

A new REML (PX)EM algorithm for linear mixed models and factor analytic mixed models

Simon Malcolm Diffey

April 2012

A thesis submitted for the degree of Doctor of Philosophy
of the Australian National University



Declaration

The work in this thesis is my own except where otherwise stated.

A handwritten signature in black ink, appearing to read 'S. Diffey'. The signature is fluid and cursive, with a horizontal line under the 'S' and a stylized 'D'.

Simon Malcolm Diffey

In memory of Andrew John Diffey (30/7/1976 to 30/12/2010) who died unexpectedly of an asthma attack. A hard worker who would have appreciated the amount of work this thesis represents.

Acknowledgements

This PhD would not have been possible without the financial support of the Grains Research Development Corporation and the New South Wales Department of Primary Industries. I would like to thank Brian Cullis for organising this support and his encouragement and advice throughout.

In his role as supervisor I would like to thank Alan Welsh for his guidance, patience, and seeing it through to the end. I would also like to thank Robin Thompson for his interest and thought provoking comments and Alison Smith for helping to clarify technical issues.

Finally, I would like to thank my family for their love, support, and understanding. In particular I wish to thank my wife Kristy. During the course of this PhD we have had two sons, bought a farm, and lost a brother. She has handled it all with grace, and where appropriate, good humour.

Abstract

Linear mixed models and factor analytic mixed models are routinely applied to biological data arising from designed experiments. The preferred method for estimating the parameters associated with these models is residual maximum likelihood (REML). Most statistical software packages available for the REML estimation of parameters associated with linear mixed models and factor analytic mixed models implement a Newton-Raphson type algorithm such as the expected information algorithm or the average information algorithm. There are two problems with these algorithms. Firstly, successive iterations of these algorithms are not guaranteed to increase the residual log-likelihood function. Secondly, parameter updates may not remain in their parameter space. Either may result in the algorithm failing to converge to a solution.

The REML expectation maximisation (REML EM) algorithm and the parameter expanded version of this algorithm (REML PX EM) are alternatives to Newton-Raphson type algorithms. Features of these two algorithms are that the residual log-likelihood may not decrease with successive iterations and parameter updates remain in their parameter space. Before the REML EM or REML PX EM algorithm can be considered practical alternatives to Newton-Raphson type algorithms two issues need to be addressed. Firstly, they can be notoriously slow to converge, particularly the REML EM algorithm. Secondly, compared to the average information algorithm, current implementations of these two algorithms are computationally more expensive at each iterate. This increased computational expense relates to calculating the trace of a matrix of the same order as the length of the observed data vector.

This thesis addresses the issues of speed and computational efficiency of the REML EM and REML PX EM algorithms for linear mixed models and factor analytic mixed models. The REML EM and REML PX EM algorithms require specification of the incomplete and missing data. A new specification of the incomplete data for linear mixed models and factor analytic mixed models is introduced which is shown to be

computationally more efficient and we describe the conditions under which this new specification will have a faster rate of convergence. For factor analytic mixed models, model specification, a new parameter expansion, a new missing data specification, and the efficacy of using a less stringent stopping rule are considered. In an example plant breeding data set and in a simulation study it is shown that these innovations can drastically reduce the number of iterations to convergence.

The improvements to the REML EM and REML PX EM algorithms presented in this thesis make these two algorithms, particularly the latter, more likely to be implemented alongside Newton-Raphson type algorithms in statistical software packages for linear mixed models and factor analytic mixed models. In such a situation this would provide users of these models a viable alternative in the event of a Newton-Raphson type algorithm failing.

Contents

Acknowledgements	vii
Abstract	ix
List of main notations	xix
1 EM algorithm introduction	1
1.1 A brief history of the EM algorithm	1
1.2 EM algorithm derivation	2
1.3 Score statistic	7
1.4 Missing information principle and observed information matrix	8
1.5 Rate of convergence	11
2 The linear mixed model as an incomplete data problem and a conditional derivation of REML	15
2.1 Linear mixed model specification	15
2.2 Henderson’s mixed model equations	16
2.3 Missing data specification for the linear mixed model	18
2.4 Conditional derivation of REML	19
3 Current REML EM algorithms for the linear mixed model	23
3.1 Linear mixed model specification - random effect approach	24

3.2	Complete data specification - random effect approach	25
3.3	Linear mixed model specification - fixed effect approach	26
3.4	Complete data specification - fixed effect approach	27
3.5	REML EM algorithm E-step	28
3.6	REML EM algorithm M-step	29
3.6.1	Update for σ^2	30
3.6.2	Update for γ_{gh}	30
3.6.3	Update for ϕ_{ij}	30
3.7	Complete data information matrix	31
3.7.1	Element $\mathcal{I}_{c_{(\sigma^2)_2}}^{(P)}$	32
3.7.2	Elements of $\mathcal{I}_{c_{\gamma,\gamma}}^{(P)}$	32
3.7.3	Elements of $\mathcal{I}_{c_{\phi,\phi}}^{(P)}$	33
3.7.4	Elements of $\mathcal{I}_{c_{\sigma^2,\phi}}^{(P)}$	33
3.8	Missing data information matrix	34
3.8.1	Element $\mathcal{I}_{m_{(\sigma^2)_2}}^{(P)}$	35
3.8.2	Elements of $\mathcal{I}_{m_{\gamma,\gamma}}^{(P)}$	35
3.8.3	Elements of $\mathcal{I}_{m_{\phi,\phi}}^{(P)}$	35
3.8.4	Elements of $\mathcal{I}_{m_{\sigma^2,\gamma}}^{(P)}$	36
3.8.5	Elements of $\mathcal{I}_{m_{\sigma^2,\phi}}^{(P)}$	36
3.8.6	Elements of $\mathcal{I}_{m_{\gamma,\phi}}^{(P)}$	37
4	A new REML EM algorithm for linear mixed models	39
4.1	New complete data specification	39
4.2	REML EM algorithm E-step	41
4.3	REML EM algorithm M-step	42
4.3.1	Update for σ^2	42
4.3.2	Update for γ_{gh}	43

4.3.3	Update for ϕ_{ij}	43
4.4	Complete data information matrix	44
4.4.1	Element $\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)}$	44
4.4.2	Elements of $\mathcal{I}_{c_{\gamma,\gamma}}^{(U)}$	45
4.4.3	Elements of $\mathcal{I}_{c_{\phi,\phi}}^{(U)}$	45
4.4.4	Elements of $\mathcal{I}_{c_{\sigma^2,\phi}}^{(U)}$	46
4.5	Missing data information matrix	46
4.5.1	Element $\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)}$	47
4.5.2	Elements of $\mathcal{I}_{m_{\gamma,\gamma}}^{(U)}$	48
4.5.3	Elements of $\mathcal{I}_{m_{\phi,\phi}}^{(U)}$	48
4.5.4	Elements of $\mathcal{I}_{m_{\sigma^2,\gamma}}^{(U)}$	48
4.5.5	Elements of $\mathcal{I}_{m_{\sigma^2,\phi}}^{(U)}$	49
4.5.6	Elements of $\mathcal{I}_{m_{\gamma,\phi}}^{(U)}$	49
4.6	Comparing rates of convergence between $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$	50
4.7	Example: lamb weight data	52
5	Two REML PX EM algorithms for the linear mixed model	57
5.1	Linear mixed model specification	57
5.2	REML PX EM algorithm for $\mathbf{y}_c^{(P)}$	59
5.2.1	Complete data log density function	59
5.2.2	REML PX EM algorithm E-step	60
5.2.3	REML PX EM algorithm M-step	61
5.2.4	Rate matrix	63
5.3	REML PX EM algorithm for $\mathbf{y}_c^{(U)}$	65
5.3.1	Complete data log density function	65
5.3.2	REML PX EM algorithm E-step	66

5.3.3	REML PX EM algorithm M-step	67
5.3.4	Rate matrix	69
5.4	Example: lamb weight data	71
6	Two model specifications for a factor analytic mixed model	75
6.1	Model specification	77
6.2	Complete data density function	79
6.3	REML EM algorithm E-step for Model $\mathcal{M}$	80
6.4	REML EM algorithm M-step for Model $\mathcal{M}$	82
6.4.1	Update for σ^2	82
6.4.2	Update for λ_{jr}	82
6.4.3	Update for ψ	84
6.4.4	Update for d_{qr}	85
6.5	Example: National Variety Trial data	87
7	REML PX EM algorithms for two factor analytic mixed model specifications	91
7.1	REML PX EM algorithm - Model $\mathcal{X}$	92
7.1.1	Complete data density function	92
7.2	REML PX EM algorithm E-step for Model $\mathcal{X}$	93
7.3	REML PX EM algorithm M-step for Model $\mathcal{X}$	94
7.3.1	Update for σ_{*}^2	94
7.3.2	Update for λ_{*jr}	94
7.3.3	Update for ψ_{*}	95
7.3.4	Update for d_{qr}	95
7.4	REML PX EM algorithm - Model $\mathcal{Z}$	95
7.4.1	Complete data density function	96

7.5	REML PX EM algorithm E-step for Model $\mathcal{Z}$	97
7.6	REML PX EM algorithm M-step for Model $\mathcal{Z}$	97
7.6.1	Update for σ_*^2	98
7.6.2	Update for λ_{*jr}	98
7.6.3	Update for ψ_*	99
7.6.4	Update for d_{*qr}	99
7.6.5	Update for ω	99
7.7	Example: National Variety Trial data	99
8	A new missing data specification for factor analytic mixed models	103
8.1	REML EM algorithm - Model $\mathcal{M}$	104
8.1.1	Complete data density function for $\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$. . .	105
8.1.2	Equivalence of $\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$ and $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$.	107
8.2	REML EM algorithm E-step for Model $\mathcal{M}$ using $\mathbf{y}_c^{(D)}$	108
8.3	REML EM algorithm M-step for Model $\mathcal{M}$ using $\mathbf{y}_c^{(D)}$	109
8.3.1	Update for σ^2	109
8.3.2	Update for λ_{jr}	109
8.3.3	Update for ψ	111
8.3.4	Update for d_{qr}	111
8.4	REML PX EM algorithm - Model $\mathcal{Z}$	111
8.4.1	Complete data density function	112
8.5	REML PX EM algorithm E-step for Model $\mathcal{Z}$ using $\mathbf{y}_c^{(D)}$	114
8.6	REML PX EM algorithm M-step for Model $\mathcal{Z}$ using $\mathbf{y}_c^{(D)}$	114
8.6.1	Update for σ_*^2	114
8.6.2	Update for λ_{*jr}	115
8.6.3	Update for ψ_*	115

8.6.4	Update for d_{*qr}	116
8.6.5	Update for ω	116
8.7	Example: National Variety Trial data	116
8.8	Simulation study	117
9	Conclusion	125
A	Vector and matrix results	129
A.1	Conditional distribution between two vectors	129
A.2	Covariance between two vectors	129
A.3	Matrix identities	130
A.3.1	Trace of a matrix	130
A.3.2	Quadratic forms	130
A.3.3	Determinants	130
A.3.4	Inverses of matrices of the form $\mathbf{A} + \mathbf{B}$	131
A.3.5	Inverses of matrices of the form $\mathbf{A} + \mathbf{BCD}$	132
A.3.6	Inverses of partitioned matrices	132
A.3.7	$\mathbf{GZ}^T(\mathbf{R} + \mathbf{ZGZ}^T)^{-1} = (\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1}$	133
A.3.8	$\mathbf{P} = \mathbf{L}_2(\mathbf{L}_2^T\mathbf{H}\mathbf{L}_2)^{-1}\mathbf{L}_2^T = \mathbf{H}^{-1} - \mathbf{H}^{-1}\mathbf{X}(\mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-1}$	133
A.3.9	$\mathbf{P} = \mathbf{R}^{-1} - \mathbf{R}^{-1}\mathbf{WC}^{-1}\mathbf{W}^T\mathbf{R}^{-1}$	134
A.3.10	$\mathbf{P} = \mathbf{S} - \mathbf{SZC}^{ZZ}\mathbf{Z}^T\mathbf{S}$	135
A.3.11	$\mathbf{PHP} = \mathbf{P}$	136
A.3.12	$\mathbf{PHS} = \mathbf{S}$	136
A.3.13	$\mathbf{PRS} = \mathbf{P}$	137
A.3.14	$\text{tr}(\mathbf{HP}) = n - t$	137
A.4	Matrix differentiation results	138
A.4.1	Derivatives of determinants	138

<i>CONTENTS</i>	xvii
A.4.2 Derivatives of inverses	138
A.4.3 Derivatives of $\boldsymbol{P}$	138
B Derivation of identity in Oakes (1999)	141
C Loewner partial ordering of symmetric positive semidefinite matrices	145
Bibliography	147

List of main notations

a	scalar
$\mathbf{a}$	vector
$\mathbf{A}$	matrix
$\mathbf{I}_n$	$n \times n$ identity matrix
$\mathbf{a}^T$	transpose of the vector $\mathbf{a}$
$\mathbf{A}^T$	transpose of the matrix $\mathbf{A}$
$\log$	natural logarithm
$\det$	determinant
tr	trace
E	expectation
V	variance
Cov	covariance
N	Gaussian distribution
L	likelihood function
ℓ	log-likelihood function
$\frac{\partial}{\partial \boldsymbol{\theta}}$	partial derivative with respect to the parameter vector $\boldsymbol{\theta}$
$\left\langle \cdot \right\rangle_{\boldsymbol{\theta} = \boldsymbol{\theta}_*}$	evaluate at $\boldsymbol{\theta} = \boldsymbol{\theta}_*$
$ \cdot $	absolute value
$\ \cdot\ $	vector norm
$\ \cdot\ _F$	Frobenius matrix norm

Chapter 1

EM algorithm introduction

In this chapter we describe the EM algorithm in general terms, i.e., in terms which are not specific to linear mixed models or factor analytic mixed models. In particular we provide a derivation of the EM algorithm, introduce the missing information principle, describe a method for calculating the observed information matrix and for computing the rate of convergence of an EM algorithm. The rate of convergence will be used to compare two different REML EM algorithms for the linear mixed model in a later chapter.

