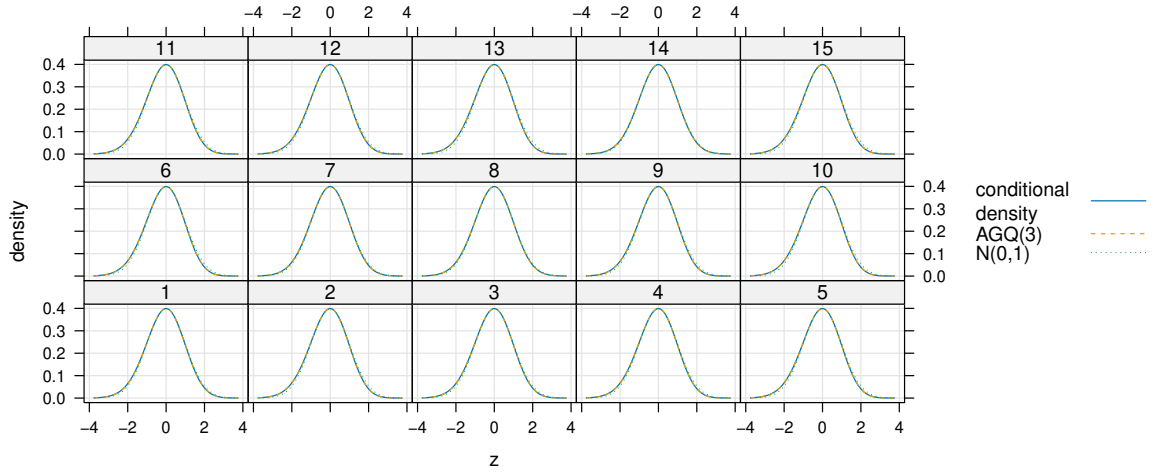



Figure 1: Comparison of univariate integrands (solid, blue line); 3-point Gauss-Hermite quadrature (dashed, yellow); and standard normal density function (green, dotted) for the CBPP model. For this model, the three approximations are nearly indistinguishable.

```
+      cbpp, binomial), corr=FALSE)

Generalized linear mixed model fit by maximum likelihood (Laplace
Approximation) [glmerMod]
Family: binomial ( logit )
Formula: cbind(incidence, size - incidence) ~ period + (1 | herd)
Data: cbpp
      AIC      BIC    logLik -2*log(L)  df.resid
 194.053   204.180   -92.027   184.053      51
Random effects:
Groups Name      Std.Dev.
herd  (Intercept) 0.642
Number of obs: 56, groups:  herd, 15
Fixed Effects:
(Intercept)    period2    period3    period4
    -1.398      -0.992     -1.128     -1.580
```

229 This model has been fit by minimizing the Laplace approximation to the deviance. We
 230 can assess the quality of this approximation by evaluating the unscaled conditional density
 231 at $u_j(z) = \tilde{u}_j + z/\mathbf{L}_{j,j}$ and comparing the ratio, $f_j(\mathbf{y}, u)/(\tilde{f}_j\sqrt{2\pi})$, to the standard normal
 232 density, $\phi(z) = e^{-z^2/2}/\sqrt{2\pi}$, as shown in Figure 1.

233 As Figure 1 shows, the univariate integrands are very close to the standard normal density,
 234 indicating that the Laplace approximation to the deviance is a good approximation in this

235 case.

3. Adaptive Gauss-Hermite quadrature for GLMMs

236 When the first integral in (18) can be expressed as a product of low-dimensional integrals,
 237 we can use Gauss-Hermite quadrature to provide a closer approximation to the integral.
 238 Univariate Gauss-Hermite quadrature evaluates the integral of a function that is multiplied
 239 by a “kernel” where the kernel is a multiple of e^{-z^2} or $e^{-z^2/2}$. For statisticians the natural
 240 candidate is the standard normal density, $\phi(z) = e^{-z^2/2}/\sqrt{(2\pi)}$. A k th-order Gauss-Hermite
 241 formula provides knots, $z_i, i = 1, \dots, k$, and weights, $w_i, i = 1, \dots, k$, such that

$$\int_{\mathbb{R}} t(z)\phi(z) dz \approx \sum_{i=1}^k w_i t(z_i) \quad (23)$$

242 The function `GHrule` in **lme4** (based on code in the **SparseGrid** package) provides knots and
 243 weights relative to the standard normal kernel for orders k from 1 to 100. For example,

```
> GHrule(5)

      z      w    ldnorm
[1,] -2.8570 0.011257 -5.00008
[2,] -1.3556 0.222076 -1.83780
[3,]  0.0000 0.533333 -0.91894
[4,]  1.3556 0.222076 -1.83780
[5,]  2.8570 0.011257 -5.00008
```

244 where **z** is the vector of knots, **w** is the vector of weights, and **ldnorm** is the log-density of the
 245 standard normal distribution at z .

246 The choice of the value of k depends on the behavior of the function $t(z)$. If $t(z)$ is a polynomial
 247 of degree $2k - 1$ then the Gauss-Hermite formula for orders k or greater provides an exact
 248 answer. The fact that we want $t(z)$ to behave like a low-order polynomial is often neglected
 249 in the formulation of a Gauss-Hermite approximation to a quadrature. The quadrature knots
 250 on the u scale are chosen as

$$u_{i,j}(z) = \tilde{u}_j + z_i/\mathbf{L}_{j,j}, \quad i = 1, \dots, k; \quad j = 1, \dots, q \quad (24)$$

251 exactly so that the function $t(z)$ should behave like a low-order polynomial over the region of
 252 interest, which is to say the region where quadrature knots with large weights are located. The
 253 term “adaptive Gauss-Hermite quadrature” reflects the fact that the approximating Gaussian
 254 density is scaled and shifted to provide a second order approximation to the logarithm of the
 255 unscaled conditional density.

256 Figure 2 shows $t(z)$ for each of the unidimensional integrals in the likelihood for the model **m1**
 257 at the parameter estimates. While this view shows the deficiency of the Laplace approximation
 258 (deviation from a horizontal reference line at $z = 1$) clearly, The AGQ(3) approximation fits
 259 extremely well over the range shown. The tails of the polynomials implied by the AGQ
 260 approximation can fluctuate widely, but these fluctuations are suppressed by multiplying by

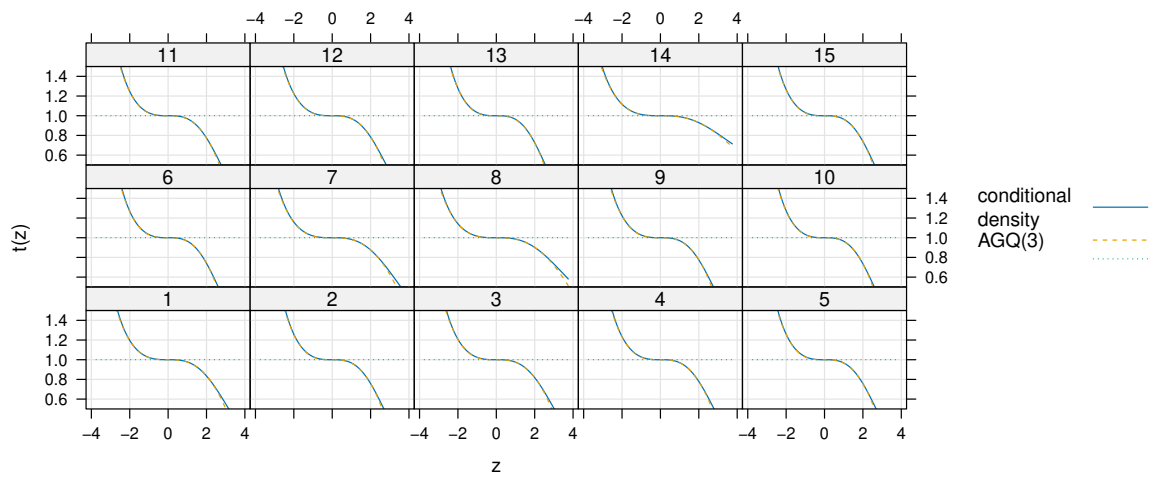


Figure 2: The function $t(z)$, which is the ratio of the normalized unscaled conditional density to the standard normal density, for each of the univariate integrals in the evaluation of the deviance for model `m1` (blue, solid line). As in Figure 1, the dashed yellow line shows the approximation for 3-point adaptive Gauss-Hermite quadrature. The horizontal green reference line ($z = 1$) gives the reference value that would apply for the Laplace approximation, which assumes the conditional density is exactly equal to the standard normal. The y -axis limits are truncated to (0.5, 1.5).

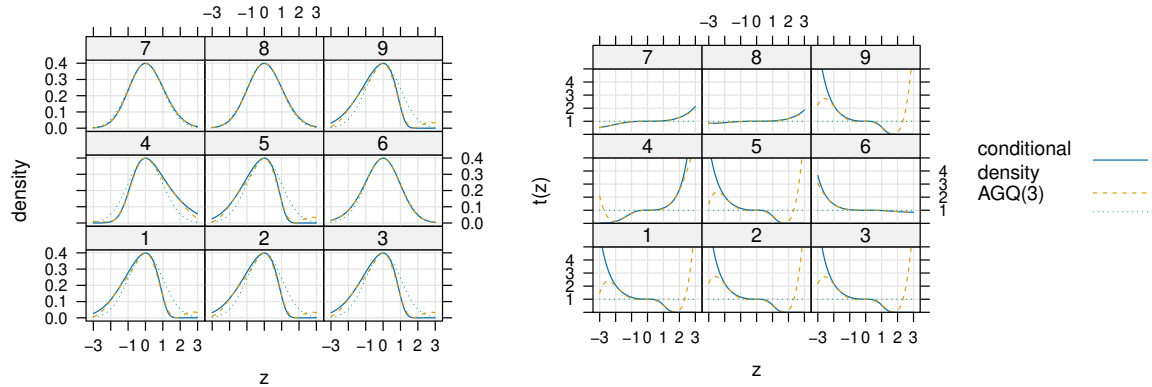


Figure 3: Normalized unscaled conditional density (left) and ratio of density to the standard normal density (right) for a random sample of 9 patients from the toenail onychomycosis data set. The y -axis limits in the right panel ($t(z)$) are truncated to $(0, 5)$.

the thin tails of the Gaussian distribution (so that in Figure 1 the deviations of the conditional density from the standard Normal are barely visible).

The CBPP data set is a relatively well-behaved data set, where Laplace approximation works well. In contrast, a widely used data set on toenail onychomycosis (De Backer, De Vroey, Lesaffre, Scheys, and De Keyser 1998), which has a very low effective sample size per cluster — an average of about 6.5 binary observations (“moderate or severe” vs “none or mild” disease) per patient — represents an example where Gauss-Hermite quadrature is necessary for reliable results. In this case the conditional densities depart much more clearly from the standard normal (Figure 3; note the scale of the density ratios goes from 0 to 10, in contrast the maximum of 4 in Figure 2). Stringer, Bilodeau, and Tang (2026) and Stringer (2024) further explore the limitations of Laplace approximation.

To use adaptive Gauss-Hermite quadrature for model fitting in `glmer` models, users would set the argument `nAGQ`, the number of quadrature points (i.e. k in eq. 23), to a value greater than 1.

Increasing the number of nodes generally improves the accuracy of the likelihood approximation at the expense of computation time — although rounding errors may accumulate when using large numbers of quadrature points. A reasonable strategy for choosing the number of nodes is to re-fit the model with increasing values of `nAGQ` until the parameter estimates

stabilize (Tuerlinckx, Rijmen, Verbeke, and De Boeck 2006); a “warm-start” procedure using the parameter estimates from the previous fit as starting values (and using `control = glmerControl(nAGQ0initStep = FALSE)` to disable the `nAGQ=0` preliminary fit) will speed up this procedure.

At present, AGQ is only available for models with a single scalar random effect. (The GLMMadaptive package implements AGQ for vector-valued random effect models, although it is still restricted to models with a single random effect; Rabe-Hesketh, Skrondal, and Pickles (2005) describe methods for AGQ for models with nested random effects.)

4. Model fitting

Once we can calculate the deviance by PIRLS for specified values of θ and β (or only θ if profiling out the fixed effects via `nAGQ=0`), we then estimate the parameters by nonlinear optimization. This procedure largely follows the description in Bates *et al.* (2015), using derivative-free optimizers with box constraints to prevent non-positive-(semi)definite covariance matrices. Specifically, the elements of θ corresponding to the diagonal of Λ_θ are currently constrained to be non-negative.

The only difference is that by default `glmer` uses a two-step fitting procedure, using `nAGQ=0` at the first stage to get preliminary estimates which are then used as starting points for a second optimization with Laplace approximation or Gauss-Hermite quadrature as specified by the user. Different nonlinear optimizers can be used at each stage: the current default, as specified in `glmerControl`, is to use Powell’s BOBYQA (Powell 2009) followed by a box-constrained variant of the Nelder-Mead simplex algorithm. For faster, approximate fitting, the second stage can be omitted; in the rare cases where the initial `nAGQ=0` fit gives poor results, the first stage can be skipped via `glmerControl(nAGQ0initStep = FALSE)`.

However, because the profiled log-likelihood is an even function of the diagonal elements of θ and $\Sigma_\theta = \Lambda_\theta \Lambda_\theta^\top$ depends on them only through their squares, positive and negative values yield identical likelihoods. This symmetry means that constrained optimization is not strictly necessary — an unconstrained optimizer will converge to a correct solution, approaching zero from either side in the boundary case of a singular random-effects covariance matrix. (Once a solution is found, we can map it to a unique solution where the diagonal elements are all non-negative.) A similar approach to removing constraints could work for structured covariance matrices where correlation parameters are constrained to $(-1, 1)$, e.g. by parameterizing the model in terms of a phase parameter p where $\rho = \sin(p)$ (and then mapping p to $(0, 2\pi)$). Removing these constraint would permit the use of a broader class of unconstrained optimizers, though this feature has not yet been incorporated into `lme4`.

5. Examples

5.1. CBPP

The `?cbpp` help page describes the CBPP data set (Lesnoff *et al.* 2004) as follows:

Contagious bovine pleuropneumonia (CBPP) is a major disease of cattle in Africa, caused by a mycoplasma. This dataset describes the serological incidence of CBPP

in zebu cattle during a follow-up survey implemented in 15 commercial herds located in the Boji district of Ethiopia. The goal of the survey was to study the within-herd spread of CBPP in newly infected herds. Blood samples were quarterly collected from all animals of these herds to determine their CBPP status. These data were used to compute the serological incidence of CBPP (new cases occurring during a given time period). Some data are missing (lost to follow-up).

Lesnoff *et al.* (2004) estimated the effects of different treatments using (1) ordinary logistic regression incorporating a variance-inflation factor, also known as a quasi-binomial model (“logistic regression” is sometimes used specifically to describe analyses of Bernoulli responses, but in this case there are multiple trials per observation [cows that could become seropositive], and so a dispersion or scale parameter can be estimated); (2) a GLMM implemented in `lme4`; and a (3) Markov chain Monte Carlo algorithm (Zeger and Karim 1991), which as they state allows for a non-parametric rather than a Normal model for the random effects. The authors did not find any significant effects of treatment, ascribing the null results to “a lack of power in the statistical analyses or to a quality problem for the medications used (and more generally, for health-care delivery in the Boji district).”

(Note that Table 1 of Lesnoff *et al.* (2004) contains a known typographical error for herd 6. Consequently, results obtained using the `cbpp` data set may not exactly reproduce some of the findings reported in that paper.)

The `lme4` package includes two variants of the `cbpp` data set. The second variant, `cbpp2`, contains corrected values corresponding to Table 1 of Lesnoff *et al.* (2004). Although the precise provenance of these data sets is unclear, the `cbpp` data set matches the version held by the corresponding author of Lesnoff *et al.* (2004).

We model proportional incidence as a binomial response depending on the additive fixed effects of period, treatment and average herd size; to account for repeated measures we fit a model with a random effect of herd. In practice we might also be interested in a period by treatment interaction, but we neglect that term here.

As in `glm`, we can specify a binomial response as a proportion and use the `weights` argument to specify the sample size, instead of the more typical two-column `cbind(successes, failures)` format:

```
> gml <- glmer(incidence/size ~ period + treatment + avg_size + (1 | herd),
+             family = binomial,
+             data = cbpp2, weights = size)
```

It is also worth considering adding an observation-level random effect to the model (Harrison 2015), which we can do by creating a new factor based on observation number and using `update()` on the previous model (we switch the optimizer to BOBYQA for both phases for one of the updates):

```
> cbpp2 <- transform(cbpp2, obs=factor(seq(nrow(cbpp2))))
```

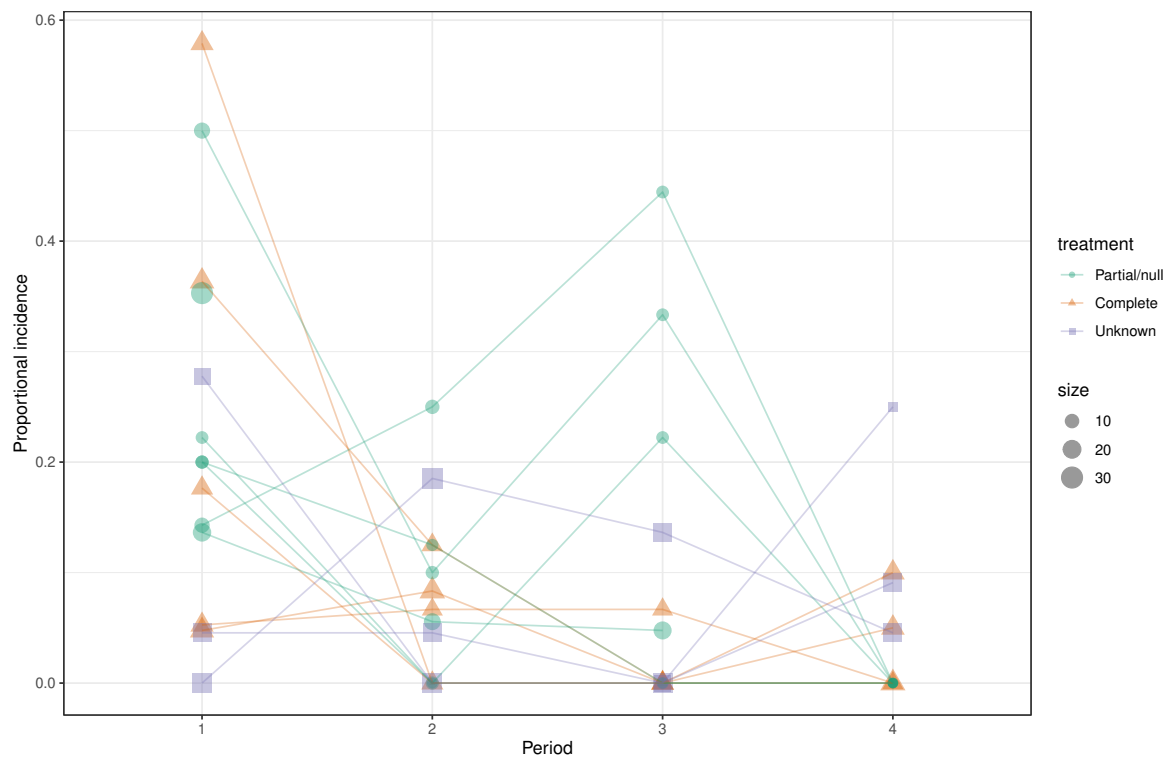


Figure 4: Incidence (proportion of cows becoming seropositive per observation period) vs. period. Colours show treatment category for each herd; point sizes reflect the number of seronegative cows at the start of each period. Lines connect the sets of observations from each herd.

```

> bob_opt <- glmerControl(optimizer = "bobyqa")
> ## herd and observation-level REs
> gm2 <- update(gm1, .~.(1|obs), control = bob_opt)
> ## observation-level REs only
> gm3 <- update(gm1, .~.(1|herd)+(1|obs))