1.1 A brief history of the EM algorithm

The term “EM algorithm” first appeared in the Dempster, Laird, and Rubin paper published in the Royal Statistical Society Series B journal in 1977 (Dempster et al., 1977). However, as noted by Dempster et al. (1977) a number of other authors had contributed to the underlying theory and method prior to their famous paper.

This history of papers with an EM flavour earlier than the Dempster et al. (1977) paper is hardly surprising as the EM algorithm was developed to compute maximum likelihood estimates from incomplete data - a problem of great practical interest. However, this previous work does not detract from the importance of the Dempster et al. (1977) paper. McLachlan and Krishnan (1997) note that it was in this paper that

ideas were synthesized, a general formulation of the EM algorithm established, its properties investigated, and a host of traditional and nontraditional appli-

cations indicated.

The EM algorithm’s popularity can be attributed to the number of problems that fall under the “incomplete data” umbrella. Besides the obvious cases where there are missing values in the data, a large number of problems can be formulated as incomplete data problems even though the incompleteness of the data is not immediately obvious. Linear mixed models, which are the subject of this thesis, are one such example.

1.2 EM algorithm derivation

The EM algorithm is an iterative optimisation method for computing maximum likelihood estimates (and residual or restricted maximum likelihood estimates) in cases where the observed data can be considered incomplete. The EM algorithm derives its name from two steps that are performed at each iteration. These are an expectation step (E-step) which is followed by a maximisation step (M-step).

There are a number of ways in which the EM algorithm can be derived. We present the traditional approach presented in the Dempster et al. (1977) paper. We will consider the observed data vector $\mathbf{y}_o = (y_{o_1}, \dots, y_{o_n})^T$ as incomplete and a random sample from a distribution having a probability density function that is parameterised by $\boldsymbol{\theta} = (\theta_1, \dots, \theta_d)^T$, a vector of unknown parameters with parameter space Θ . It is assumed that there exists a complete data vector $\mathbf{y}_c$ which consists of the observed data vector and the missing data vector $\mathbf{y}_m$, i.e., $\mathbf{y}_c = (\mathbf{y}_o^T, \mathbf{y}_m^T)^T$.

Define the density function for the complete data vector as

$$f(\mathbf{y}_c; \boldsymbol{\theta}) = f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta}).$$

If we let $S(\mathbf{y}_m)$ denote the sample space of $\mathbf{y}_m$, the density function for the observed data vector can be expressed as

$$g(\mathbf{y}_o; \boldsymbol{\theta}) = \int_{S(\mathbf{y}_m)} f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta}) d\mathbf{y}_m.$$

The density function for the observed data vector can also be defined in terms of the complete data vector and a conditional density function, i.e.,

$$g(\mathbf{y}_o; \boldsymbol{\theta}) = \frac{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})}.$$

Taking natural logarithms we have

$$\log\{g(\mathbf{y}_o; \boldsymbol{\theta})\} = \log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} - \log\{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\} \quad (1.1)$$

$$\ell(\boldsymbol{\theta}; \mathbf{y}_o) = \log\{f(\mathbf{y}_c; \boldsymbol{\theta})\} - \log\{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\}, \quad (1.2)$$

where $\ell(\boldsymbol{\theta}; \mathbf{y}_o) = \log\{g(\mathbf{y}_o; \boldsymbol{\theta})\}$ is a log-likelihood function based on the observed data and which will be referred to as the observed data log-likelihood function. Note that the log-density function $\log\{f(\mathbf{y}_c; \boldsymbol{\theta})\}$ cannot be referred to as a log-likelihood function as $\mathbf{y}_m$ is not observed. Maximum likelihood estimates for $\boldsymbol{\theta}$ are obtained by finding values of $\boldsymbol{\theta}$ that maximize $\ell(\boldsymbol{\theta}; \mathbf{y}_o)$, i.e.,

$$\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta}} \ell(\boldsymbol{\theta}; \mathbf{y}_o).$$

The problem is that $\hat{\boldsymbol{\theta}}$ is often difficult to compute. The EM algorithm approaches this problem by finding the maximum likelihood estimate of $\boldsymbol{\theta}$ indirectly by considering the complete data log-density function, i.e., $\log\{f(\mathbf{y}_c; \boldsymbol{\theta})\}$. As this function is unobservable it is replaced with the current conditional expectation of $\log\{f(\mathbf{y}_c; \boldsymbol{\theta})\}$ given the observed data $\mathbf{y}_o$ and the current estimate of $\boldsymbol{\theta}$.

Let the w -th iterate of $\boldsymbol{\theta}$ be denoted $\boldsymbol{\theta}^{(w)}$. Multiplying both sides of (1.2) by $k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})$ and integrating with respect to $\mathbf{y}_m$ we have

$$\ell(\boldsymbol{\theta}; \mathbf{y}_o) = Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)}), \quad (1.3)$$

where

$$\begin{aligned} Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)}) &= \int_{S(\mathbf{y}_m)} \log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m, \\ H(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)}) &= \int_{S(\mathbf{y}_m)} \log\{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m. \end{aligned}$$

The E and M-steps of the EM algorithm can be written

$$\begin{aligned}
 \text{E-step: Calculate } Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)}) &= \int_{S(\mathbf{y}_m)} \log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &= E[\log\{f(\mathbf{y}_c; \boldsymbol{\theta})\} | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}] \\
 \text{M-step: } \boldsymbol{\theta}^{(w+1)} &= \arg \max_{\boldsymbol{\theta}} Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(w)}).
 \end{aligned}$$

Although we have ignored the $H(\cdot)$ function we can form a bound on the difference $H(\boldsymbol{\theta}^{(w+1)}, \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w)}, \boldsymbol{\theta}^{(w)})$ through the use of Jensen's inequality. Forming this bound enables us to show that each successive iteration of the EM algorithm does not decrease the observed data log-likelihood function and that iteratively maximising $Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})$ will maximise $\ell(\boldsymbol{\theta}; \mathbf{y}_o)$. This attribute of not decreasing the observed data log-likelihood with successive iterations is known as monotonic convergence. For the EM algorithm

$$\begin{aligned}
 H(\boldsymbol{\theta}^{(w+1)}, \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w)}, \boldsymbol{\theta}^{(w)}) &= \int_{S(\mathbf{y}_m)} \log\{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w+1)})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &\quad - \int_{S(\mathbf{y}_m)} \log\{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &= \int_{S(\mathbf{y}_m)} \log \left\{ \frac{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w+1)})}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})} \right\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &\leq \log \left\{ \int_{S(\mathbf{y}_m)} \frac{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w+1)})}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \right\} \\
 &= \log \left\{ \int_{S(\mathbf{y}_m)} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w+1)}) d\mathbf{y}_m \right\} \\
 &= 0,
 \end{aligned} \tag{1.4}$$

where the inequality in (1.4) is a result of Jensen's inequality. It follows that

$$-\{H(\boldsymbol{\theta}^{(w+1)}, \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w)}, \boldsymbol{\theta}^{(w)})\} \geq 0. \tag{1.5}$$

The inequality (1.5) becomes relevant if we consider finding $\boldsymbol{\theta}^{(w+1)}$ such that

$$\begin{aligned}
 \ell(\boldsymbol{\theta}^{(w+1)}; \mathbf{y}_o) &\geq \ell(\boldsymbol{\theta}^{(w)}; \mathbf{y}_o) \\
 Q(\boldsymbol{\theta}^{(w+1)}; \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w+1)}; \boldsymbol{\theta}^{(w)}) &\geq Q(\boldsymbol{\theta}^{(w)}; \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w)}; \boldsymbol{\theta}^{(w)}) \\
 Q(\boldsymbol{\theta}^{(w+1)}; \boldsymbol{\theta}^{(w)}) - Q(\boldsymbol{\theta}^{(w)}; \boldsymbol{\theta}^{(w)}) - \{H(\boldsymbol{\theta}^{(w+1)}; \boldsymbol{\theta}^{(w)}) - H(\boldsymbol{\theta}^{(w)}; \boldsymbol{\theta}^{(w)})\} &\geq 0.
 \end{aligned}
 \tag{1.6}$$

The inequality presented in (1.6) is positive for any $\boldsymbol{\theta}^{(w+1)}$ if

$$Q(\boldsymbol{\theta}^{(w+1)}; \boldsymbol{\theta}^{(w)}) - Q(\boldsymbol{\theta}^{(w)}; \boldsymbol{\theta}^{(w)}) \geq 0.$$

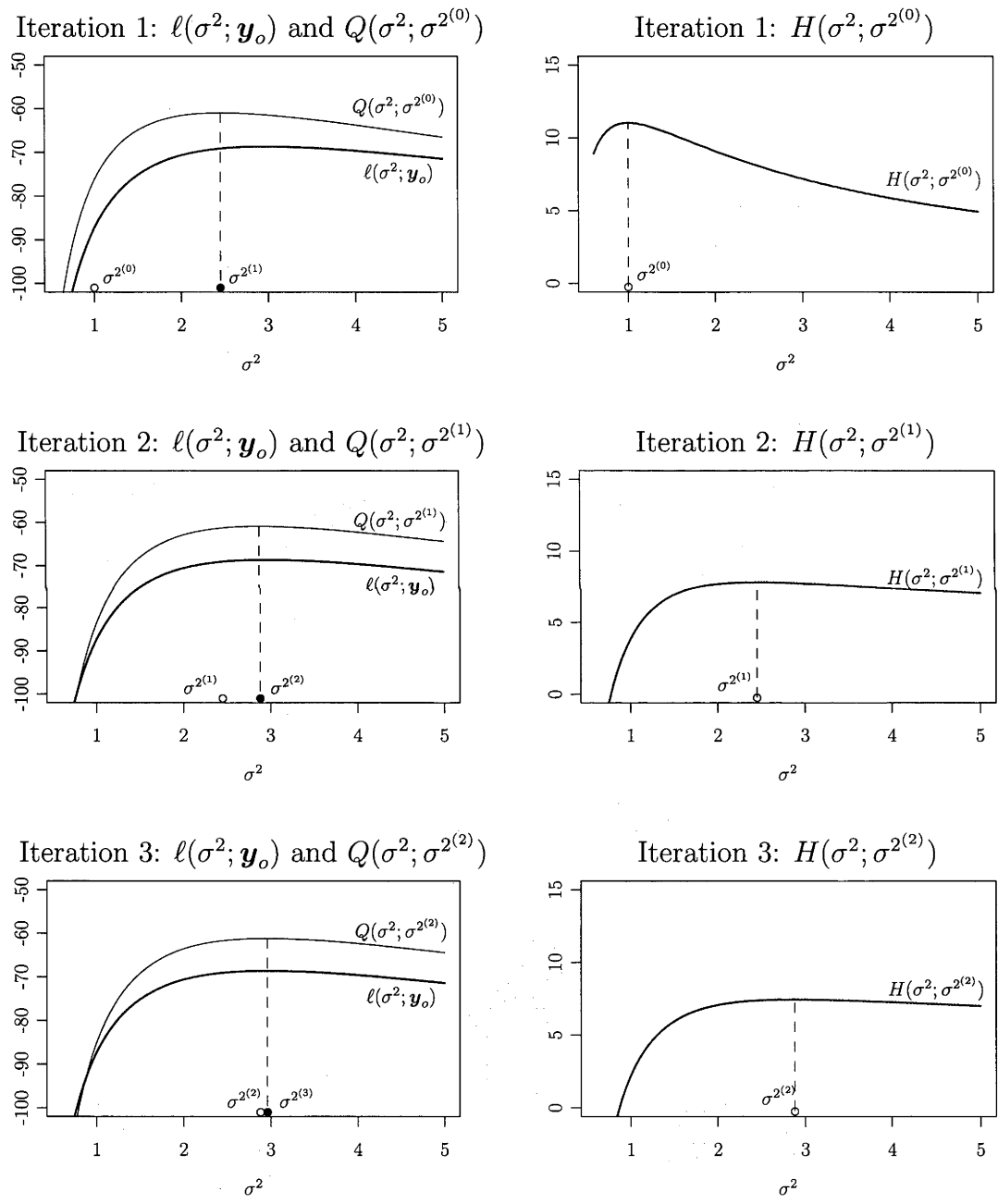
Choosing $\boldsymbol{\theta}^{(w+1)}$ such that $Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})$ is maximised for fixed $\boldsymbol{\theta}^{(w)}$ is the M-step of the EM algorithm and it follows that

$$\ell(\boldsymbol{\theta}^{(w+1)}; \mathbf{y}_o) \geq \ell(\boldsymbol{\theta}^{(w)}; \mathbf{y}_o),$$

establishing the result that each iteration of the EM algorithm increases the observed data log-likelihood and iteratively maximising $Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})$ will maximise $\ell(\boldsymbol{\theta}; \mathbf{y}_o)$. Jensen's inequality also appears in Wald's proof of the consistency of maximum likelihood estimators (Wald, 1949) and is therefore fundamental to maximum likelihood estimation.

A graphical illustration of iteratively maximising $Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})$ is provided in Figure 1.1. The main purpose of Figure 1.1 is to graphically illustrate that $H(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})$ is maximised at $\boldsymbol{\theta} = \boldsymbol{\theta}^{(w)}$. This result is fundamental to the EM algorithm having monotonic convergence in $\ell(\boldsymbol{\theta}; \mathbf{y}_o)$.

Figure 1.1: Graphical illustration of the EM algorithm being used to compute a maximum likelihood estimate of $\theta = \sigma^2$ by iteratively maximising $Q(\theta; \theta^{(w)})$ (for 3 iterations) and of $H(\theta; \theta^{(w)})$ being maximised at $\theta = \theta^{(w)}$.