```

350 *Model summary*

351 The first part of the summary reiterates the family and link function used, the model formula,
 352 and gives various summary statistics (log-likelihood etc.), as well as quantiles of the scaled
 353 (Pearson) residuals:

```

Generalized linear mixed model fit by maximum likelihood (Laplace
  Approximation) [glmerMod]
Family: binomial ( logit )
Formula: incidence/size ~ period + treatment + avg_size + (1 | herd)
Data: cbpp2
Weights: size

      AIC      BIC    logLik -2*log(L)  df.resid
    197.8    214.0    -90.9    181.8      48

Scaled residuals:
   Min     1Q  Median     3Q      Max
-2.231 -0.797 -0.373  0.469  2.756

```

354 These quantities are also accessible via standard accessors (`AIC()`, `BIC()`, `logLik()`).

355 The next chunk of `summary()` describes the random effects and the number of levels associated
 356 with each grouping factor (the latter is useful for checking that random-effects formulae have
 357 been specified correctly):

```

Random effects:
Groups Name      Variance Std.Dev.
herd   (Intercept) 0.312    0.558
Number of obs: 56, groups:  herd, 15

```

358 This information is also accessible via `VarCorr()`, which returns a list of variance-covariance
 359 matrices (the `print` method for `VarCorr` objects allows control of whether the variance, or
 360 standard deviation, or both, are printed).

361 Next come the estimates of the fixed effects, along with Wald estimates of the standard error,
 362 *Z* statistic, and *p*-value:

```

Fixed effects:

```


	Estimate	Std. Error	z value	Pr(> z)
(Intercept)	-1.00562	0.70842	-1.42	0.15574
period2	-0.98628	0.30338	-3.25	0.00115
period3	-1.12515	0.32314	-3.48	0.00050
period4	-1.56110	0.42263	-3.69	0.00022
treatmentComplete	-0.37623	0.50300	-0.75	0.45448
treatmentUnknown	-0.68325	0.64518	-1.06	0.28960
avg_size	-0.00613	0.04561	-0.13	0.89300

363 One can use `coef(summary())` to retrieve this information, and optionally format it with
 364 `printCoefmat()`.

365 The last component of `summary()` gives the estimated correlations among the fixed-effect
 366 parameters, which can be useful for assessing multicollinearity (it can also be overwhelming:
 367 it is suppressed by default for models with more than 20 fixed-effect parameters, and can also
 368 be suppressed by using `print(summary(.), correlation=FALSE)`).

```
Correlation of Fixed Effects:
              (Intr) perid2 perid3 perid4 trtmnC trtmnU
period2      -0.135
period3      -0.130  0.278
period4      -0.086  0.210  0.195
trtmnCmplt   0.289 -0.017 -0.021 -0.059
trtmnUnknw   0.431 -0.053 -0.045 -0.043  0.588
avg_size     -0.910  0.026  0.028  0.020 -0.547 -0.649
```

369 *Diagnostics*

370 A range of graphical diagnostic tools is available for `merMod` objects. The plot methods in
 371 the `lme4` package are inspired by those in the `nlme` package, using `lattice` plots to provide
 372 a reasonable blend of convenience and flexibility.

373 `merMod` objects are also compatible with the `performance` package and the `DHARMA` package,
 374 both commonly used for model checking.

375 The following code produces a standard range of diagnostic plots (Figure 5), similar to the
 376 ones in base R's `plot.lm` method. These diagnostics will generally be useful for models
 377 where the conditional density is approximately normal (but heteroscedastic) — e.g., Poisson
 378 responses with large mean or binomial responses with large numbers of successes — and less
 379 so otherwise, e.g. for binary responses.

```
> ## basic residual plot
```

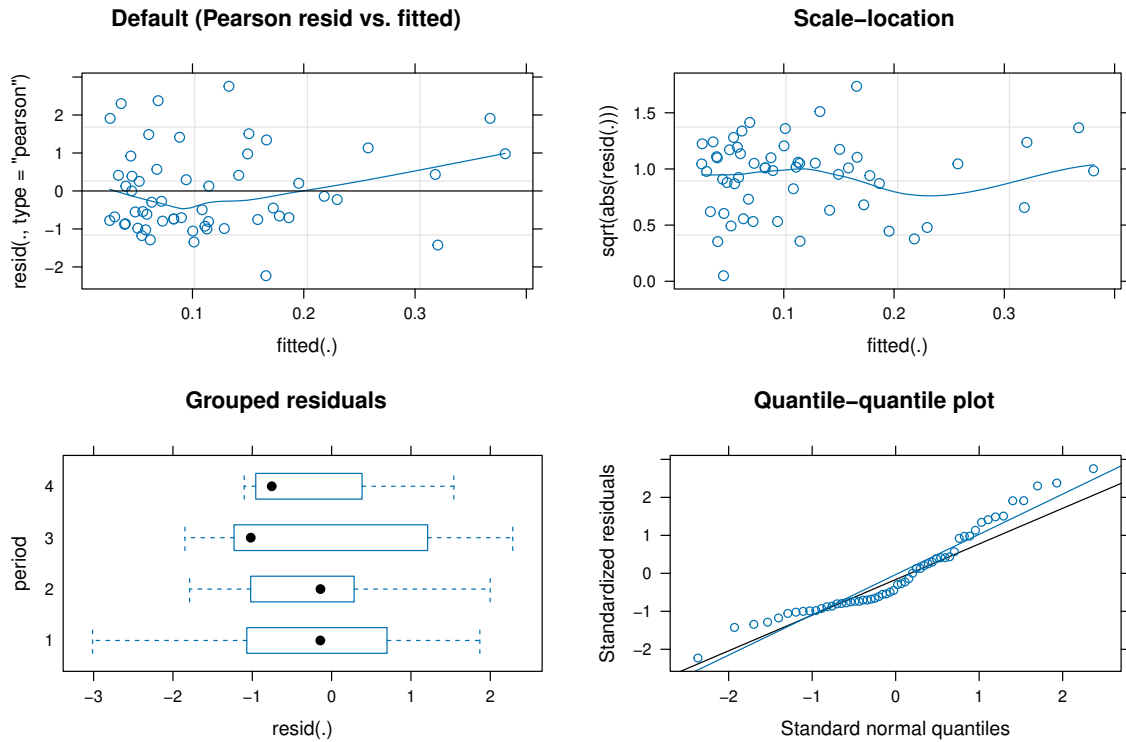


Figure 5: Graphical diagnostics (using package built-in functions)

```

> plot(gm1)
> ## scale-location plot
> plot(gm1, sqrt(abs(resid(.)))~fitted(., type=c("p", "smooth")))
> ## boxplot of residuals grouped by a categorical predictor
> plot(gm1, period~resid(.))
> ## Q-Q plot
> qqmath(gm1)

```

380 The `ranef()` accessor extracts the conditional modes; the argument `condVar=TRUE` additionally
 381 extracts the variances of the conditional modes, which are stored as an attribute labelled
 382 `"postVar"`² — a three-dimensional array that gives the variance-covariance matrix of the
 383 conditional modes for each level of the grouping variable. The plotting methods `dotplot()`
 384 and `qqmath()` return lists of graphical objects showing *caterpillar plots* (ordered values of the
 385 random effects with confidence bars); in the case of the Q-Q plot (`qqmath`) the *y*-axis shows
 386 corresponding values of the standard normal quantiles (Figure 6).

387 `performance::check_model()` checks a variety of classical assumptions of GLMs. For mixed-
 388 effect models, it also assesses the normality of the distribution of conditional modes (Figure 7).

389 The `DHARMa` package is also compatible with `merMod` objects. `DHARMa` generates simulation-
 390 based residuals for generalized linear (mixed) models and uses them for graphical and statis-

²In earlier versions of the software, these elements were called “posterior variances” rather than “conditional variances”, based on the close correspondence between mixed models and hierarchical Bayesian models.

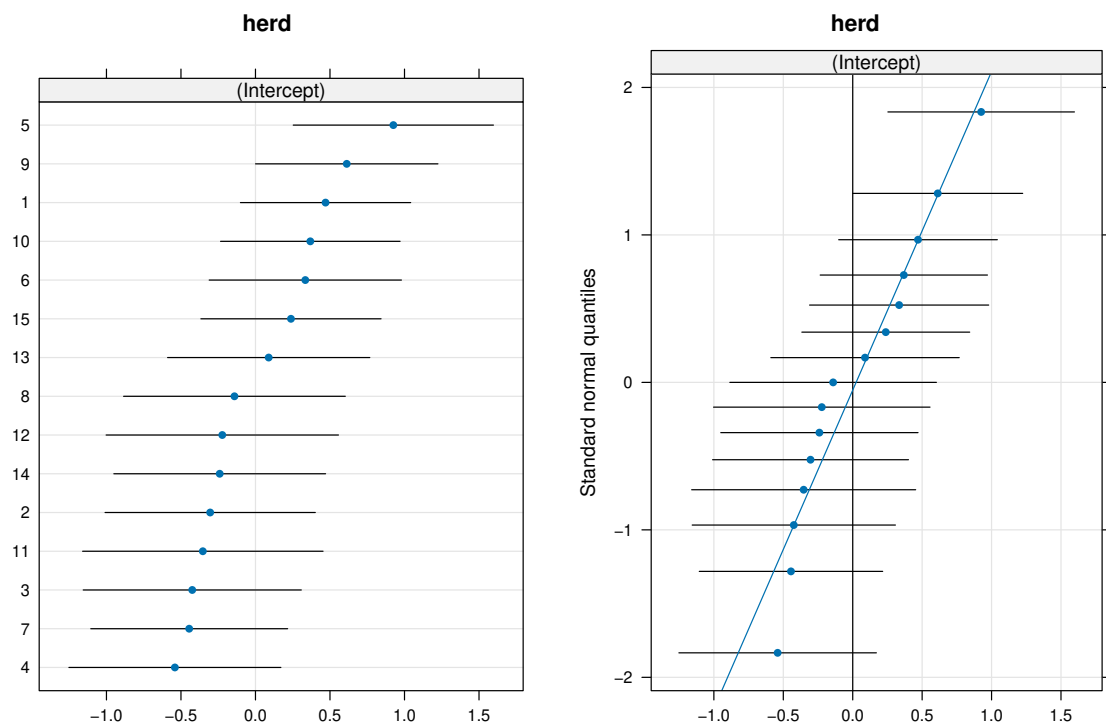
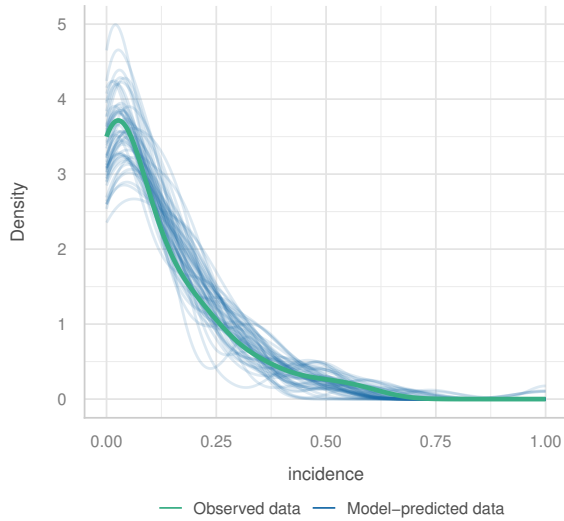


Figure 6: Graphical display of random effects. *Left:* conditional modes $\pm 1.96 \times$ conditional standard deviation, ordered by magnitude. *Right:* quantile-quantile plot, with linear regression line overlaid.

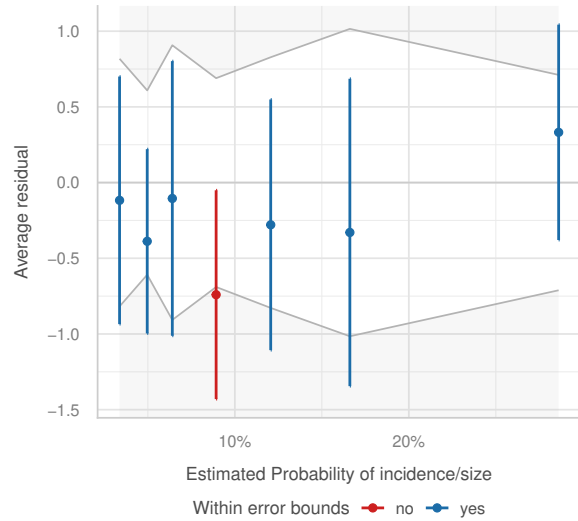
Posterior Predictive Check

Model-predicted lines should resemble observed data line



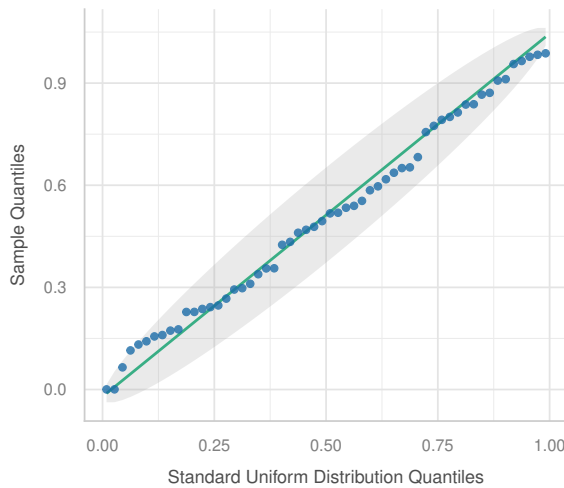
Binned Residuals

Points should be within error bounds



Distribution of Quantile Residuals

Dots should fall along the line



Normality of Random Effects (herd)

Dots should be plotted along the line

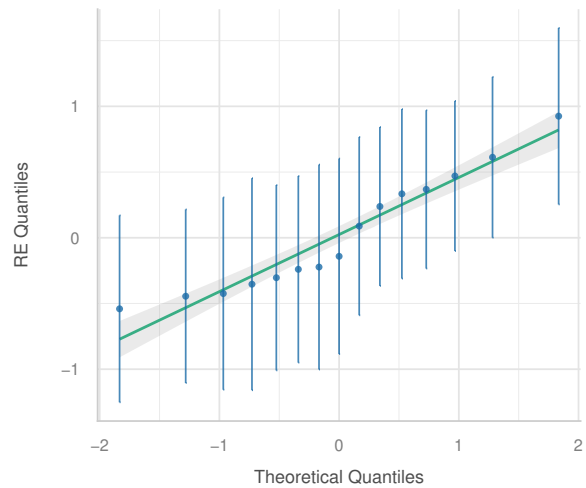


Figure 7: Model diagnostics using `performance::check_model`. The plot showcasing the normality of random effects (lower right panel) is the transpose of the right panel in Figure 6.

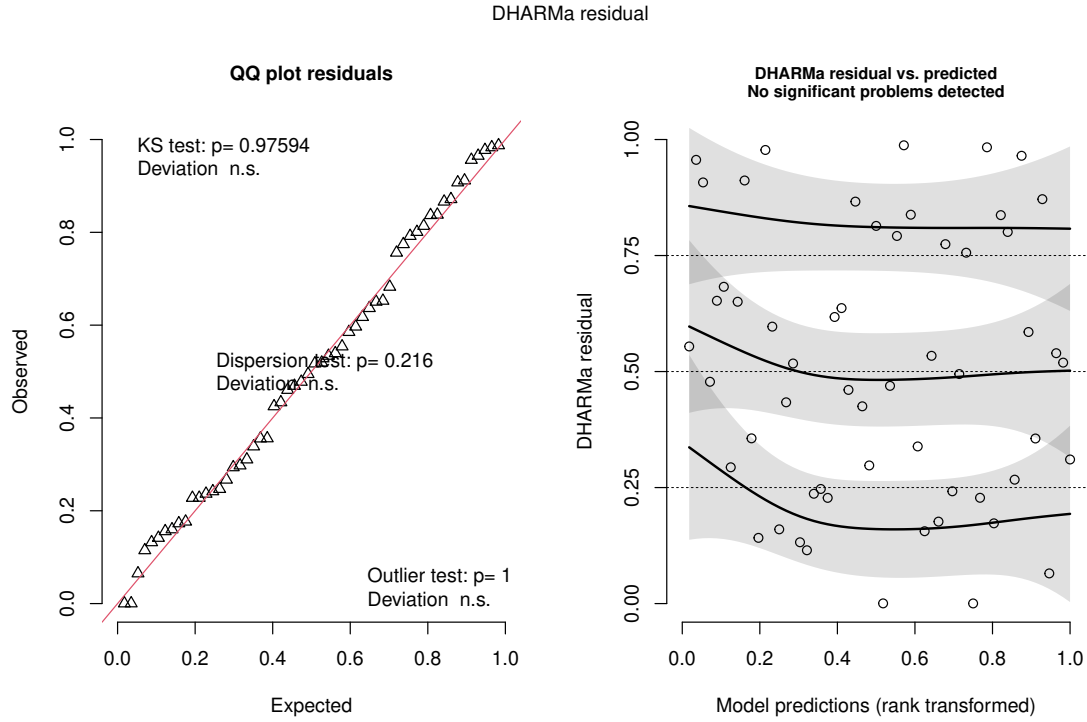


Figure 8: Model assumption checks using the DHARMa package. **n.s.** denotes “not significant”.

tical tests of model assumptions (Figure 8).

Having checked the diagnostics, we would now like to compare the three models we have fitted. Inspecting the **VarCorr** components, we see that when we fit both herd- and observation-level random effects, the among-herd variance is estimated as zero. The appropriate procedure at this point (e.g. whether one drops non-significant terms, or those with small scaled magnitudes, or those that worsen the AIC or BIC of the model) depends on the goals of the analysis and one’s philosophy of model-building (Barr, Levy, Scheepers, and Tily 2013; Matuschek, Kliegl, Vasishth, Baayen, and Bates 2017; Scandola and Tidoni 2024).

One might either stick with the full model, or continue with the reduced model with observation-level random effects only (as it has exactly the same likelihood as the full model but uses an additional parameter, it would be chosen according to either an information-theoretic or a hypothesis-testing model selection framework).

Here we will start by computing likelihood profiles and confidence intervals (CIs) for the model incorporating both random effects; although it has the same point estimates and maximum likelihood as the reduced model, confidence intervals that incorporate non-local information (i.e. profile- or parametric bootstrap-based) will give different, more conservative results for the full model.