1.3 Score statistic

The gradient vector of the observed data log-likelihood function is

$$\mathcal{S}_o(\mathbf{y}_o; \boldsymbol{\theta}) = \frac{\partial \log\{g(\mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta}}$$

and the gradient vector of the complete data log-density function is

$$\mathcal{S}_c(\mathbf{y}_c; \boldsymbol{\theta}) = \frac{\partial \log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta}}.$$

The observed data score statistic can be expressed as the conditional expectation of the complete data score statistic (Orchard and Woodbury, 1972), i.e.,

$$\mathcal{S}_o(\mathbf{y}_o; \boldsymbol{\theta}) = E\{\mathcal{S}_c(\mathbf{y}_c; \boldsymbol{\theta}) | \mathbf{y}_o; \boldsymbol{\theta}\}.$$

This result is developed below, assuming regularity conditions that permit the interchange of differentiation and integration hold. We can write

$$\begin{aligned} \mathcal{S}_o(\mathbf{y}_o; \boldsymbol{\theta}) &= \frac{\partial \log\{g(\mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta}} \\ &= \frac{\frac{\partial}{\partial \boldsymbol{\theta}} \{g(\mathbf{y}_o; \boldsymbol{\theta})\}}{g(\mathbf{y}_o; \boldsymbol{\theta})} \\ &= \frac{1}{g(\mathbf{y}_o; \boldsymbol{\theta})} \frac{\partial}{\partial \boldsymbol{\theta}} \left\{ \int_{S(\mathbf{y}_m)} f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta}) d\mathbf{y}_m \right\} \\ &= \frac{1}{g(\mathbf{y}_o; \boldsymbol{\theta})} \int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}} \left\{ f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta}) \right\} d\mathbf{y}_m \\ &= \int_{S(\mathbf{y}_m)} \frac{\frac{\partial}{\partial \boldsymbol{\theta}} \{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\}}{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})} \frac{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})}{g(\mathbf{y}_o; \boldsymbol{\theta})} d\mathbf{y}_m \\ &= \int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} \right] k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}) d\mathbf{y}_m \\ &= E\{\mathcal{S}_c(\mathbf{y}_c; \boldsymbol{\theta}) | \mathbf{y}_o; \boldsymbol{\theta}\}. \end{aligned}$$

We can also express the observed data score statistic as

$$\begin{aligned}
 \mathcal{S}_o(\mathbf{y}_o; \boldsymbol{\theta}) &= \left\langle \frac{\partial Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta}} \right\rangle_{\boldsymbol{\theta}^{(w)} = \boldsymbol{\theta}} \quad (1.7) \\
 &= \left\langle \int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right] d\mathbf{y}_m \right\rangle_{\boldsymbol{\theta}^{(w)} = \boldsymbol{\theta}} \\
 &= \left\langle \int_{S(\mathbf{y}_m)} \frac{\frac{\partial}{\partial \boldsymbol{\theta}} \{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\}}{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \right\rangle_{\boldsymbol{\theta}^{(w)} = \boldsymbol{\theta}} \\
 &= \left\langle \frac{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})}{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})} \frac{\partial}{\partial \boldsymbol{\theta}} \left\{ \int_{S(\mathbf{y}_m)} f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta}) d\mathbf{y}_m \right\} \right\rangle_{\boldsymbol{\theta}^{(w)} = \boldsymbol{\theta}} \\
 &= \left\langle \frac{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})}{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})} \frac{\partial}{\partial \boldsymbol{\theta}} \left\{ g(\mathbf{y}_o; \boldsymbol{\theta}) \right\} \right\rangle_{\boldsymbol{\theta}^{(w)} = \boldsymbol{\theta}} \\
 &= \frac{\frac{\partial}{\partial \boldsymbol{\theta}} \{g(\mathbf{y}_o; \boldsymbol{\theta})\}}{g(\mathbf{y}_o; \boldsymbol{\theta})} \\
 &= \frac{\partial \log\{g(\mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta}},
 \end{aligned}$$

as required. This result is referred to as the self-consistency property of the EM algorithm.

1.4 Missing information principle and observed information matrix

Often the justification for using the EM algorithm instead of another algorithm is that the complete data log-likelihood function is easier to consider than the observed data log-likelihood function. However, an early criticism of the EM algorithm was that it didn't provide an estimate of the observed information matrix and, to estimate this matrix, derivatives of the observed data log-likelihood function had to be considered anyway. Two notable papers that address the issue of finding the observed information matrix when using the EM algorithm are Louis (1982) and Oakes (1999). We will consider the derivation of the observed information matrix presented by Oakes (1999) where the observed information matrix can be estimated using only derivatives of the $Q(\cdot)$ function. Before proceeding with the Oakes (1999) derivation of the observed information matrix it is useful to consider the so-called missing information principle. Beginning with the relation in (1.1) we can write

$$-\frac{\partial^2}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \left[\log \{g(\mathbf{y}_o; \boldsymbol{\theta})\} \right] = -\frac{\partial^2}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \left[\log \{f(\mathbf{y}_o, \mathbf{y}_m; \boldsymbol{\theta})\} \right] + \frac{\partial^2}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \left[\log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\} \right].$$

Taking expectations on both sides over $k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})$ we can write

$$E\{I_o(\boldsymbol{\theta}; \mathbf{y}_o) | \mathbf{y}_o; \boldsymbol{\theta}\} = E\{I_c(\boldsymbol{\theta}; \mathbf{y}_c) | \mathbf{y}_o; \boldsymbol{\theta}\} - E\{I_m(\boldsymbol{\theta}; \mathbf{y}_m) | \mathbf{y}_o; \boldsymbol{\theta}\}$$

which can be re-expressed as

$$I_o(\boldsymbol{\theta}; \mathbf{y}_o) = \mathcal{I}_c(\boldsymbol{\theta}; \mathbf{y}_c) - \mathcal{I}_m(\boldsymbol{\theta}; \mathbf{y}_m),$$

where the symbol $\mathcal{I}$ indicates that it is a conditional expectation of an information matrix that is being considered. The above identity has been interpreted as

$$\text{observed information} = \text{complete information} - \text{missing information}$$

which has been called the missing information principle and can be attributed to Orchard and Woodbury (1972). We now present the Oakes (1999) derivation of the observed information matrix using only derivatives of the $Q(\cdot)$ function. In the following development the regularity conditions required for the interchange of differentiation and integration are assumed. Consider the second order partial derivatives,

$$\begin{aligned} \frac{\partial^2 \ell(\boldsymbol{\theta}; \mathbf{y}_o)}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} &= \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} - \frac{\partial^2 H(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \\ &= \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} - E \left[\frac{\partial^2 \log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \end{aligned} \quad (1.8)$$

and

$$\begin{aligned} \frac{\partial^2 \ell(\boldsymbol{\theta}; \mathbf{y}_o)}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} &= \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} - \frac{\partial^2 H(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \\ \mathbf{0} &= \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} - E \left[\frac{\partial^2 \log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right]. \end{aligned} \quad (1.9)$$

The second term on the right hand side of (1.9) can be written

$$\begin{aligned}
 & E \left[\frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
 &= \frac{\partial}{\partial \boldsymbol{\theta}} \left(\int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\} k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right] d\mathbf{y}_m \right) \\
 &= \int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\} \frac{\partial k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta}^{(w)T}} \right] d\mathbf{y}_m \\
 &= \int_{S(\mathbf{y}_m)} \frac{\frac{\partial}{\partial \boldsymbol{\theta}} \{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})} \frac{\frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})} k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &= \int_{S(\mathbf{y}_m)} \frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
 &= E \left(\frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right). \quad (1.10)
 \end{aligned}$$

Substituting (1.10) into (1.9) and adding (1.8) and (1.9) we have

$$\begin{aligned}
 \frac{\partial^2 \ell(\boldsymbol{\theta}; \mathbf{y}_o)}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} &= \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} + \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \\
 &\quad - E \left[\frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
 &\quad - E \left(\frac{\partial}{\partial \boldsymbol{\theta}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right). \quad (1.11)
 \end{aligned}$$

A key result in Oakes (1999) is the identity

$$\begin{aligned}
 & E \left[\frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
 &= -E \left(\frac{\partial}{\partial \boldsymbol{\theta}^{(w)}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right) \quad (1.12)
 \end{aligned}$$

which is derived in Appendix B. Evaluating (1.11) at $\boldsymbol{\theta} = \boldsymbol{\theta}^{(w)}$ and using the identity in (1.12) we have

$$\frac{\partial^2 \ell(\boldsymbol{\theta}^{(w)}; \mathbf{y}_o)}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} = \left\langle \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} + \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \right\rangle_{\boldsymbol{\theta}=\boldsymbol{\theta}^{(w)}} \quad (1.13)$$

which is the result presented by Oakes (1999). It follows that the complete and missing information matrices are given by

$$\begin{aligned} \mathcal{I}_c(\boldsymbol{\theta}^{(w)}; \mathbf{y}_c) &= - \left\langle \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T} \right\rangle_{\boldsymbol{\theta}=\boldsymbol{\theta}^{(w)}}, \\ \mathcal{I}_m(\boldsymbol{\theta}^{(w)}; \mathbf{y}_m) &= \left\langle \frac{\partial^2 Q(\boldsymbol{\theta}; \boldsymbol{\theta}^{(w)})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^{(w)T}} \right\rangle_{\boldsymbol{\theta}=\boldsymbol{\theta}^{(w)}}. \end{aligned}$$

As noted by Oakes (1999) some noteworthy features of (1.13) are that the relation is valid for all $\boldsymbol{\theta}^{(w)}$, the second term on the right hand side can be interpreted as the missing information, and the partitioning of the observed information matrix into two parts may be of some benefit when the inverse of the observed information matrix is required.

1.5 Rate of convergence

In this thesis we will be comparing EM algorithms using different complete data specifications. One measure of algorithm performance is the number of iterations the algorithm takes to converge. A more rigorous approach to measuring algorithm performance is to measure the rate of convergence. For θ_i , $i = 1, \dots, d$ the rate of convergence can be measured empirically using

$$r_i = \lim_{w \rightarrow \infty} \frac{|\theta_i^{(w+1)} - \theta_i^{(w)}|}{|\theta_i^{(w)} - \theta_i^{(w-1)}|}.$$

The global rate of convergence can be measured empirically using

$$r = \lim_{w \rightarrow \infty} \frac{\|\boldsymbol{\theta}^{(w+1)} - \boldsymbol{\theta}^{(w)}\|}{\|\boldsymbol{\theta}^{(w)} - \boldsymbol{\theta}^{(w-1)}\|}$$

or by

$$r = \max_i r_i \quad i = 1, \dots, d.$$

Empirical rates of convergence are only useful when considering an EM algorithm applied to specific examples. A more general formulation of the global rate of convergence can be derived by considering each iteration of the EM algorithm to involve a mapping $\mathcal{M}$ from a set of values $\boldsymbol{\theta}^{(w)}$ in Θ to another set of values $\boldsymbol{\theta}^{(w+1)}$ also in Θ , i.e.,

$$\boldsymbol{\theta}^{(w+1)} = \mathcal{M}(\boldsymbol{\theta}^{(w)}) \quad w = 0, 1, 2, \dots$$

Following the development in McLachlan and Krishnan (1997) we can say that if $\boldsymbol{\theta}^{(w)}$ converges to some point $\boldsymbol{\theta}^*$ then $\boldsymbol{\theta}^*$ must satisfy

$$\boldsymbol{\theta}^* = \mathcal{M}(\boldsymbol{\theta}^*).$$

A Taylor series expansion of $\boldsymbol{\theta}^{(w+1)} = \mathcal{M}(\boldsymbol{\theta}^{(w)})$ about $\boldsymbol{\theta}^*$ gives

$$\boldsymbol{\theta}^{(w+1)} - \boldsymbol{\theta}^* \approx \mathcal{J}(\boldsymbol{\theta}^*)(\boldsymbol{\theta}^{(w)} - \boldsymbol{\theta}^*),$$

where $\mathcal{J}(\boldsymbol{\theta}^*)$ is the $d \times d$ Jacobian matrix for $\mathcal{M}(\boldsymbol{\theta}) = \{\mathcal{M}_1(\boldsymbol{\theta}), \dots, \mathcal{M}_d(\boldsymbol{\theta})\}^T$ evaluated at $\boldsymbol{\theta} = \boldsymbol{\theta}^*$, i.e.,

$$\mathcal{J}(\boldsymbol{\theta}^*) = \left\langle \begin{pmatrix} \frac{\partial \mathcal{M}_1(\boldsymbol{\theta})}{\partial \theta_1} & \dots & \frac{\partial \mathcal{M}_1(\boldsymbol{\theta})}{\partial \theta_d} \\ \vdots & \ddots & \vdots \\ \frac{\partial \mathcal{M}_d(\boldsymbol{\theta})}{\partial \theta_1} & \dots & \frac{\partial \mathcal{M}_d(\boldsymbol{\theta})}{\partial \theta_d} \end{pmatrix} \right\rangle_{\boldsymbol{\theta} = \boldsymbol{\theta}^*}.$$

The EM algorithm is a first order algorithm having linear convergence with rate matrix $\mathcal{J}(\boldsymbol{\theta}^*)$. The global rate of convergence is the largest eigenvalue of $\mathcal{J}(\boldsymbol{\theta}^*)$, denoted $\lambda_{\max}$. A small value of $\lambda_{\max}$ implies fast convergence. Dempster et al. (1977) showed that $\mathcal{J}(\boldsymbol{\theta}^*)$ can be expressed as

$$\begin{aligned} \mathcal{J}(\boldsymbol{\theta}^*) &= \mathbf{I}_d - \mathcal{I}_c^{-1}(\boldsymbol{\theta}^*; \mathbf{y}_o) \mathcal{I}_o(\boldsymbol{\theta}^*; \mathbf{y}_o) \\ &= \mathcal{I}_c^{-1}(\boldsymbol{\theta}^*; \mathbf{y}_o) \mathcal{I}_m(\boldsymbol{\theta}^*; \mathbf{y}_o) \end{aligned}$$

i.e., the rate of convergence is determined by the amount of missing information and the greater the proportion of missing information the slower the rate of convergence. Using $\lambda_{\max}$ as a measure of the global rate of convergence allows us to make general

comparisons regarding EM algorithms using different complete data specifications. If we can show that $\lambda_{\max}$ for one specification is always less than that for another specification, then all else being equal we will prefer the specification with smaller $\lambda_{\max}$.