The **profile** method computes profile likelihoods. The computation can be slow, since complete profiling for a model with p random- and fixed-effect parameters requires fitting p profiles, each of which requires many $p - 1$ -dimensional optimizations. The machinery for

generating likelihood profiles for GLMMs is similar to that for LMMs (see Bates *et al.* (2015), § 5.1).

The `profile` method returns an object of class `thpr` — a data frame containing the profiles, augmented with attributes containing interpolation splines for each parameter profile and their inverses (using `splines::interpSpline` and `splines::backSpline`); the latter are used for plotting profiles and computing confidence intervals. An `as.data.frame` method adds `.focal` and `.par` variables to the data frame, useful for customized plots.

Profiles can be used for univariate (`xyplot`) and bivariate (`splo`) profile plots, and to compute profile confidence intervals (`confint`). (`confint` applied to a `glmer` fit will first fit the profile, then use it to compute profile confidence intervals. Given the computational cost of profiling, it makes sense to compute and save the profile as an intermediate step if one plans to do anything other than computing confidence intervals.)

Two other common methods for computing confidence intervals are parametric bootstrapping (`method = "boot"`) and the classical Wald approximation (`method = "Wald"`). Parametric bootstrapping is much slower, but more accurate (and, via the `FUN` argument, can generate confidence intervals for any quantity that can be derived from a fitted model). The Wald approximation is faster and less accurate than profile confidence intervals. By default `glmer` only returns estimates for the fixed-effects parameters, as the assumptions of the Wald approximation are often violated badly for random-effects (co)variances and correlations. In the examples below we use the finite-difference Hessian (second derivative matrix of the estimated parameters) and the delta method to compute Wald confidence intervals for random-effects standard deviations and correlations, when possible.

Figure 9 compares all three of these confidence intervals across all three of the models fitted. If the default optimizers (BOBYQA followed by Nelder-Mead) do not perform well, one could attempt to re-fit the model with a variety of different optimizers using `allFit()`. The command `allFit(show.meth.tab=TRUE)` lists the optimizers `lme4` currently supports.

To use `allFit()`, supply the initial fitted model as input. Useful results are obtained from `summary(allFit())`; the components below (other than `$which.OK`) apply only to the optimizers that succeeded:

- `$which.OK` — which optimizers worked
- `$l1lik` — log-likelihoods
- `$fixef` — fixed-effect estimates
- `$sdcor` — random-effect standard deviations and correlations
- `$theta` — random-effect parameters on the Cholesky scale

We will show the summary results for `$sdcor`; the other components are similar.

```
> gm_all <- allFit(gm1)
```

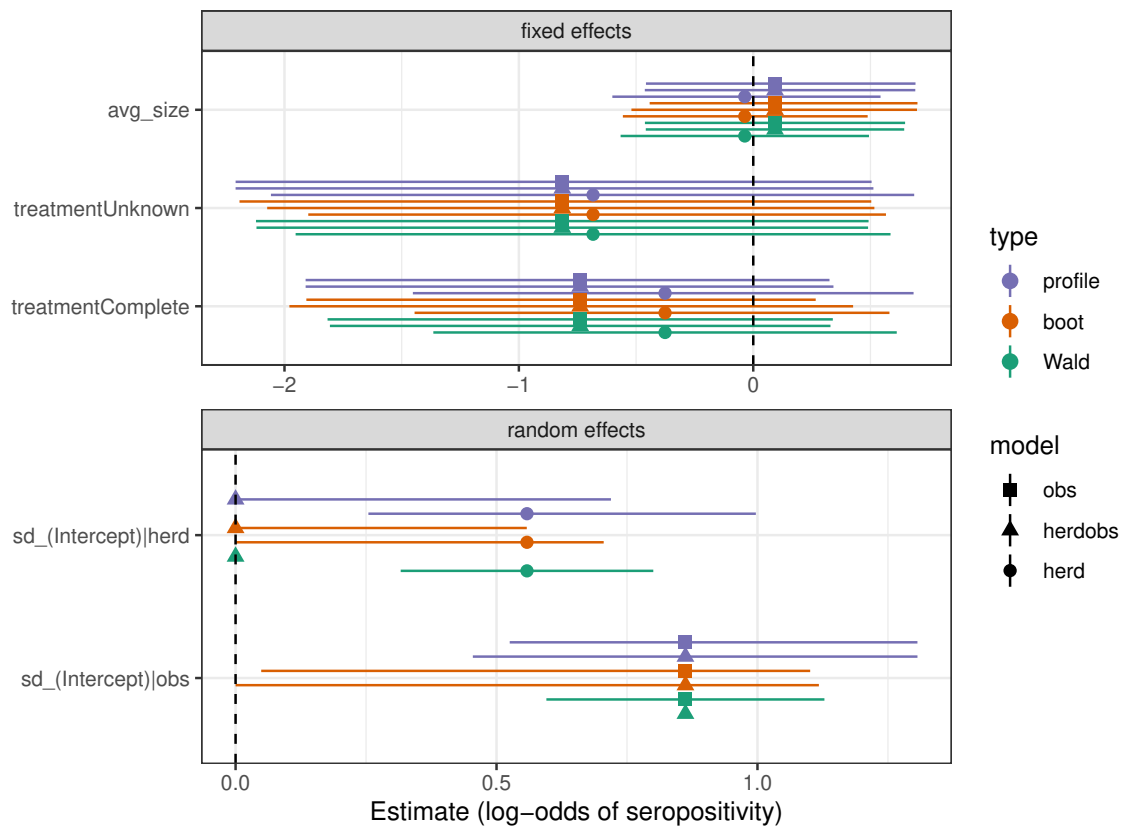


Figure 9: CBPP example: comparison of point and confidence interval estimation for different methods. Wald CIs are missing for random-effects parameters when the fitted model is singular.

```

bobyqa : [OK]
Nelder_Mead : [OK]
nlminbwrap : [OK]
nmkbw :
[OK]
optimx.L-BFGS-B : [OK]
nloptwrap.NLOPT_LN_NELDERMEAD : [OK]
nloptwrap.NLOPT_LN_BOBYQA : [OK]

```

```

> ss <- summary(gm_all)
> ss$sdcov

```

	herd.(Intercept)
bobyqa	0.55817
Nelder_Mead	0.55818
nlminbwrap	0.55818
nmkbw	0.55801
optimx.L-BFGS-B	0.55818
nloptwrap.NLOPT_LN_NELDERMEAD	0.55814
nloptwrap.NLOPT_LN_BOBYQA	0.55816

As of version 2.0, `lme4` can also specify structured variance-covariance matrices for (generalized) linear mixed models. `lme4` now supports unstructured (general positive definite), diagonal, compound symmetry, and first-order autoregressive (AR1) structures. By default, AR1 models assume a homogeneous-variance model (the variance is the same for all time steps), while the other models assume heterogeneous-variance models (variances differ for every level of the varying term); users can adjust this with the `hom` argument (e.g. `ar1(..., hom = FALSE)`).

The unstructured covariance structure is the default for mixed models. Here we illustrate fitting an AR1 model; the next example will show a compound symmetric model.

```

> gm.ar1 <- glmer(incidence/size ~ ar1(1 + herd | period),
+               family = binomial,
+               data = cbpp, weights = size)
> print(VarCorr(gm.ar1))

```

Groups Name	Std.Dev.	Corr
period (Intercept)	0.79	-0.04 (ar1)

For this particular model, a heterogeneous AR1 model (`ar1(..., hom = FALSE)`) results in a singular fit.

5.2. Contraception

Huq and Cleland (1990) use multilevel models to analyze data from a fertility survey of

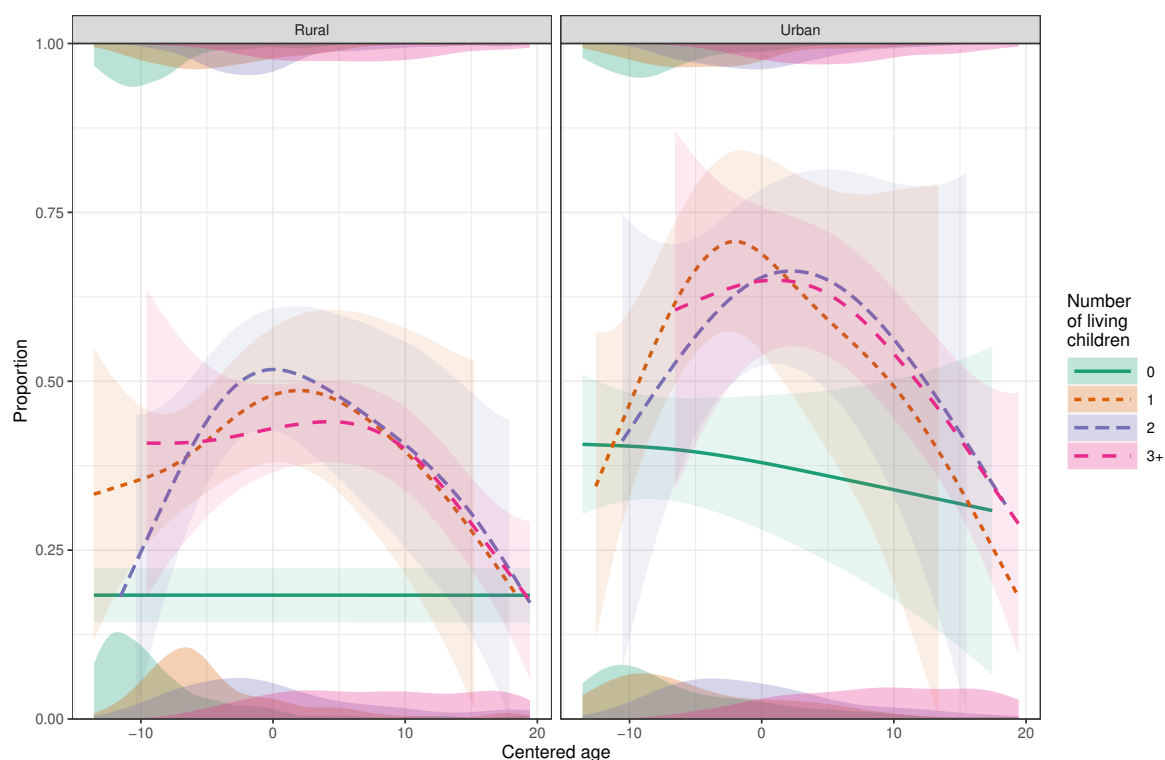


Figure 10: Contraception example: proportion of contraceptive use by centered age, number of living children (line type), and urban/rural residence (facet). Curves fitted using generalized additive models with cubic spline smoothing. Density plots along bottom and top margins show the age distributions of women not using contraception (bottom) and using contraception (top).

women in Bangladesh. These data are available as the `Contraception` object in the **mlmRev** package. The response variable is binary and indicates whether or not each woman was using contraception at the time of the survey. Covariates included the woman's age, the number of live children she had, whether she lived in an urban or rural setting, and the district in which she lived.