Chapter 2

The linear mixed model as an incomplete data problem and a conditional derivation of REML

Many authors, including Dempster et al. (1977), note that it can be unclear how estimating the parameters of a linear mixed model can be viewed as an incomplete data problem. We present Henderson's mixed model equations as we believe it provides a natural starting point from which to consider an EM algorithm for the linear mixed model. Furthermore, we use the coefficient matrix associated with Henderson's mixed model equations extensively in later chapters. This thesis is concerned with the development of a REML EM algorithm for the linear mixed model. We present the conditional derivation of REML derived by Verbyla (1990). This derivation provides the basis for a new specification of a REML EM algorithm for the linear mixed model considered in a later chapter and presents terminology that will be used throughout the rest of this thesis. We begin this chapter by considering a linear mixed model specification.

2.1 Linear mixed model specification

The linear mixed model can be written as

$$\mathbf{y}_o = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u} + \mathbf{e}, \tag{2.1}$$

CHAPTER 2. THE LINEAR MIXED MODEL AS AN INCOMPLETE DATA PROBLEM AND A CONDITIONAL DERIVATION OF REML

where $\mathbf{y}_o$ is a $n \times 1$ vector of observed data, $\boldsymbol{\beta}$ is a $t \times 1$ vector of fixed effects, $\mathbf{X}$ is a $n \times t$ design matrix of full column rank, $\mathbf{u}$ is a $b \times 1$ vector of random effects, $\mathbf{Z}$ is a $n \times b$ design matrix, and $\mathbf{e}$ is a $n \times 1$ vector of residuals. It is assumed that $\mathbf{u}$ and $\mathbf{e}$ are realisations of a random variable with the joint distribution

$$\begin{pmatrix} \mathbf{u} \\ \mathbf{e} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \sigma_H^2 \begin{pmatrix} \mathbf{G} & \mathbf{0} \\ \mathbf{0} & \mathbf{R} \end{pmatrix} \right\},$$

where $\mathbf{G} = \mathbf{G}(\boldsymbol{\gamma})$, $\mathbf{R} = \sigma^2 \boldsymbol{\Sigma}$, and $\boldsymbol{\Sigma} = \boldsymbol{\Sigma}(\boldsymbol{\phi})$. The vectors $\boldsymbol{\gamma}$ and $\boldsymbol{\phi}$ are vectors of variance parameters associated with random effects and residuals respectively. The variance matrices $\mathbf{G}$, $\mathbf{R}$, and $\boldsymbol{\Sigma}$ are assumed to be symmetric positive definite. The distribution of the data vector $\mathbf{y}_o$ is therefore Gaussian with mean $\mathbf{X}\boldsymbol{\beta}$ and variance $\sigma_H^2 \mathbf{H}$ where $\mathbf{H} = \mathbf{Z}\mathbf{G}\mathbf{Z}^T + \mathbf{R}$. The variance matrix $\mathbf{H}$ is also assumed to be symmetric positive definite.

This specification of the linear mixed model is completely general. However, when σ^2 and σ_H^2 are both included in the model these two parameters are unidentifiable, i.e., they cannot both be estimated. They are both included in the mixed model specification to allow for different parameterisations, namely the parameter vectors $\boldsymbol{\gamma}$ and $\boldsymbol{\phi}$ can be considered as either variance components or variance ratios. The problem of identifiability is overcome by setting at least one of them to 1. In this thesis we will consider $\sigma_H^2 = 1$ so that the elements of $\boldsymbol{\gamma}$ are variance components and when $\boldsymbol{\Sigma}$ is a variance-covariance matrix the elements of $\boldsymbol{\phi}$ are ratios relative to σ^2 . The joint distribution of $\mathbf{y}_o$, $\mathbf{u}$, and $\mathbf{e}$ is

$$\begin{pmatrix} \mathbf{y}_o \\ \mathbf{u} \\ \mathbf{e} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{X}\boldsymbol{\beta} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{H} & \mathbf{Z}\mathbf{G} & \mathbf{R} \\ \mathbf{G}\mathbf{Z}^T & \mathbf{G} & \mathbf{0} \\ \mathbf{R} & \mathbf{0} & \mathbf{R} \end{pmatrix} \right\} \quad (2.2)$$

Fitting the model given in (2.1) involves estimating a vector of variance parameters $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$ and the vector of fixed effects $\boldsymbol{\beta}$.

2.2 Henderson's mixed model equations

Charles Roy Henderson revolutionised the application of linear mixed models in the biological sciences in two important papers in the 1950's. In the first (Henderson, 1953) three methods, known as Method I, II, and III were presented for finding

maximum likelihood estimates of the parameters associated with a linear mixed model. The second (Henderson et al., 1959) presents the now famous mixed model equations. Henderson's mixed model equations are relevant to this thesis in two ways. Firstly, both Henderson's mixed model equations and a ML EM algorithm for linear mixed models are derived by considering the same log-density function. Secondly, the coefficient matrix associated with Henderson's mixed model equations and its inverse are used to compare the computational requirements of alternative REML EM algorithms.

Henderson's mixed model equations are derived by considering the log joint density function of $\mathbf{y}_o$ and $\mathbf{u}$, i.e.

$$\log\{f(\mathbf{y}_o, \mathbf{u})\} = \log\{f(\mathbf{y}_o|\mathbf{u})\} + \log\{f(\mathbf{u})\}.$$

Obtaining the conditional distribution $f(\mathbf{y}_o|\mathbf{u})$ requires a well known result in multivariate normal statistics (Result A.1) and using the joint distribution given in (2.2). The conditional distribution $f(\mathbf{y}_o|\mathbf{u})$ is

$$\mathbf{y}_o|\mathbf{u} \sim N(\mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u}, \mathbf{R}).$$

We define the vector of unknown scale and location parameters as $\boldsymbol{\theta} = (\boldsymbol{\beta}^T, \boldsymbol{\kappa}^T)^T$ where $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$. The log joint density function of $\mathbf{y}_o$ and $\mathbf{u}$ is therefore (excluding constants)

$$\begin{aligned} \log\{f(\mathbf{y}_o, \mathbf{u}; \boldsymbol{\theta})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \log\{\det(\mathbf{G})\} \right. \\ &\quad \left. + (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u})^T \mathbf{R}^{-1} (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u}) + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right]. \end{aligned} \quad (2.3)$$

Maximising (2.3) with respect to $\boldsymbol{\beta}$ and $\mathbf{u}$ and equating to zero gives

$$\begin{aligned} \mathbf{X}^T \mathbf{R}^{-1} \mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{X}^T \mathbf{R}^{-1} \mathbf{Z} \tilde{\mathbf{u}} &= \mathbf{X}^T \mathbf{R}^{-1} \mathbf{y}_o \\ \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{X} \hat{\boldsymbol{\beta}} - (\mathbf{Z}^T \mathbf{R}^{-1} \mathbf{Z} + \mathbf{G}^{-1}) \tilde{\mathbf{u}} &= \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{y}_o, \end{aligned}$$

which are known as Henderson's mixed model equations. They can be written in matrix form as

$$\begin{pmatrix} \mathbf{X}^T \mathbf{R}^{-1} \mathbf{X} & \mathbf{X}^T \mathbf{R}^{-1} \mathbf{Z} \\ \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{X} & \mathbf{G}^{-1} + \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{Z} \end{pmatrix} \begin{pmatrix} \hat{\boldsymbol{\beta}} \\ \tilde{\mathbf{u}} \end{pmatrix} = \begin{pmatrix} \mathbf{X}^T \mathbf{R}^{-1} \mathbf{y}_o \\ \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{y}_o \end{pmatrix},$$

where $\hat{\boldsymbol{\beta}}$ is the best linear unbiased estimator (BLUE) of $\boldsymbol{\beta}$ and $\tilde{\mathbf{u}}$ is the best linear unbiased predictor (BLUP) of $\mathbf{u}$. It is convenient to define the coefficient matrix of Henderson's mixed model equations as

$$\begin{aligned} \mathbf{C} &= \begin{pmatrix} \mathbf{X}^T \mathbf{R}^{-1} \mathbf{X} & \mathbf{X}^T \mathbf{R}^{-1} \mathbf{Z} \\ \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{X} & \mathbf{Z}^T \mathbf{R}^{-1} \mathbf{Z} + \mathbf{G}^{-1} \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{C}_{XX} & \mathbf{C}_{XZ} \\ \mathbf{C}_{ZX} & \mathbf{C}_{ZZ} \end{pmatrix} \end{aligned}$$

and the inverse of the coefficient matrix as

$$\mathbf{C}^{-1} = \begin{pmatrix} \mathbf{C}^{XX} & \mathbf{C}^{XZ} \\ \mathbf{C}^{ZX} & \mathbf{C}^{ZZ} \end{pmatrix}.$$

Using Result A.3.6 it can be shown that

$$\begin{aligned} \mathbf{C}^{XX} &= (\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1}, & \text{using Result A.3.7,} \\ \mathbf{C}^{XZ} &= \mathbf{C}^{ZX^T} = -(\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{H}^{-1} \mathbf{Z} \mathbf{G}, & \text{using Result A.3.7,} \\ \mathbf{C}^{ZZ} &= (\mathbf{Z}^T \mathbf{S} \mathbf{Z} + \mathbf{G}^{-1})^{-1}, & \text{using Result A.3.5,} \end{aligned}$$

where $\mathbf{S} = \mathbf{R}^{-1} - \mathbf{R}^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{R}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}^{-1}$. Throughout this thesis we will define the conditional variances required for computation of the E-step in the EM algorithm for linear mixed models in terms related to the coefficient matrix of Henderson's mixed model equations.

2.3 Missing data specification for the linear mixed model

Dempster et al. (1977) in Section 4.4 of their paper provide an example of an EM algorithm for a linear mixed model. They note that

in the context of linear models...the relevance of incomplete-data concepts is at first sight remote.

However, since Henderson's mixed model equations are based on the joint density function of $\mathbf{y}_o$ and $\mathbf{u}$, a natural choice for the complete data when formulating an EM algorithm for a linear mixed model is $\mathbf{y}_c = (\mathbf{y}_o^T, \mathbf{u}^T)^T$ where $\mathbf{y}_o$ and $\mathbf{u}$ are the incomplete and missing data vectors respectively. Using EM algorithm terminology, equation (2.3) is referred to as the complete data log-density function.

It is well known that maximum likelihood estimates of variance components are biased in small samples. Therefore, we do not present an EM algorithm for a linear mixed model in this thesis. However, it is worth noting that considering the vector of random effects $\mathbf{u}$ as the missing data is the approach taken later when we consider two REML EM algorithms for the linear mixed model. The second of these is a new complete data specification based on the Verbyla (1990) conditional derivation of REML.

2.4 Conditional derivation of REML

The preferred method for estimating the parameters associated with a linear mixed model is residual or restricted maximum likelihood (REML). The original reference for REML is the paper by Patterson and Thompson (1971), however the new complete data specification we consider later is based on the conditional derivation of REML presented in Verbyla (1990). The conditional derivation of REML begins by considering the transformation

$$\mathbf{L}^T \mathbf{y}_o = \begin{pmatrix} \mathbf{L}_1^T \mathbf{y}_o \\ \mathbf{L}_2^T \mathbf{y}_o \end{pmatrix} = \begin{pmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix},$$

where $\mathbf{L} = (\mathbf{L}_1, \mathbf{L}_2)$ is a non-singular matrix, with $\mathbf{L}_1$ and $\mathbf{L}_2$ $n \times t$ and $n \times (n - t)$ matrices respectively, both of full-column rank and chosen to satisfy $\mathbf{L}_1^T \mathbf{X} = \mathbf{I}_t$ and $\mathbf{L}_2^T \mathbf{X} = \mathbf{0}$. It is worth noting that the $n - t \times 1$ vector $\mathbf{y}_2 = \mathbf{L}_2^T \mathbf{y}_o$ will play a key role in the development of a new complete data specification for linear mixed models and factor analytic mixed models. The distribution of the transformed data is

$$\begin{pmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix} \sim N \left\{ \begin{pmatrix} \boldsymbol{\beta} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_1^T \mathbf{H} \mathbf{L}_1 & \mathbf{L}_1^T \mathbf{H} \mathbf{L}_2 \\ \mathbf{L}_2^T \mathbf{H} \mathbf{L}_1 & \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2 \end{pmatrix} \right\}.$$

Following the example of Verbyla (1990) it can be shown that

$$\begin{aligned} \mathbf{L}^T \mathbf{y}_o &\sim N(\mathbf{L}^T \mathbf{X} \boldsymbol{\beta}, \mathbf{L}^T \mathbf{H} \mathbf{L}), \\ \mathbf{y}_1 | \mathbf{y}_2 &\sim N\{\boldsymbol{\beta}, (\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1}\}, \\ \mathbf{y}_2 &\sim N(\mathbf{0}, \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2). \end{aligned}$$

Based on the relation

$$f(\mathbf{L}^T \mathbf{y}_o; \boldsymbol{\theta}) = f(\mathbf{y}_2; \boldsymbol{\kappa}) f(\mathbf{y}_1 | \mathbf{y}_2; \boldsymbol{\theta}),$$

we can equate the determinants, i.e.,

$$\begin{aligned} \det(\mathbf{L}^T \mathbf{H} \mathbf{L}) &= \det\{(\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1}\} \det(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2), \\ \det(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2) &= \det(\mathbf{L}^T \mathbf{L}) \det(\mathbf{H}) \det(\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X}). \end{aligned} \quad (2.4)$$

The identity (2.4) will be useful later in the development of the residual log-likelihood function that is most commonly presented. The log-likelihood function of $\mathbf{L}^T \mathbf{y}_o$ can be expressed as

$$\ell(\boldsymbol{\theta}; \mathbf{L}^T \mathbf{y}_o) = \ell(\boldsymbol{\kappa}; \mathbf{y}_2) + \ell(\boldsymbol{\theta}; \mathbf{y}_1 | \mathbf{y}_2). \quad (2.5)$$

The conditional log-likelihood function $\ell(\boldsymbol{\theta}; \mathbf{y}_1 | \mathbf{y}_2)$ is used to estimate the fixed effects, after which there is no information left for the estimation of variance parameters. Therefore, the log-likelihood function associated with the marginal density of $\mathbf{y}_2$, i.e., $\ell(\boldsymbol{\kappa}; \mathbf{y}_2)$ is used for the REML estimation of variance parameters and can be expressed (excluding constants) as

$$\ell(\boldsymbol{\kappa}; \mathbf{y}_2) = -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)\} + \mathbf{y}_2^T (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{y}_2 \right] \quad (2.6)$$

$$= -\frac{1}{2} \left[\log\{\det(\mathbf{H})\} + \log\{\det(\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})\} + \mathbf{y}_o^T \mathbf{P} \mathbf{y}_o \right], \quad (2.7)$$

where going from (2.6) to (2.7) makes use of (2.4) and Result A.3.8, i.e.,

$$\begin{aligned} \mathbf{P} &= \mathbf{L}_2(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \\ &= \mathbf{H}^{-1} - \mathbf{H}^{-1} \mathbf{X}(\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{H}^{-1}. \end{aligned}$$

The expression given in (2.7) is the form of the residual log-likelihood function most commonly presented.

Although we have concentrated on the conditional derivation of REML presented in Verbyla (1990) it is worth noting that much earlier Kalbfleisch and Sprott (1970) considered factorising a distribution, parameterised by two parameters of interest, into a marginal and conditional distribution in much the same way that the likelihood is factorised in Verbyla (1990).

Chapter 3

Current REML EM algorithms for the linear mixed model

There are two specifications of a REML EM algorithm for the linear mixed model currently in the literature. The first and most widely used is the random effect specification which assumes β is a vector of random effects with variance tending to infinity. The random effect approach was suggested by Dempster et al. (1977) and followed by others including Laird and Ware (1982) and Foulley et al. (2000). The other approach is based on the expectation computed in the E-step being conditional on the transformed observed data vector that is commonly associated with the residual log-likelihood function. This fixed effect approach was taken by Cullis et al. (2004) and Knight (2008). It will be shown that from a practical implementation point of view the two approaches are equivalent.

This chapter proceeds by considering the linear mixed model specification, the complete data specification, the complete data log-density function, and the conditional distributions required in the E-step for both the random effect and fixed effect approaches. We show that the complete data log-density function and the conditional distributions required are exactly the same for both approaches. This has not been shown in the literature previously. This chapter ends by considering the E and M-steps for the fixed effect approach in more detail and the complete and missing data information matrices associated with this approach.

3.1 Linear mixed model specification - random effect approach

The linear mixed model is specified in the same way as in Section 2.1 except that β is considered to be a $t \times 1$ vector of random effects. We define the joint distribution of $\mathbf{y}_o$, β , $\mathbf{u}$, and $\mathbf{e}$ as

$$\begin{pmatrix} \mathbf{y}_o \\ \beta \\ \mathbf{u} \\ \mathbf{e} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{F} & \mathbf{XB} & \mathbf{ZG} & \mathbf{R} \\ \mathbf{BX}^T & \mathbf{B} & \mathbf{0} & \mathbf{0} \\ \mathbf{GZ}^T & \mathbf{0} & \mathbf{G} & \mathbf{0} \\ \mathbf{R} & \mathbf{0} & \mathbf{0} & \mathbf{R} \end{pmatrix} \right\}, \quad (3.1)$$

where

$$\begin{aligned} \mathbf{F} &= \mathbf{XBX}^T + \mathbf{ZGZ}^T + \mathbf{R} \\ &= \mathbf{XBX}^T + \mathbf{H} \end{aligned}$$

and

$$\mathbf{F}^{-1} = \mathbf{H}^{-1} - \mathbf{H}^{-1}\mathbf{X}(\mathbf{B}^{-1} + \mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-1}, \quad \text{using Result A.3.5.}$$

The random effect approach assumes that the variance of β tends to infinity, i.e., $\mathbf{B}^{-1} \rightarrow \mathbf{0}$. Therefore

$$\begin{aligned} \mathbf{F}^{-1} &\rightarrow \mathbf{H}^{-1} - \mathbf{H}^{-1}\mathbf{X}(\mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-1}, \\ \mathbf{F}^{-1} &\rightarrow \mathbf{P}. \end{aligned} \quad (3.2)$$

From here on we will assume that if $\mathbf{B}^{-1} \rightarrow \mathbf{0}$ then $\mathbf{F}^{-1} = \mathbf{P}$.

3.2 Complete data specification - random effect approach

Under the random effect approach the complete data is defined as $\mathbf{y}_c = (\mathbf{y}_o^T, \boldsymbol{\beta}^T, \mathbf{u}^T)^T$. We define $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$ and note that the complete data density function can be expressed as

$$f(\mathbf{y}_c; \boldsymbol{\kappa}) = f(\mathbf{y}_o | \boldsymbol{\beta}, \mathbf{u}; \sigma^2, \boldsymbol{\phi}) f(\boldsymbol{\beta}, \mathbf{u}; \boldsymbol{\gamma}).$$

Since it is assumed $\boldsymbol{\beta}$ and $\mathbf{u}$ are independent and that $\mathbf{B}^{-1} \rightarrow \mathbf{0}$, then

$$f(\boldsymbol{\beta}, \mathbf{u}; \boldsymbol{\gamma}) \propto f(\mathbf{u}; \boldsymbol{\gamma}).$$

Using (3.1) and Result A.1 it can be shown that

$$\mathbf{y}_o | \boldsymbol{\beta}, \mathbf{u} \sim N(\mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u}, \mathbf{R}).$$

The complete data log-density function is therefore

$$\begin{aligned} \log\{f(\mathbf{y}_c; \boldsymbol{\kappa})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \log\{\det(\mathbf{G})\} + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right. \\ &\quad \left. + (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u})^T \mathbf{R}^{-1} (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u}) \right] \\ &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \log\{\det(\mathbf{G})\} \right. \\ &\quad \left. + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} + \mathbf{e}^T \mathbf{R}^{-1} \mathbf{e} \right], \end{aligned} \tag{3.3}$$

which is equivalent to the log joint density function of $\mathbf{y}_o$ and $\mathbf{u}$ used in deriving Henderson's mixed model equations. The E-step of a REML EM algorithm using the complete data specification $\mathbf{y}_c = (\mathbf{y}_o^T, \boldsymbol{\beta}^T, \mathbf{u}^T)^T$ requires the conditional distributions $\mathbf{u} | \mathbf{y}_o$ and $\mathbf{e} | \mathbf{y}_o$. Using (3.1) and Result A.1 we have

$$\begin{aligned}
 \mathbf{u}|\mathbf{y}_o &\sim N(\mathbf{GZ}^T\mathbf{F}^{-1}\mathbf{y}_o, \mathbf{G} - \mathbf{GZ}^T\mathbf{F}^{-1}\mathbf{ZG}), \\
 &\sim N(\mathbf{GZ}^T\mathbf{P}\mathbf{y}_o, \mathbf{G} - \mathbf{GZ}^T\mathbf{PZG}) \quad \text{using (3.2),} \\
 &\sim N(\mathbf{GZ}^T\mathbf{P}\mathbf{y}_o, \mathbf{C}^{ZZ}) \quad \text{using Result A.3.10}
 \end{aligned} \tag{3.4}$$

and

$$\begin{aligned}
 \mathbf{e}|\mathbf{y}_o &\sim N(\mathbf{RF}^{-1}\mathbf{y}_o, \mathbf{R} - \mathbf{RF}^{-1}\mathbf{R}), \\
 &\sim N(\mathbf{RP}\mathbf{y}_o, \mathbf{R} - \mathbf{RPR}) \quad \text{using (3.2),} \\
 &\sim N(\mathbf{RP}\mathbf{y}_o, \mathbf{WC}^{-1}\mathbf{W}^T) \quad \text{using Result A.3.9,}
 \end{aligned} \tag{3.5}$$

where $\mathbf{W} = (\mathbf{X} \ \mathbf{Z})$. An important feature of (3.4) and (3.5) is that the conditional variances associated with $\mathbf{u}|\mathbf{y}_o$ and $\mathbf{e}|\mathbf{y}_o$ are $b \times b$ and $n \times n$ matrices respectively. Therefore, the conditional variance of $\mathbf{e}|\mathbf{y}_o$ can require significantly more computation time than the conditional variance of $\mathbf{u}|\mathbf{y}_o$.

3.3 Linear mixed model specification - fixed effect approach

The linear mixed model is specified in the same way as in Section 2.1. Under the fixed effect approach we require the joint distribution of $\mathbf{y}_o$, $\mathbf{u}$, and $\mathbf{e}$ given in (2.2) to form the complete data log-density function. To form the conditional distributions necessary for computation of the E-step we require the joint distribution of $\mathbf{y}_2$, $\mathbf{u}$, and $\mathbf{e}$ which can be shown to be

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{u} \\ \mathbf{e} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2 & \mathbf{L}_2^T \mathbf{Z} \mathbf{G} & \mathbf{L}_2^T \mathbf{R} \\ \mathbf{G} \mathbf{Z}^T \mathbf{L}_2 & \mathbf{G} & \mathbf{0} \\ \mathbf{R} \mathbf{L}_2 & \mathbf{0} & \mathbf{R} \end{pmatrix} \right\} \tag{3.6}$$

3.4 Complete data specification - fixed effect approach

Under the fixed effect approach the complete data is defined as $\mathbf{y}_c = (\mathbf{y}_o^T, \mathbf{u}^T)^T$. We define $\boldsymbol{\theta} = (\boldsymbol{\beta}^T, \boldsymbol{\kappa}^T)^T$ and $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$ and note that the complete data density function can be expressed as

$$f(\mathbf{y}_c; \boldsymbol{\theta}) = f(\mathbf{y}_o | \mathbf{u}; \boldsymbol{\theta}) f(\mathbf{u}; \boldsymbol{\gamma}).$$

Using (2.2) and Result A.1 it can be shown that

$$\mathbf{y}_o | \mathbf{u} \sim N(\mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u}, \mathbf{R}). \quad (3.7)$$

The complete data log-density function is therefore

$$\begin{aligned} \log\{f(\mathbf{y}_c; \boldsymbol{\theta})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \log\{\det(\mathbf{G})\} + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right. \\ &\quad \left. + (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u})^T \mathbf{R}^{-1} (\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u}) \right] \\ &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \log\{\det(\mathbf{G})\} \right. \\ &\quad \left. + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} + \mathbf{e}^T \mathbf{R}^{-1} \mathbf{e} \right], \end{aligned} \quad (3.8)$$

which is exactly the same as the random effect complete data specification in (3.3). Note that the difference between the parameterisation of the complete data log-density under the two approaches is due to the different treatment of $\boldsymbol{\beta}$. Under the random effect approach, the complete data log-likelihood function is parameterised by $\boldsymbol{\kappa}$ and under the fixed effect approach by $\boldsymbol{\theta}$, including when we make the substitution $\mathbf{e} = \mathbf{y}_o - \mathbf{X}\boldsymbol{\beta} - \mathbf{Z}\mathbf{u}$.

The E-step of the REML EM algorithm using the fixed effect approach requires the conditional distributions $\mathbf{u} | \mathbf{y}_2$ and $\mathbf{e} | \mathbf{y}_2$. Using (3.6) and Result A.1 we have

$$\begin{aligned}
\mathbf{u}|\mathbf{y}_2 &\sim N \left\{ \mathbf{GZ}^T \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \mathbf{y}_o, \mathbf{G} - \mathbf{GZ}^T \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \mathbf{ZG} \right\}, \\
&\sim N (\mathbf{GZ}^T \mathbf{P} \mathbf{y}_o, \mathbf{G} - \mathbf{GZ}^T \mathbf{PZG}) && \text{using Result A.3.8,} \\
&\sim N (\mathbf{GZ}^T \mathbf{P} \mathbf{y}_o, \mathbf{C}^{ZZ}) && \text{using Result A.3.10}
\end{aligned} \tag{3.9}$$

and

$$\begin{aligned}
\mathbf{e}|\mathbf{y}_2 &\sim N \left\{ \mathbf{RL}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \mathbf{y}_o, \mathbf{R} - \mathbf{RL}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \mathbf{R} \right\}, \\
&\sim N (\mathbf{RP} \mathbf{y}_o, \mathbf{R} - \mathbf{RPR}) && \text{using Result A.3.8,} \\
&\sim N (\mathbf{RP} \mathbf{y}_o, \mathbf{WC}^{-1} \mathbf{W}^T) && \text{using Result A.3.9.}
\end{aligned} \tag{3.10}$$

Note that (3.9) and (3.10) are exactly the same as (3.4) and (3.5) respectively. Since the complete data specification and the conditional distributions necessary for computation of the E-step are exactly the same for the two approaches the random effect and fixed effect approach are equivalent from a practical implementation point of view. From here on we will only consider the fixed effect approach.

3.5 REML EM algorithm E-step

We use the superscript (P) to define $\mathbf{y}_c^{(P)} = (\mathbf{y}_o^T, \mathbf{u}^T)^T$ to distinguish this complete data specification and its associated $Q(\cdot)$ function from the new complete data specification considered later. The log-density function associated with $\mathbf{y}_c^{(P)}$ can be partitioned as

$$\log\{f(\mathbf{y}_c^{(P)}; \boldsymbol{\theta})\} = \log\{f(\mathbf{e}; \boldsymbol{\theta})\} + \log\{f(\mathbf{u}; \boldsymbol{\gamma})\}, \tag{3.11}$$

where

$$\begin{aligned}
\log\{f(\mathbf{e}; \boldsymbol{\theta})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \mathbf{e}^T \mathbf{R}^{-1} \mathbf{e} \right], \\
\log\{f(\mathbf{u}; \boldsymbol{\gamma})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{G})\} + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right].
\end{aligned}$$

The E-step involves taking the expectation of (3.11) over $k(\mathbf{u}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)})$ where $\boldsymbol{\kappa}^{(w)}$ is the w -th iterate of $\boldsymbol{\kappa}$. We can write

$$\begin{aligned} Q^{(P)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)}) &= E[\log\{f(\mathbf{y}_c^{(P)}; \boldsymbol{\theta})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= E[\log\{f(\mathbf{e}; \boldsymbol{\theta})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] + E[\log\{f(\mathbf{u}; \boldsymbol{\gamma})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)}) + Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)}). \end{aligned}$$

Upon taking the expectation of $\log\{f(\mathbf{y}_c^{(P)}; \boldsymbol{\theta})\}$ over $k(\mathbf{u}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)})$ the associated $Q(\cdot)$ function is now parameterised by $\boldsymbol{\kappa}$ (and $\boldsymbol{\kappa}^{(w)}$). Using Result A.3.2 we can write

$$\begin{aligned} Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{R})\} + \tilde{\mathbf{e}}^{(w)T} \mathbf{R}^{-1} \tilde{\mathbf{e}}^{(w)} + \text{tr}(\mathbf{R}^{-1} \mathbf{W} \mathbf{C}^{-1(w)} \mathbf{W}^T) \right], \\ Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{G})\} + \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1} \tilde{\mathbf{u}}^{(w)} + \text{tr}(\mathbf{G}^{-1} \mathbf{C}^{ZZ(w)}) \right], \end{aligned}$$

where

$$\tilde{\mathbf{e}}^{(w)} = E(\mathbf{e}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}) = \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o, \quad (3.12)$$

$$\tilde{\mathbf{u}}^{(w)} = E(\mathbf{u}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}) = \mathbf{G}^{(w)} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o. \quad (3.13)$$

This completes the E-step.