Figure 10 shows exploratory smooth curves of contraceptive use as a function of centered age, stratified by number of living children and urban/rural residence. In rural areas, women with two living children show the highest rates of contraceptive use, peaking near 50% at the average age, while women in urban areas with one living child show the highest rates near the average age but with a more pronounced decline at older ages. In both settings, women with no living children consistently show the lowest rates of contraceptive use. The nonmonotonic effect of age on contraceptive use and the apparent dependence of this age trend on child status motivate the model specifications explored below. (These nonmonotonic trends were not noticed in the original analysis of the data, which concluded that there was no significant effect of age.)

We construct six `glmer` models with varying choices of explanatory variables and random effects structure, comparing them via `anova`. The models considered different variables.

	df	$\Delta\text{negloglik}$	ΔAIC
binary_child $\times$ age + (1 district:urban)	7	0.472	0.00
binary_child $\times$ age + (1 district/urban)	8	0.467	1.99
binary_child $\times$ age + (1 + urban district)	9	0.000	3.06
binary_child $\times$ age + (1 district)	7	5.825	10.71
binary_child + age + (1 district)	6	9.828	16.71
int_child + age + (1 district)	8	9.599	20.25

Table 1: Model comparison for Contraception fits. Note that for likelihood ratio tests, or for AIC comparisons restricted to nested models (Ripley 2004), the nesting sequence for the random-effects models is (1+urban|district) > (1|district/urban) > {(1|district), (1|district:urban)}.

For instance, there are two different ways to encode the number of living children: `livch` is a four-level factor distinguishing 0, 1, 2, or 3+ children, while other models use `ch` instead, which is a binary indicator for whether the woman has any living children (in Table 1 and Figure 12 we label this variable as `binary_child`).

Second, we vary whether child status interacts with the woman’s age: two models include `ch` (or `livch`) and `age` as additive terms, while the other four include a `ch` and `age` interaction. A quadratic effect of age ($I(\text{age}^2)$) was added to account for the nonlinear effect of age.

Third, we explore different random effects structures at the district level. Three of them use a single random intercept per district (1 | district). One of them extends this with a random slope for urban status (urban | district), allowing the urban/rural difference to vary by district. The next uses nested random intercepts for district and site within district (1 | district/urban) separating district-level from urban-within-district variation. (Scandola and Tidoni (2024) refer to this formulation as a “complex random intercepts” model). The remaining model uses only the `urban:district` grouping. To reduce convergence warnings and facilitate interpretability, age (already approximately centered in the data set) was standardized by scaling by twice the standard deviation (Gelman 2008).

```

> cm1 <- glmer(use ~ age_s + I(age_s^2) + urban + livch + (1|district),
+               Contraception, binomial)
> ## switch from livch (ordinal) to ch (binary)
> cm2 <- update(cm1, . ~ . - livch + ch)
> ## add age by children interaction
> cm3 <- update(cm2, . ~ . + age_s:ch)
> ## allow urban effect to vary across districts (correlated)
> cm4 <- update(cm3, . ~ . - (1|district) + (1+urban|district))
> ## compound symmetric/nested formulation
> cm5 <- update(cm3, . ~ . - (1|district) + (1 | district/urban))
> ## as above but drop district effect
> cm6 <- update(cm3, . ~ . - (1|district) + (1 | district:urban))

```

Table 1 shows that the top three models, all of which include a child-by-age interaction and some effect of urbanization, fit approximately equally well (Δ negative log-likelihood < 0.5). The (1|district/urban) model barely improves on the fit of (1|district:urban)

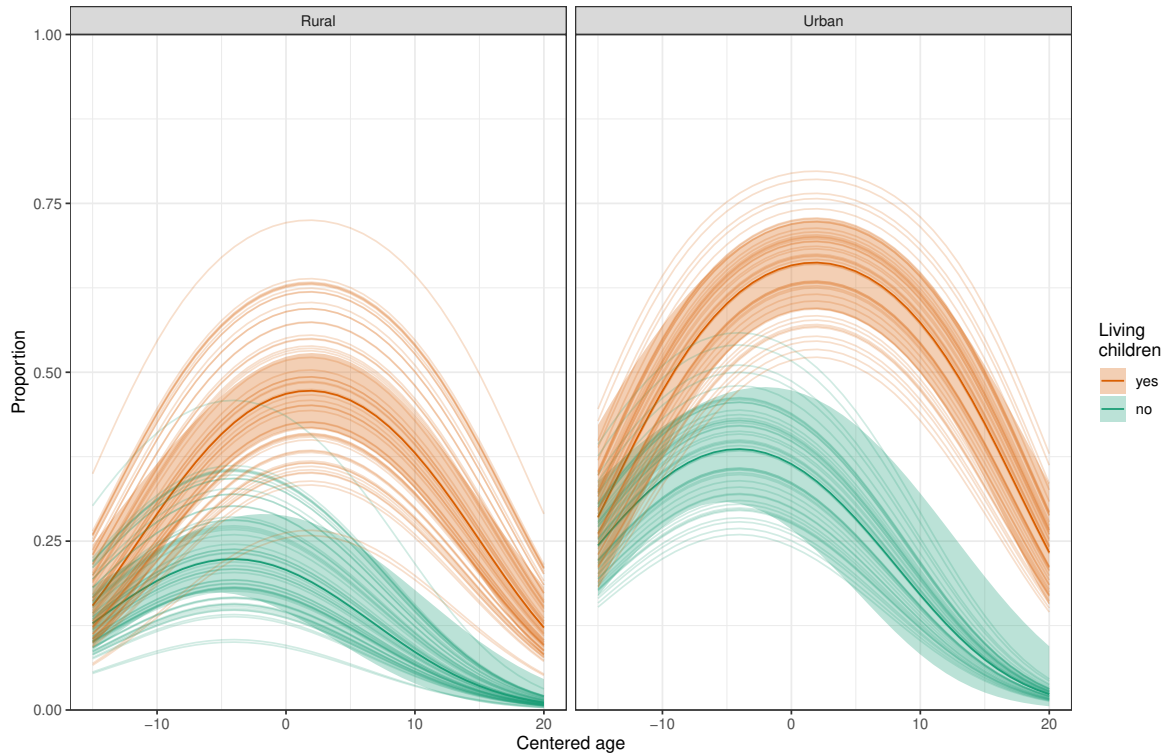


Figure 11: Predictions for contraception fit ((1|district/urban) model). Heavy lines and ribbons show population-level predictions; light lines show district-level predictions.

(0.005 log-likelihood units), at the cost of an extra variance parameter, so it is almost 2 AIC units worse. $(1 + \text{urban} \mid \text{district})$ is a bit better (≈ 0.5 log-likelihood units), but includes a covariance parameter. Models that drop the interaction between child status and age or replace the binary child indicator with an integer count perform substantially worse ($\Delta\text{AIC} > 10$).

Overall, the results suggest that accounting for urban/rural variation at the district level and including the age and child interaction are both important, but the precise random effects structure for urban/rural variation is relatively unimportant.

As with the CBPP data set, we can also compute profile, Wald, and parametric bootstrap confidence intervals for all of the models to understand the effects of each variable and visualize the among-model variation.

The point estimates and confidence intervals of the explanatory variables are similar across the six models, and across different methods for confidence interval construction, with a few exceptions.

Urban residence, presence of living children, and age were all strong predictors of contraceptive use among women in Bangladesh (because we are fitting a quadratic model for the effect of age, the non-significant effect of `age_s` simply means that the marginal effect of age *at the mean age* is not clearly negative or positive).