3.6 REML EM algorithm M-step

The M-step involves maximising $Q^{(P)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ with respect to $\boldsymbol{\kappa}$ to find an update for $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$. The partitioning of $Q^{(P)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ into $Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})$ and $Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)})$ simplifies the process. We have not provided an updating equation for $\boldsymbol{\beta}$ as the maximum likelihood estimate for $\boldsymbol{\beta}$ can be obtained at the end of the REML EM algorithm by using the generalised least squares estimate of $\boldsymbol{\beta}$. We use the notation in (3.12) and (3.13) when considering the updating equations for σ^2 , γ_{ij} and ϕ_{gh} .

3.6.1 Update for σ^2

Noting that $\mathbf{R} = \sigma^2 \mathbf{\Sigma}$ and differentiating $Q^{(P)}(\sigma^2, \phi; \kappa^{(w)})$ with respect to σ^2 we have

$$\frac{\partial Q^{(P)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \sigma^2} = -\frac{1}{2} \left\{ \frac{n}{\sigma^2} - \frac{1}{(\sigma^2)^2} \tilde{\mathbf{e}}^{(w)T} \mathbf{\Sigma}^{-1} \tilde{\mathbf{e}}^{(w)} - \frac{1}{(\sigma^2)^2} \text{tr}(\mathbf{\Sigma}^{-1} \mathbf{W} \mathbf{C}^{-1(w)} \mathbf{W}^T) \right\}.$$

Equating the derivative to zero we have

$$\sigma^{2(w+1)} = \frac{1}{n} \left\{ \tilde{\mathbf{e}}^{(w)T} \mathbf{\Sigma}^{-1(w)} \tilde{\mathbf{e}}^{(w)} + \text{tr}(\mathbf{\Sigma}^{-1(w)} \mathbf{W} \mathbf{C}^{-1(w)} \mathbf{W}^T) \right\}. \quad (3.14)$$

3.6.2 Update for γ_{gh}

To find an update for γ_{gh} , $g = 1, \dots, b$, $h = 1, \dots, b$, $i \geq j$ involves differentiating $Q^{(P)}(\gamma; \kappa^{(w)})$ with respect to γ_{gh} . Using Result A.4.1 and Result A.4.2 we have

$$\frac{\partial Q^{(P)}(\gamma; \kappa^{(w)})}{\partial \gamma_{gh}} = -\frac{1}{2} \left\{ \text{tr} \left(\mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_{gh}} \right) - \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_{gh}} \mathbf{G}^{-1} \tilde{\mathbf{u}}^{(w)} - \text{tr} \left(\mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_{gh}} \mathbf{G}^{-1} \mathbf{C}^{ZZ(w)} \right) \right\}.$$

The updating formula for γ_{gh} depends on the specific form of $\mathbf{G}$. For a one-way random effects model $\mathbf{G}$ has the simple form $\mathbf{G} = \sigma_u^2 \mathbf{I}_b$. In this case the updating equation is

$$\sigma_u^{2(w+1)} = \frac{1}{b} \left\{ \tilde{\mathbf{u}}^{(w)T} \tilde{\mathbf{u}}^{(w)} + \text{tr}(\mathbf{C}^{ZZ(w)}) \right\}.$$

3.6.3 Update for ϕ_{ij}

To find an update for ϕ_{ij} , $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ involves differentiating $Q^{(P)}(\sigma^2, \phi; \kappa^{(w)})$ with respect to ϕ_{ij} . Noting that $\mathbf{R} = \sigma^2 \mathbf{\Sigma}$ and using Result A.4.1 and Result A.4.2 we have

$$\begin{aligned} \frac{\partial Q^{(P)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \phi_{ij}} &= -\frac{1}{2} \left\{ \text{tr} \left(\Sigma^{-1} \frac{\partial \Sigma}{\partial \phi_{ij}} \right) - \frac{1}{\sigma^2} \tilde{\mathbf{e}}^{(w)T} \Sigma^{-1} \frac{\partial \Sigma}{\partial \phi_{ij}} \Sigma^{-1} \tilde{\mathbf{e}}^{(w)} \right. \\ &\quad \left. - \frac{1}{\sigma^2} \text{tr} \left(\Sigma^{-1} \frac{\partial \Sigma}{\partial \phi_{ij}} \Sigma^{-1} \mathbf{W} \mathbf{C}^{-1(w)} \mathbf{W}^T \right) \right\}. \end{aligned}$$

The updating equation for ϕ_{ij} depends on the specific form of Σ . In this thesis we only consider models where $\Sigma = \mathbf{I}_n$.

This completes the M-step.

3.7 Complete data information matrix

We use the superscript (P) when considering the complete and missing data information matrices associated with $\mathbf{y}_c^{(P)}$. This is done to distinguish these information matrices from the complete and missing data information matrices associated with the new complete data specification we consider later. The complete and missing data information matrices are symmetric and in the interests of parsimony we only present the upper right triangle for each.

The complete data information matrix based on $\mathbf{y}_c^{(P)}$ is

$$\mathcal{I}_c^{(P)}(\kappa^{(w)}; \mathbf{y}_c^{(P)}) = \begin{pmatrix} \mathcal{I}_{c(\sigma^2)^2}^{(P)} & \mathbf{0} & \mathcal{I}_{c\sigma^2\phi}^{(P)} \\ & \mathcal{I}_{c\gamma,\gamma}^{(P)} & \mathbf{0} \\ & & \mathcal{I}_{c\phi,\phi}^{(P)} \end{pmatrix},$$

where the zeroes in the complete data information matrix are due to the partitioning of $Q^{(P)}(\kappa; \kappa^{(w)})$ into $Q^{(P)}(\sigma^2, \phi; \kappa^{(w)})$ and $Q^{(P)}(\gamma; \kappa^{(w)})$. Using the Oakes (1999) result presented in (1.13) we have

$$\begin{aligned}
 \mathcal{I}_{c_{(\sigma^2)^2}}^{(P)} &= - \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial (\sigma^2)^2} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\
 \mathcal{I}_{c_{\gamma, \gamma}}^{(P)} &= - \left\langle \frac{\partial^2 Q^{(P)}(\gamma; \boldsymbol{\kappa}^{(w)})}{\partial \gamma \partial \gamma^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\
 \mathcal{I}_{c_{\phi, \phi}}^{(P)} &= - \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial \phi \partial \phi^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\
 \mathcal{I}_{c_{\sigma^2, \phi}}^{(P)} &= - \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial \sigma^2 \partial \phi^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}.
 \end{aligned}$$

We present each in turn. The linear algebra involved is straightforward but tedious and therefore is omitted.

3.7.1 Element $\mathcal{I}_{c_{(\sigma^2)^2}}^{(P)}$

It can be shown that

$$\mathcal{I}_{c_{(\sigma^2)^2}}^{(P)} = \frac{n}{2(\sigma^{2(w)})^2} + \frac{1}{(\sigma^{2(w)})^2} \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o - \frac{1}{(\sigma^{2(w)})^2} \text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)}). \quad (3.15)$$

3.7.2 Elements of $\mathcal{I}_{c_{\gamma, \gamma}}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$, $g \geq h$ and $q = 1, \dots, b$, $r = 1, \dots, b$, $q \geq r$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{c_{\gamma_{gh}, \gamma_{qr}}}^{(P)} = & \frac{1}{2} \left\{ \text{tr} \left(\mathbf{G}^{-1(w)} \frac{\partial^2 \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)} \partial \gamma_{qr}^{(w)}} \right) - \text{tr} \left(\mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \right) \right. \\
& + \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \tilde{\mathbf{u}}^{(w)} \\
& + \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} \tilde{\mathbf{u}}^{(w)} \\
& - \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1(w)} \frac{\partial^2 \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)} \partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \tilde{\mathbf{u}}^{(w)} - \text{tr} \left(\mathbf{G}^{-1(w)} \frac{\partial^2 \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)} \partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \mathbf{C}^{ZZ(w)} \right) \\
& + \text{tr} \left(\mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \mathbf{C}^{ZZ(w)} \right) \\
& \left. + \text{tr} \left(\mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} \mathbf{C}^{ZZ(w)} \right) \right\}. \tag{3.16}
\end{aligned}$$

3.7.3 Elements of $\mathcal{I}_{c_{\phi, \phi}}^{(P)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ and $k = 1, \dots, n$, $l = 1, \dots, n$, $k \geq l$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{c_{\phi_{ij}, \phi_{kl}}}^{(P)} = & \frac{1}{2} \left\{ (\sigma^{2(w)})^2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
& + (\sigma^{2(w)})^2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
& - \sigma^{2(w)} \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial^2 \Sigma^{(w)}}{\partial \phi_{ij}^{(w)} \partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o + \sigma^{2(w)} \text{tr} \left(\frac{\partial^2 \Sigma^{(w)}}{\partial \phi_{ij}^{(w)} \partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right) \\
& + (\sigma^{2(w)})^2 \text{tr} \left\{ \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{R}^{-1(w)} - \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right)^T \left(\mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} - \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\} \\
& \left. - (\sigma^{2(w)})^2 \text{tr} \left(\mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\}. \tag{3.17}
\end{aligned}$$

3.7.4 Elements of $\mathcal{I}_{c_{\sigma^2, \phi}}^{(P)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ it can be shown that

$$\begin{aligned} \mathcal{I}_{\sigma^2(w)\phi_{ij}^{(w)}}^{(P)} &= \frac{1}{2} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ &\quad \left. + \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{R}^{-1(w)} \right) - \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \right) \right\}. \end{aligned} \quad (3.18)$$

This concludes the section on the complete data information matrix.

3.8 Missing data information matrix

The missing data information matrix based on $\mathbf{y}_c^{(P)}$ and using the Oakes (1999) result presented in (1.13) is

$$\mathcal{I}_m^{(P)}(\boldsymbol{\kappa}^{(w)}; \mathbf{y}_c^{(P)}) = \begin{pmatrix} \mathcal{I}_{m_{(\sigma^2)^2}}^{(P)} & \mathcal{I}_{m_{\sigma^2, \gamma}}^{(P)} & \mathcal{I}_{m_{\sigma^2, \phi}}^{(P)} \\ & \mathcal{I}_{m_{\gamma, \gamma}}^{(P)} & \mathcal{I}_{m_{\sigma^2, \gamma}}^{(P)} \\ & & \mathcal{I}_{m_{\phi, \phi}}^{(P)} \end{pmatrix},$$

where

$$\begin{aligned} \mathcal{I}_{m_{(\sigma^2)^2}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})}{\partial \sigma^2 \partial \sigma^2} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\ \mathcal{I}_{m_{\gamma, \gamma}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)})}{\partial \boldsymbol{\gamma} \partial \boldsymbol{\gamma}^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\ \mathcal{I}_{m_{\phi, \phi}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})}{\partial \boldsymbol{\phi} \partial \boldsymbol{\phi}^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\ \mathcal{I}_{m_{\sigma^2, \gamma}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})}{\partial \sigma^2 \partial \boldsymbol{\gamma}^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\ \mathcal{I}_{m_{\sigma^2, \phi}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})}{\partial \sigma^2 \partial \boldsymbol{\phi}^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}, \\ \mathcal{I}_{m_{\gamma, \phi}}^{(P)} &= \left\langle \frac{\partial^2 Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)})}{\partial \boldsymbol{\gamma} \partial \boldsymbol{\phi}^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}. \end{aligned}$$

We make use of Result A.4.3 and present each in turn.

3.8.1 Element $\mathcal{I}_{m_{(\sigma^2)^2}}^{(P)}$

It can be shown that

$$\begin{aligned} \mathcal{I}_{m_{(\sigma^2)^2}}^{(P)} = & -\frac{1}{2(\sigma^{2(w)})^2} \left\{ -n - 2\mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ & + 2\mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \\ & \left. + 2\text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)}) - \text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{R}^{(w)}) \right\}. \end{aligned} \quad (3.19)$$

3.8.2 Elements of $\mathcal{I}_{m_{\gamma, \gamma}}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$, $g \geq h$ and $q = 1, \dots, b$, $r = 1, \dots, b$, $q \geq r$ it can be shown that

$$\begin{aligned} \mathcal{I}_{m_{\gamma gh, \gamma qr}}^{(P)} = & -\frac{1}{2} \left[\mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ & + \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \\ & - \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \\ & - \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \\ & \left. - \text{tr} \left\{ \left(\frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{G}^{-1(w)} - \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{Z} \right)^T \left(\mathbf{G}^{-1(w)} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} - \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{qr}^{(w)}} \right) \right\} \right]. \end{aligned} \quad (3.20)$$

3.8.3 Elements of $\mathcal{I}_{m_{\phi, \phi}}^{(P)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ and $k = 1, \dots, n$, $l = 1, \dots, n$, $k \geq l$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\phi_{ij}, \phi_{kl}}}^{(P)} &= -\frac{(\sigma^{2(w)})^2}{2} \left[\mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
&\quad + \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
&\quad - \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
&\quad - \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
&\quad \left. - \text{tr} \left\{ \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{R}^{-1(w)} - \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right)^T \left(\mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} - \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\} \right]. \tag{3.21}
\end{aligned}$$

3.8.4 Elements of $\mathcal{I}_{m_{\sigma^2, \gamma}}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$, $g \geq h$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\sigma^2, \gamma_{gh}}}^{(P)} &= -\frac{1}{2\sigma^{2(w)}} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
&\quad + \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \\
&\quad \left. - \text{tr} \left(\mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \right) \right\}. \tag{3.22}
\end{aligned}$$

3.8.5 Elements of $\mathcal{I}_{m_{\sigma^2, \phi}}^{(P)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\sigma^2, \phi_{ij}}}^{(P)} = & -\frac{1}{2} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
& + \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o - 2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
& - \text{tr} \left(\mathbf{R}^{-1(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) + 2 \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \right) \\
& \left. - \text{tr} \left(\mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \right) \right\}. \tag{3.23}
\end{aligned}$$

3.8.6 Elements of $\mathcal{I}_{m_{\gamma, \phi}}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$, $g \geq h$ and $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\gamma_{gh}, \phi_{ij}}}^{(P)} = & -\frac{\sigma^{2(w)}}{2} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
& + \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
& \left. + \text{tr} \left(\mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \right) \right\}. \tag{3.24}
\end{aligned}$$

This concludes the section on the missing data information matrix.