Returning to the structured covariance matrices introduced in Section 5.1, we can replace the unstructured random effect $(1 + \text{urban} \mid \text{district})$ with diagonal (`diag`) or compound

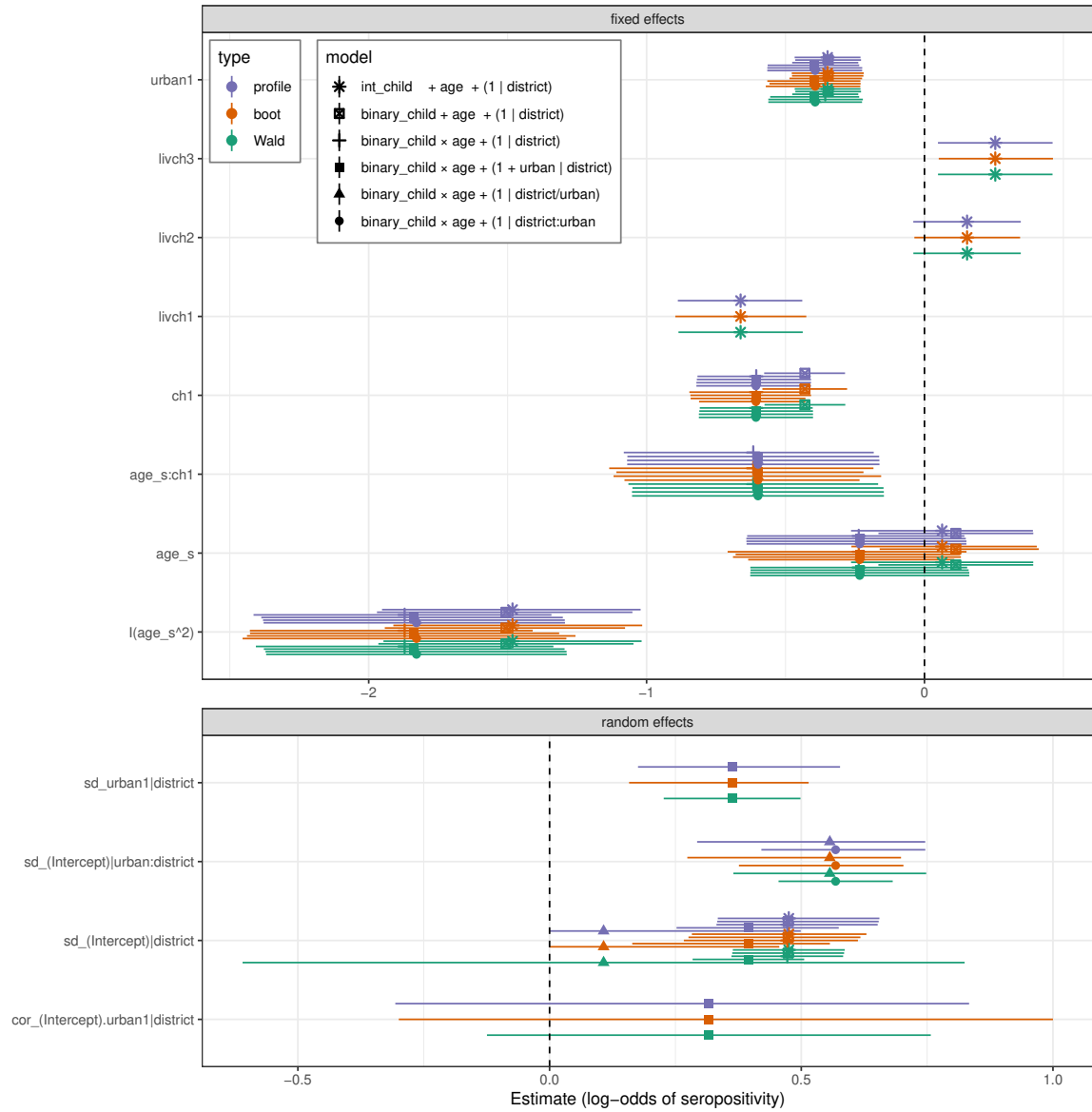


Figure 12: Contraception example: comparison of point and confidence interval estimation for different methods. (Note that the Wald CI for variation in the intercept across districts, in the (1|district/urban) model, includes negative values.)

515 symmetric (`cs`) structures, with either heterogeneous or homogeneous variances (via the `hom`
 516 argument):

517 For example, the compound symmetric model is fitted as follows:

```
> cm.cs <- glmer(use ~ age*ch + I(age^2) + urban + cs(1 + urban | district),
+               Contraception, binomial, control = cc)
```

518 We can use `VarCorr` and other accessor methods as we would for models with default, un-
 519 structured covariance matrices, e.g.:

```
> print(VarCorr(cm.cs))

Groups   Name             Std.Dev. Corr
district (Intercept) 0.615      -0.79 (cs)
          urbanY      0.725
```

Acknowledgements

520 We would like to thank Steve Walker, for early work on `glmer`, and Alex Stringer, for helpful
 521 comments on the manuscript.

References

- 522 Barr DJ, Levy R, Scheepers C, Tily HJ (2013). “Random Effects Structure for Confirmatory
 523 Hypothesis Testing: Keep It Maximal.” *Journal of Memory and Language*, **68**(3), 255–278.
 524 doi:10.1016/j.jml.2012.11.001.
- 525 Bates D, Mächler M, Bolker B, Walker S (2015). “Fitting linear mixed-effects models using
 526 lme4.” *Journal of Statistical Software*, **67**, 1–48. doi:10.18637/jss.v067.i01.
- 527 De Backer M, De Vroey C, Lesaffre E, Scheys I, De Keyser P (1998). “Twelve weeks of
 528 continuous oral therapy for toenail onychomycosis caused by dermatophytes: a double-
 529 blind comparative trial of terbinafine 250 mg/day versus itraconazole 200 mg/day.” *Journal*
 530 *of the American Academy of Dermatology*, **38**(5, Supplement 2), S57–S63. doi:10.1016/
 531 S0190-9622(98)70486-4.
- 532 Gelman A (2008). “Scaling regression inputs by dividing by two standard deviations.” *Statis-*
 533 *tics in medicine*, **27**(15), 2865–2873. doi:10.1002/sim.3107.
- 534 Harrison XA (2015). “A Comparison of Observation-Level Random Effect and Beta-Binomial
 535 Models for Modelling Overdispersion in Binomial Data in Ecology and Evolution.” *PeerJ*,
 536 **3**, e1114. ISSN 2167-8359. doi:10.7717/peerj.1114.
- 537 Huq N, Cleland J (1990). *Bangladesh fertility survey 1989*. National Institute of Population
 538 Research and Training, Dhaka.

- Kristensen K, Nielsen A, Berg CW, Skaug H, Bell BM (2016). “TMB : Automatic Differentiation and Laplace Approximation.” *Journal of Statistical Software*, **70**(5). doi:10.18637/jss.v070.i05.
- Lesnoff M, Laval G, Bonnet P, Abdicho S, Workalemahu A, Kifle D, Peyraud A, Lancelot R, Thiaucourt F (2004). “Within-herd spread of contagious bovine pleuropneumonia in Ethiopian highlands.” *Preventive Veterinary Medicine*, **64**(1), 27–40. doi:10.1016/j.prevetmed.2004.03.005.
- Madsen H, Thyregod P (2011). *Introduction to General and Generalized Linear Models*. CRC Press. ISBN 978-1-4200-9155-7.
- Matuschek H, Kliegl R, Vasishth S, Baayen H, Bates D (2017). “Balancing Type I Error and Power in Linear Mixed Models.” *Journal of Memory and Language*, **94**, 305–315. doi:10.1016/j.jml.2017.01.001.
- McCullagh P, Nelder JA (1989). *Generalized Linear Models*. Chapman and Hall, London.
- Powell MJD (2009). “The BOBYQA Algorithm for Bound Constrained Optimization without Derivatives.” *Technical Report DAMTP 2009/NA06*, Centre for Mathematical Sciences, University of Cambridge, Cambridge, England. URL http://www.damtp.cam.ac.uk/user/na/NA_papers/NA2009_06.pdf.
- Rabe-Hesketh S, Skrondal A, Pickles A (2005). “Maximum Likelihood Estimation of Limited and Discrete Dependent Variable Models with Nested Random Effects.” *Journal of Econometrics*, **128**(2), 301–323. ISSN 0304-4076. doi:10.1016/j.jeconom.2004.08.017.
- Ripley BD (2004). “Selecting amongst Large Classes of Models.” In NM Adams, M Crowder, D Hand, D Stephens (eds.), *Methods and models in statistics: In honor of Professor John Nelder, FRS*, pp. 155–170. Imperial College Press.
- Scandola M, Tidoni E (2024). “Reliability and Feasibility of Linear Mixed Models in Fully Crossed Experimental Designs.” *Advances in Methods and Practices in Psychological Science*, **7**(1), 25152459231214454. doi:10.1177/25152459231214454.
- Stringer A (2024). “Exact Gradient Evaluation for Adaptive Quadrature Approximate Marginal Likelihood in Mixed Models for Grouped Data.” *Statistics and Computing*, **35**(1), 4. ISSN 1573-1375. doi:10.1007/s11222-024-10536-z.
- Stringer A, Bilodeau B, Tang Y (2026). “Asymptotics of Numerical Integration for Two-Level Mixed Models.” *Bernoulli*. (forthcoming).
- Tuerlinckx F, Rijmen F, Verbeke G, De Boeck P (2006). “Statistical Inference in Generalized Linear Mixed Models: A Review.” *British Journal of Mathematical and Statistical Psychology*, **59**(2), 225–255. ISSN 2044-8317. doi:10.1348/000711005X79857.
- Zeger SL, Karim MR (1991). “Generalized linear models with random effects: a Gibbs sampling approach.” *Journal of the American Statistical Association*, **86**(413), 79–86. doi:10.1080/01621459.1991.10475006.

Package versions used

Compiled with R Under development (unstable) (2026-07-15 r90257) and package versions
lme4: 2.0.6, performance: 0.17.1, DHARMa: 0.5.0, see: 0.14.1.