An example data set where a REML EM algorithm based on the complete data specification $\mathbf{y}_c^{(P)} = (\mathbf{y}_o^T, \mathbf{u}^T)^T$ is considered in the next chapter where the complete data specification $\mathbf{y}_c^{(P)}$ is compared to a new complete data specification.

Chapter 4

A new REML EM algorithm for linear mixed models

In this chapter we present a new REML EM algorithm for linear mixed models based on a new specification of the complete data. This new complete data specification is based upon the conditional derivation of REML presented by Verbyla (1990). This specification is computationally more efficient than the complete data specification considered in the previous chapter as the E-step under the new specification involves computing the trace of two matrices each of order b compared to a matrix of order b and a matrix of order n under the previously published complete data specification. Furthermore, we provide the conditions under which the new specification has a faster rate of convergence than the previously published specification. This chapter concludes by comparing two REML EM algorithms applied to an example data set. One algorithm is based on the new complete data specification and the other on the previously published specification.

4.1 New complete data specification

We consider a linear mixed model specified in the same way as in Section 2.1. We define a new complete data specification $\mathbf{y}_c^{(U)} = (\mathbf{y}_2^T, \mathbf{u}^T)^T$. This differs from the previously published specification $\mathbf{y}_c^{(P)} = (\mathbf{y}_o^T, \mathbf{u}^T)^T$ by considering the incomplete data to be the transformed observed data vector associated with the marginal log-density function in (2.5) which is used for the REML estimation of variance parameters. There are a number of advantages in using the complete data specification

$\mathbf{y}_c^{(U)}$. Firstly, to form the complete data log-density function only requires the joint distribution of $\mathbf{y}_2$ and $\mathbf{u}$ which can be ascertained from (3.6) but is reproduced here for convenience

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{u} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2 & \mathbf{L}_2^T \mathbf{Z} \mathbf{G} \\ \mathbf{G} \mathbf{Z}^T \mathbf{L}_2 & \mathbf{G} \end{pmatrix} \right\} \quad (4.1)$$

Using (2.2) and Result A.1 or by simple modification of the expectation and variance associated with the distribution of $\mathbf{y}_o | \mathbf{u}$ presented in (3.7) it can be shown that

$$\mathbf{y}_2 | \mathbf{u} \sim N(\mathbf{L}_2^T \mathbf{Z} \mathbf{u}, \mathbf{L}_2^T \mathbf{R} \mathbf{L}_2).$$

The complete data log density function is therefore (excluding constants)

$$\begin{aligned} \log\{f(\mathbf{y}_c^{(U)}; \boldsymbol{\kappa})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} + \log\{\det(\mathbf{G})\} \right. \\ &\quad \left. + (\mathbf{y}_o - \mathbf{Z} \mathbf{u})^T \mathbf{S} (\mathbf{y}_o - \mathbf{Z} \mathbf{u}) + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right], \end{aligned} \quad (4.2)$$

where

$$\begin{aligned} \mathbf{S} &= \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \\ &= \mathbf{R}^{-1} - \mathbf{R}^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{R}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}^{-1}. \end{aligned} \quad \text{using Result A.3.8}$$

The relationship between the complete data log density functions based on $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ can be shown to be

$$\log\{f(\mathbf{y}_c^{(P)}; \boldsymbol{\theta})\} = \log\{f(\mathbf{y}_c^{(U)}; \boldsymbol{\kappa})\} + \log\{f(\mathbf{y}_1 | \mathbf{y}_2; \boldsymbol{\theta})\},$$

which is similar to the relation presented in (2.5) based on the Verbyla (1990) conditional derivation of REML, i.e.,

$$\ell(\boldsymbol{\theta}; \mathbf{L}^T \mathbf{y}_o) = \ell(\boldsymbol{\kappa}; \mathbf{y}_2) + \ell(\boldsymbol{\theta}; \mathbf{y}_1 | \mathbf{y}_2).$$

This suggests that using the complete data specification $\mathbf{y}_c^{(P)}$ would be similar to using $\ell(\boldsymbol{\theta}; \mathbf{L}^T \mathbf{y}_o)$ for the REML estimation of variance parameters. While it is

possible to use both for the REML estimation of variance parameters the task is simplified when considering the complete data specification $\mathbf{y}_c^{(U)}$ or using the residual log-likelihood function $\ell(\boldsymbol{\kappa}; \mathbf{y}_2)$. The advantages of using $\mathbf{y}_c^{(U)}$ compared to $\mathbf{y}_c^{(P)}$ are more apparent when considering the E-step of a REML EM algorithm based on the complete data specification $\mathbf{y}_c^{(U)}$.

4.2 REML EM algorithm E-step

The complete data log-density function for $\mathbf{y}_c^{(U)}$ can be partitioned as

$$\log\{f(\mathbf{y}_c^{(U)}; \boldsymbol{\kappa})\} = \log\{f(\mathbf{L}_2^T \mathbf{e}; \sigma^2, \boldsymbol{\phi})\} + \log\{f(\mathbf{u}; \boldsymbol{\gamma})\},$$

where

$$\begin{aligned} \log\{f(\mathbf{L}_2^T \mathbf{e}; \sigma^2, \boldsymbol{\phi})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} + (\mathbf{y}_o - \mathbf{Z}\mathbf{u})^T \mathbf{S}(\mathbf{y}_o - \mathbf{Z}\mathbf{u}) \right], \\ \log\{f(\mathbf{u}; \boldsymbol{\gamma})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{G})\} + \mathbf{u}^T \mathbf{G}^{-1} \mathbf{u} \right]. \end{aligned}$$

The E-step involves taking the conditional expectation of (4.2) over $k(\mathbf{u}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)})$ where $\boldsymbol{\kappa}^{(w)}$ is the w -th iterate of $\boldsymbol{\kappa}$. We can write

$$\begin{aligned} Q^{(U)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)}) &= E[\log\{f(\mathbf{y}_c^{(U)}; \boldsymbol{\kappa})\} | \mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= E[\log\{f(\mathbf{L}_2^T \mathbf{e}; \sigma^2, \boldsymbol{\phi})\} | \mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] + E[\log\{f(\mathbf{u}; \boldsymbol{\gamma})\} | \mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= Q^{(U)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)}) + Q^{(U)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)}). \end{aligned}$$

Using Result A.3.2 we have

$$\begin{aligned} Q^{(U)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} + (\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)})^T \mathbf{S}(\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)}) \right. \\ &\quad \left. + \text{tr}(\mathbf{Z}^T \mathbf{S} \mathbf{Z} \mathbf{C}^{ZZ^{(w)}}) \right], \end{aligned} \tag{4.3}$$

$$\begin{aligned} Q^{(U)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{G})\} + \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1} \tilde{\mathbf{u}}^{(w)} + \text{tr}(\mathbf{G}^{-1} \mathbf{C}^{ZZ^{(w)}}) \right], \end{aligned} \tag{4.4}$$

where

$$\tilde{\mathbf{u}}^{(w)} = E(\mathbf{u}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}) = \mathbf{G}^{(w)} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o.$$

It is worth noting that $Q^{(U)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)})$ is equivalent to $Q^{(P)}(\boldsymbol{\gamma}; \boldsymbol{\kappa}^{(w)})$.

A feature of $Q^{(U)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ is that it requires knowledge of only one conditional distribution, namely $\mathbf{u}|\mathbf{y}_2$. This differs from the $Q(\cdot)$ function associated with the complete data specification $\mathbf{y}_c^{(P)}$ which requires knowledge of both $\mathbf{u}|\mathbf{y}_2$ and $\mathbf{e}|\mathbf{y}_2$. This difference is notable because the $Q(\cdot)$ function associated with $\mathbf{y}_c^{(U)}$ involves computing the traces of two matrices both of order b whereas the $Q(\cdot)$ function associated with $\mathbf{y}_c^{(P)}$ also involves computing the trace of two matrices, one of order b but the other of order n . Depending on the linear mixed model being fitted the computing time associated with $\mathbf{y}_c^{(U)}$ can be significantly less than that associated with $\mathbf{y}_c^{(P)}$.

Although we do not consider in this thesis models where the E-step is intractable and where a stochastic variant of a REML EM algorithm is used to approximate the $Q(\cdot)$ function, the use of the complete data specification $\mathbf{y}_c^{(U)}$ in these cases might be of some interest as approximating the $Q(\cdot)$ function would only require drawing samples from one conditional distribution rather than the two required when using $\mathbf{y}_c^{(P)}$.

4.3 REML EM algorithm M-step

We consider the updating equations for σ^2 , γ_{kl} , and ϕ_{ij} in turn.

4.3.1 Update for σ^2

Noting that $\mathbf{R} = \sigma^2 \boldsymbol{\Sigma}$ and differentiating $Q^{(U)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})$ with respect to σ^2 we have

$$\begin{aligned} \frac{\partial Q^{(U)}(\sigma^2, \boldsymbol{\phi}; \boldsymbol{\kappa}^{(w)})}{\partial \sigma^2} &= -\frac{1}{2} \left\{ \frac{n-t}{\sigma^2} - \frac{1}{(\sigma^2)^2} (\mathbf{y}_o - \mathbf{Z} \tilde{\mathbf{u}}^{(w)})^T \mathbf{U} (\mathbf{y}_o - \mathbf{Z} \tilde{\mathbf{u}}^{(w)}) \right. \\ &\quad \left. - \frac{1}{(\sigma^2)^2} \text{tr}(\mathbf{Z}^T \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ(w)}) \right\}, \end{aligned}$$

where $\mathbf{U} = \Sigma^{-1} - \Sigma^{-1}\mathbf{X}(\mathbf{X}^T\Sigma^{-1}\mathbf{X})^{-1}\mathbf{X}^T\Sigma^{-1}$. Equating to zero we have

$$\sigma^{2(w+1)} = \frac{1}{n-t} \left\{ (\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)})^T \mathbf{U} (\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)}) + \text{tr}(\mathbf{Z}^T \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ(w)}) \right\}. \quad (4.5)$$

4.3.2 Update for γ_{gh}

Since $Q^{(U)}(\gamma; \kappa^{(w)})$ is equivalent to $Q^{(P)}(\gamma; \kappa^{(w)})$ the update for γ_{gh} , $g = 1, \dots, b$, $h = 1, \dots, b$, $g \geq h$ is the same as that presented in Section 3.6.2. It is reproduced here for convenience. Differentiating $Q^{(P)}(\gamma; \kappa^{(w)})$ with respect to γ_{gh} we have

$$\begin{aligned} \frac{\partial Q^{(P)}(\gamma; \kappa^{(w)})}{\partial \gamma_{gh}} &= -\frac{1}{2} \left\{ \text{tr} \left(\mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_{gh}} \right) - \tilde{\mathbf{u}}^{(w)T} \mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_{gh}} \mathbf{G}^{-1} \tilde{\mathbf{u}}^{(w)} \right. \\ &\quad \left. - \text{tr} \left(\mathbf{G}^{-1} \frac{\partial \mathbf{G}}{\partial \gamma_i} \mathbf{G}^{-1} \mathbf{C}^{ZZ(w)} \right) \right\}. \end{aligned}$$

The updating formula for γ_{gh} depends on the specific form of $\mathbf{G}$. For a one-way random effects model $\mathbf{G}$ has the simple form $\mathbf{G} = \sigma_u^2 \mathbf{I}_b$. In this case the updating equation is

$$\sigma_u^{2(w+1)} = \frac{1}{b} \left\{ \tilde{\mathbf{u}}^{(w)T} \tilde{\mathbf{u}}^{(w)} + \text{tr}(\mathbf{C}^{ZZ(w)}) \right\}. \quad (4.6)$$

4.3.3 Update for ϕ_{ij}

To find an update for ϕ_{ij} , $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ involves differentiating $Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})$ with respect to ϕ_{ij} . Noting that $\mathbf{R} = \sigma^2 \Sigma$ and using Result A.4.1 and Result A.4.3 we have

$$\begin{aligned} \frac{\partial Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \phi_{ij}} &= -\frac{1}{2} \left\{ \text{tr} \left(\mathbf{U} \frac{\partial \Sigma}{\partial \phi_{ij}} \right) - \frac{1}{\sigma^2} (\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)})^T \mathbf{U} \frac{\partial \Sigma}{\partial \phi_{ij}} \mathbf{U} (\mathbf{y}_o - \mathbf{Z}\tilde{\mathbf{u}}^{(w)}) \right. \\ &\quad \left. - \frac{1}{\sigma^2} \text{tr} \left(\mathbf{Z}^T \mathbf{U} \frac{\partial \Sigma}{\partial \phi_{ij}} \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ(w)} \right) \right\}, \end{aligned}$$

where $\mathbf{U} = \Sigma^{-1} - \Sigma^{-1}\mathbf{X}(\mathbf{X}^T\Sigma^{-1}\mathbf{X})^{-1}\mathbf{X}^T\Sigma^{-1}$. The updating formula for ϕ_{ij} depends on the specific form of Σ . In this thesis we only consider models where

$$\Sigma = I_n.$$

4.4 Complete data information matrix

We use the superscript (U) when considering the complete and missing data information matrices associated with $\mathbf{y}_c^{(U)}$. The complete and missing data information matrices are symmetric and in the interests of parsimony we only present the upper right triangle for each.