6. Appendix: derivation of PIRLS

We seek to maximize the unscaled conditional log density for a GLMM over the conditional modes, $\mathbf{u}$. This problem is very similar to maximizing the log-likelihood for a GLM, which is a very thoroughly studied problem (e.g. McCullagh and Nelder 1989). The standard algorithm for dealing with this kind of problem is iteratively reweighted least squares (IRLS). Here we modify IRLS by incorporating a penalty term that accounts for variation in the random effects; we call the resulting algorithm penalized iteratively reweighted least squares (PIRLS). The unscaled conditional log-density takes the form,

$$f(\mathbf{u}) = \log p(\mathbf{y}, \mathbf{u} | \boldsymbol{\beta}, \boldsymbol{\theta}) = \boldsymbol{\psi}^\top \mathbf{A} \mathbf{y} - \mathbf{a}^\top \boldsymbol{\phi} + \mathbf{c} - \frac{1}{2} \mathbf{u}^\top \mathbf{u} - \frac{q}{2} \log 2\pi \quad (25)$$

where $\boldsymbol{\psi}$ is the n -by-1 canonical parameter of an exponential family, $\boldsymbol{\phi}$ is the n -by-1 vector of cumulant functions, $\mathbf{c}$ an n -by-1 vector of normalizing constants, and $\mathbf{A}$ is an n -by- n diagonal matrix of prior weights, $\mathbf{a}$. Both $\mathbf{a}$ and $\mathbf{c}$ could depend on a dispersion parameter, although we ignore this possibility for now.

The canonical parameter, $\boldsymbol{\psi}$, and vector of cumulant functions, $\boldsymbol{\phi}$, depend on a linear predictor,

$$\boldsymbol{\eta} = \mathbf{o} + \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\boldsymbol{\Lambda}_\theta \mathbf{u} \quad (26)$$

where $\mathbf{o}$ is an n -by-1 vector of *a priori* offsets. The specific form of this dependency is specified by the choice of the exponential family. The mean of this distribution, $\boldsymbol{\mu}$, is the *inverse link function* g^{-1} applied to $\boldsymbol{\eta}$.

Our goal is to find the values of $\mathbf{u}$ that maximize the unscaled conditional density, for given $\boldsymbol{\theta}$ and $\boldsymbol{\beta}$ vectors. These maximizers are the conditional modes, which we require for the Laplace approximation and adaptive Gauss-Hermite quadrature. To do this maximization we use a variant of the Fisher scoring method, which is the basis of the iteratively reweighted least squares algorithm for generalized linear models. Fisher scoring is itself based on Newton's method, which we apply first.

6.1. Newton's method

To apply Newton's method, we need the gradient and the Hessian of the unscaled conditional log-likelihood. Following standard GLM theory (McCullagh and Nelder 1989), we use the chain rule,

$$\frac{dL(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}, \mathbf{u})}{d\mathbf{u}} = \frac{dL'(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}, \mathbf{u})}{d\boldsymbol{\psi}} \frac{d\boldsymbol{\psi}}{d\boldsymbol{\mu}} \frac{d\boldsymbol{\mu}}{d\boldsymbol{\eta}} \frac{d\boldsymbol{\eta}}{d\mathbf{u}} - \frac{1}{2} \frac{d(\mathbf{u}^\top \mathbf{u})}{d\mathbf{u}}$$

where L' represents the log-density in (25) exclusive of the penalty term. The first derivative in the first term's chain follows from basic results in GLM theory,

$$\frac{dL'(\boldsymbol{\beta}, \boldsymbol{\theta} | \mathbf{y}, \mathbf{u})}{d\boldsymbol{\psi}} = (\mathbf{y} - \boldsymbol{\mu})^\top \mathbf{A}.$$

Again from standard GLM theory, the next two derivatives define the inverse diagonal variance matrix,

$$\frac{d\psi}{d\boldsymbol{\mu}} = \mathbf{V}^{-1}$$

and the diagonal Jacobian matrix,

$$\frac{d\boldsymbol{\mu}}{d\boldsymbol{\eta}} = \mathbf{M} \quad .$$

Finally, because $\mathbf{u}$ affects $\boldsymbol{\eta}$ only linearly,

$$\frac{d\boldsymbol{\eta}}{d\mathbf{u}} = \mathbf{Z}\boldsymbol{\Lambda}_\theta$$

Therefore we have,

$$\frac{dL(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u}} = (\mathbf{y} - \boldsymbol{\mu})^\top \mathbf{A}\mathbf{V}^{-1}\mathbf{M}\mathbf{Z}\boldsymbol{\Lambda}_\theta - \mathbf{u}^\top \quad . \quad (27)$$

This is very similar to the gradient for GLMs with respect to fixed effects coefficients, $\boldsymbol{\beta}$. The only difference induced by including the penalty term for the random effects ($\mathbf{u}$), is the subtraction of the $\mathbf{u}^\top$ term.

Again we apply the chain rule to take the Hessian,

$$\frac{d^2L(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} = \frac{d^2L'(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\boldsymbol{\mu}} \frac{d\boldsymbol{\mu}}{d\boldsymbol{\eta}} \frac{d\boldsymbol{\eta}}{d\mathbf{u}} - \mathbf{I}_q \quad (28)$$

which leads to,

$$\frac{d^2L(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} = \frac{d^2L'(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\boldsymbol{\mu}} \mathbf{M}\mathbf{Z}\boldsymbol{\Lambda}_\theta - \mathbf{I}_q \quad (29)$$

The first derivative in this chain can be expressed as,

$$\frac{d^2L'(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\boldsymbol{\mu}} = -\boldsymbol{\Lambda}_\theta^\top \mathbf{Z}^\top \mathbf{M}\mathbf{V}^{-1} \mathbf{A} + \boldsymbol{\Lambda}_\theta^\top \mathbf{Z}^\top \left[\frac{d\mathbf{M}\mathbf{V}^{-1}}{d\boldsymbol{\mu}} \right] \mathbf{A}\mathbf{R} \quad (30)$$

where $\mathbf{R}$ is a diagonal residuals matrix with $\mathbf{y} - \boldsymbol{\mu}$ on the diagonal. The two terms arise from a type of product rule, where we first differentiate the residuals, $\mathbf{y} - \boldsymbol{\mu}$, and then the diagonal matrix, $\mathbf{M}\mathbf{V}^{-1}$, with respect to $\boldsymbol{\mu}$.

The Hessian can therefore be expressed as,

$$\frac{d^2L(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} = -\boldsymbol{\Lambda}_\theta^\top \mathbf{Z}^\top \mathbf{M}\mathbf{A}^{1/2} \mathbf{V}^{-1/2} \left(\mathbf{I}_n - \mathbf{V}\mathbf{M}^{-1} \left[\frac{d\mathbf{M}\mathbf{V}^{-1}}{d\boldsymbol{\mu}} \right] \mathbf{R} \right) \mathbf{V}^{-1/2} \mathbf{A}^{1/2} \mathbf{M}\mathbf{Z}\boldsymbol{\Lambda}_\theta - \mathbf{I}_q \quad (31)$$

This result can be simplified by expressing it in terms of a weighted random-effects design matrix, $\mathbf{U} = \mathbf{A}^{1/2} \mathbf{V}^{-1/2} \mathbf{M}\mathbf{Z}\boldsymbol{\Lambda}_\theta$,

$$\frac{d^2L(\boldsymbol{\beta}, \boldsymbol{\theta}|\mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} = -\mathbf{U}^\top \left(\mathbf{I}_n - \mathbf{V}\mathbf{M}^{-1} \left[\frac{d\mathbf{V}^{-1}\mathbf{M}}{d\boldsymbol{\mu}} \right] \mathbf{R} \right) \mathbf{U} - \mathbf{I}_q \quad . \quad (32)$$

6.2. Fisher-like scoring

There are two ways to further simplify this expression for $\mathbf{U}^\top \mathbf{U}$. The first is to use the canonical link function for the family being used. Canonical links have the property that $\mathbf{V} = \mathbf{M}$, which means that for canonical links,

$$\frac{d^2 L(\beta, \boldsymbol{\theta} | \mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} = -\mathbf{U}^\top \left(\mathbf{I}_n - \mathbf{I}_n \left[\frac{d\mathbf{I}_n}{d\boldsymbol{\mu}} \right] \mathbf{R} \right) \mathbf{U} - \mathbf{I}_q = -\mathbf{U}^\top \mathbf{U} - \mathbf{I}_q \quad (33)$$

The second way to simplify the Hessian is to take its expectation with respect to the distribution of the response, conditional on the current values of the spherical random effects coefficients, $\mathbf{u}$. The diagonal residual matrix, $\mathbf{R}$, has expectation 0. Therefore, because the response only enters into the expression for the Hessian via $\mathbf{R}$, we have that,

$$E \left(\frac{d^2 L(\beta, \boldsymbol{\theta} | \mathbf{y}, \mathbf{u})}{d\mathbf{u} d\mathbf{u}} | \mathbf{u} \right) = -\mathbf{U}^\top \left(\mathbf{I}_n - \mathbf{U} \mathbf{M}^{-1} \left[\frac{d\mathbf{V}^{-1} \mathbf{M}}{d\boldsymbol{\mu}} \right] E(\mathbf{R}) \right) \mathbf{U} - \mathbf{I}_q = -\mathbf{U}^\top \mathbf{U} - \mathbf{I}_q \quad (34)$$

Like base R's `glm`, `glmer` uses Fisher scoring throughout (which is identical to Newton's method for canonical links) to enable the use of the (P)IRLS algorithm. As usual, we minimize the conditional negative log-density rather than maximizing the conditional log-density.

Affiliation:

Anna Ly

Department of Mathematics & Statistics

McMaster University

1280 Main Street W

Hamilton, ON L8S 4K1, Canada

Email: annahuynh.ly@utoronto.ca

Rune Haubo Bojesen Christensen

Copenhagen Research Centre for Biological and Precision Psychiatry

Mental Health Centre Copenhagen, Copenhagen University Hospital - Bispebjerg and Frederiksberg

Copenhagen, Denmark

Email: Rune.Haubo@pm.me

Douglas Bates

Department of Statistics, University of Wisconsin - Madison

1205 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

658 Benjamin M. Bolker
659 Departments of Mathematics & Statistics and Biology
660 McMaster University
661 1280 Main Street W
662 Hamilton, ON L8S 4K1, Canada
663 E-mail: bolker@mcmaster.ca