The complete data information matrix based on $\mathbf{y}_c^{(U)}$ is

$$\mathcal{I}_c^{(U)}(\boldsymbol{\kappa}^{(w)}; \mathbf{y}_c^{(U)}) = \begin{pmatrix} \mathcal{I}_{c_{(\sigma^2)^2}}^{(U)} & \mathbf{0} & \mathcal{I}_{c_{\sigma^2\phi}}^{(U)} \\ & \mathcal{I}_{c_{\gamma,\gamma}}^{(U)} & \mathbf{0} \\ & & \mathcal{I}_{c_{\phi,\phi}}^{(U)} \end{pmatrix},$$

where the zeroes in the complete data information matrix are due to the partitioning of $Q^{(U)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ into $Q^{(U)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})$ and $Q^{(U)}(\gamma; \boldsymbol{\kappa}^{(w)})$. Using the Oakes (1999) result presented in (1.13) we have

$$\begin{aligned} \mathcal{I}_{c_{(\sigma^2)^2}}^{(U)} &= - \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial(\sigma^2)^2} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}} \\ \mathcal{I}_{c_{\gamma,\gamma}}^{(U)} &= - \left\langle \frac{\partial^2 Q^{(U)}(\gamma; \boldsymbol{\kappa}^{(w)})}{\partial\gamma\partial\gamma^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}} \\ \mathcal{I}_{c_{\phi,\phi}}^{(U)} &= - \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial\phi\partial\phi^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}} \\ \mathcal{I}_{c_{\sigma^2,\phi}}^{(U)} &= - \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \boldsymbol{\kappa}^{(w)})}{\partial\sigma^2\partial\phi^T} \right\rangle_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^{(w)}}. \end{aligned}$$

We present each in turn and note any difference with their counterpart associated with $\mathbf{y}_c^{(P)}$.

4.4.1 Element $\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)}$

It can be shown that

$$\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)} = \frac{n-t}{2(\sigma^{2(w)})^2} + \frac{1}{(\sigma^{2(w)})^2} \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o - \frac{1}{(\sigma^{2(w)})^2} \text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)}). \quad (4.7)$$

Comparing $\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)}$ and $\mathcal{I}_{c_{(\sigma^2)^2}}^{(P)}$ we have

$$\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)} - \mathcal{I}_{c_{(\sigma^2)^2}}^{(P)} = -\frac{t}{2(\sigma^{2(w)})^2}.$$

For $t \geq 1$ and $\sigma^{2(w)} > 0$ we have the following inequality

$$\mathcal{I}_{c_{(\sigma^2)^2}}^{(U)} < \mathcal{I}_{c_{(\sigma^2)^2}}^{(P)}.$$

4.4.2 Elements of $\mathcal{I}_{c_{\gamma,\gamma}}^{(U)}$

Since $\mathcal{I}_{c_{\gamma,\gamma}}^{(U)}$ involves the second partial differentials of $Q^{(U)}(\gamma; \kappa^{(w)})$ and $Q^{(U)}(\gamma; \kappa^{(w)})$ is equivalent to $Q^{(P)}(\gamma; \kappa^{(w)})$ then the elements for $\mathcal{I}_{c_{\gamma,\gamma}}^{(U)}$ are equivalent to those for $\mathcal{I}_{c_{\gamma,\gamma}}^{(P)}$ which are presented in (3.16).

4.4.3 Elements of $\mathcal{I}_{c_{\phi,\phi}}^{(U)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ and $k = 1, \dots, n$, $l = 1, \dots, n$, $k \geq l$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\phi_{ij}, \phi_{kl}}}^{(U)} = & \frac{1}{2} \left\{ (\sigma^{2(w)})^2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ & + (\sigma^{2(w)})^2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\ & - \sigma^{2(w)} \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial^2 \Sigma^{(w)}}{\partial \phi_{ij}^{(w)} \partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o + \sigma^{2(w)} \text{tr} \left(\frac{\partial^2 \Sigma^{(w)}}{\partial \phi_{ij}^{(w)} \partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right) \\ & + (\sigma^{2(w)})^2 \text{tr} \left\{ \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{S}^{(w)} - \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right)^T \left(\mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} - \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\} \\ & \left. - (\sigma^{2(w)})^2 \text{tr} \left(\mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\}. \end{aligned} \quad (4.8)$$

The difference between $\mathcal{I}_{c_{\phi_{ij}, \phi_{kl}}}^{(U)}$ and $\mathcal{I}_{c_{\phi_{ij}, \phi_{kl}}}^{(P)}$ is that $\mathbf{R}^{-1(w)}$ replaces $\mathbf{S}^{(w)}$ in three terms in the later.

4.4.4 Elements of $\mathcal{I}_{c_{\sigma^2, \phi}}^{(U)}$

For $i = 1, \dots, n, j = 1, \dots, n, i \geq j$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\sigma^2(w) \phi_{ij}^{(w)}}}^{(U)} &= \frac{1}{2} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ &\quad \left. + \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{S}^{(w)} \right) - \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \right) \right\}. \end{aligned} \quad (4.9)$$

The difference between $\mathcal{I}_{c_{\sigma^2(w) \phi_{ij}^{(w)}}}^{(U)}$ and $\mathcal{I}_{c_{\sigma^2(w) \phi_{ij}^{(w)}}}^{(P)}$ is that $\mathbf{R}^{-1(w)}$ replaces $\mathbf{S}^{(w)}$ in one trace term in the latter.

This concludes the section on the complete data information matrix.

4.5 Missing data information matrix

The missing data information matrix based on $\mathbf{y}_c^{(U)}$ and using the Oakes (1999) result presented in (1.13) is

$$\mathcal{I}_m^{(U)}(\boldsymbol{\kappa}^{(w)}; \mathbf{y}_c^{(U)}) = \begin{pmatrix} \mathcal{I}_{m_{(\sigma^2)^2}}^{(U)} & \mathcal{I}_{m_{\sigma^2, \gamma}}^{(U)} & \mathcal{I}_{m_{\sigma^2, \phi}}^{(U)} \\ & \mathcal{I}_{m_{\gamma, \gamma}}^{(U)} & \mathcal{I}_{m_{\sigma^2, \gamma}}^{(U)} \\ & & \mathcal{I}_{m_{\phi, \phi}}^{(U)} \end{pmatrix},$$

where

$$\begin{aligned}
\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \sigma^2 \partial \sigma^{2(w)}} \right\rangle_{\kappa = \kappa^{(w)}} \\
\mathcal{I}_{m_{\gamma, \gamma}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\gamma; \kappa^{(w)})}{\partial \gamma \partial \gamma^{(w)T}} \right\rangle_{\kappa = \kappa^{(w)}} \\
\mathcal{I}_{m_{\phi, \phi}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \phi \partial \phi^{(w)T}} \right\rangle_{\kappa = \kappa^{(w)}} \\
\mathcal{I}_{m_{\sigma^2, \gamma}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \sigma^2 \partial \gamma^{(w)T}} \right\rangle_{\kappa = \kappa^{(w)}} \\
\mathcal{I}_{m_{\sigma^2, \phi}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\sigma^2, \phi; \kappa^{(w)})}{\partial \sigma^2 \partial \phi^{(w)T}} \right\rangle_{\kappa = \kappa^{(w)}} \\
\mathcal{I}_{m_{\gamma, \phi}}^{(U)} &= \left\langle \frac{\partial^2 Q^{(U)}(\gamma; \kappa^{(w)})}{\partial \gamma \partial \phi^{(w)T}} \right\rangle_{\kappa = \kappa^{(w)}}.
\end{aligned}$$

We make use of Result A.4.3, Result A.3.11, Result A.3.12, Result A.3.13, Result A.3.14 and present each in turn and note any differences with their counterpart associated with $\mathbf{y}_c^{(P)}$.

4.5.1 Element $\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)}$

It can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{(\sigma^2)^2}}^{(P)} &= -\frac{1}{2(\sigma^{2(w)})^2} \left\{ -(n-t) - 2\mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
&\quad + 2\mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \\
&\quad \left. + 2\text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)}) - \text{tr}(\mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{R}^{(w)}) \right\}. \quad (4.10)
\end{aligned}$$

Comparing $\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)}$ and $\mathcal{I}_{m_{(\sigma^2)^2}}^{(P)}$ we have

$$\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)} - \mathcal{I}_{m_{(\sigma^2)^2}}^{(P)} = -\frac{t}{2(\sigma^{2(w)})^2}.$$

For $t \geq 1$ and $\sigma^{2(w)} > 0$ we have the following inequality

$$\mathcal{I}_{m_{(\sigma^2)^2}}^{(U)} < \mathcal{I}_{m_{(\sigma^2)^2}}^{(P)}.$$

4.5.2 Elements of $\mathcal{I}_{m_{\gamma,\gamma}}^{(U)}$

Since $\mathcal{I}_{m_{\gamma,\gamma}}^{(U)}$ involves the second partial differentials of $Q^{(U)}(\gamma; \kappa^{(w)})$ and $Q^{(U)}(\gamma; \kappa^{(w)})$ is equivalent to $Q^{(P)}(\gamma; \kappa^{(w)})$ then the elements for $\mathcal{I}_{m_{\gamma,\gamma}}^{(U)}$ are equivalent to those for $\mathcal{I}_{m_{\gamma,\gamma}}^{(P)}$ which are presented in (3.20).

4.5.3 Elements of $\mathcal{I}_{m_{\phi,\phi}}^{(U)}$

For $i = 1, \dots, n, j = 1, \dots, n, i \geq j$ and $k = 1, \dots, n, l = 1, \dots, n, k \geq l$ it can be shown that

$$\begin{aligned} \mathcal{I}_{m_{\phi_{ij}, \phi_{kl}}}^{(U)} = & -\frac{(\sigma^{2(w)})^2}{2} \left[\mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\ & + \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\ & - \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\ & - \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\ & \left. - \text{tr} \left\{ \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{S}^{(w)} - \frac{\partial \Sigma^{(w)}}{\partial \phi_{kl}^{(w)}} \mathbf{P}^{(w)} \right)^T \left(\mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} - \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) \right\} \right]. \end{aligned} \quad (4.11)$$

The difference between $\mathcal{I}_{m_{\phi_{ij}, \phi_{kl}}}^{(U)}$ and $\mathcal{I}_{m_{\phi_{ij}, \phi_{kl}}}^{(P)}$ is that $\mathbf{R}^{-1(w)}$ replaces $\mathbf{S}^{(w)}$ in three terms in the latter.

4.5.4 Elements of $\mathcal{I}_{m_{\sigma^2, \gamma}}^{(U)}$

For $g = 1, \dots, b, h = 1, \dots, b, g \geq h$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\sigma^2}, \gamma_{gh}}^{(U)} = & -\frac{1}{2\sigma^{2(w)}} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
& + \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{y}_o \\
& \left. - \text{tr} \left(\mathbf{P}^{(w)} \mathbf{Z} \frac{\partial \mathbf{G}^{(w)}}{\partial \gamma_{gh}^{(w)}} \mathbf{Z}^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \right) \right\}. \tag{4.12}
\end{aligned}$$

This is equivalent to $\mathcal{I}_{m_{\sigma^2}, \gamma_{gh}}^{(P)}$ as presented in (3.22).

4.5.5 Elements of $\mathcal{I}_{m_{\sigma^2}, \phi}^{(U)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $i \geq j$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{m_{\sigma^2}, \phi_{ij}}^{(U)} = & -\frac{1}{2} \left\{ \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \mathbf{y}_o \right. \\
& + \mathbf{y}_o^T \mathbf{P}^{(w)} \mathbf{R}^{(w)} \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o - 2 \mathbf{y}_o^T \mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{y}_o \\
& - \text{tr} \left(\mathbf{S}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \right) + 2 \text{tr} \left(\frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \right) \\
& \left. - \text{tr} \left(\mathbf{P}^{(w)} \frac{\partial \Sigma^{(w)}}{\partial \phi_{ij}^{(w)}} \mathbf{P}^{(w)} \mathbf{R}^{(w)} \right) \right\}. \tag{4.13}
\end{aligned}$$

The difference between $\mathcal{I}_{m_{\sigma^2}, \phi_{ij}}^{(U)}$ and $\mathcal{I}_{m_{\sigma^2}, \phi_{ij}}^{(P)}$ is that $\mathbf{R}^{-1(w)}$ replaces $\mathbf{S}^{(w)}$ in a single trace term in the latter.

4.5.6 Elements of $\mathcal{I}_{m_{\gamma}, \phi}^{(U)}$

Since $\mathcal{I}_{m_{\gamma}, \phi}^{(U)}$ involves the second partial differentials of $Q^{(U)}(\gamma; \kappa^{(w)})$ and $Q^{(U)}(\gamma; \kappa^{(w)})$ is equivalent to $Q^{(P)}(\gamma; \kappa^{(w)})$ then the elements for $\mathcal{I}_{m_{\gamma}, \phi}^{(U)}$ are equivalent to those for $\mathcal{I}_{m_{\gamma}, \phi}^{(P)}$ which are presented in (3.24).

4.6 Comparing rates of convergence between $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$

For ease of notation we define

$$\begin{aligned}\mathcal{I}_c^{(U)} &= \mathcal{I}_c^{(U)}(\boldsymbol{\kappa}^*; \mathbf{y}_c^{(U)}), \\ \mathcal{I}_c^{(P)} &= \mathcal{I}_c^{(P)}(\boldsymbol{\kappa}^*; \mathbf{y}_c^{(P)}), \\ \mathbf{I}_o &= \mathbf{I}_o(\boldsymbol{\kappa}^*; \mathbf{y}_o),\end{aligned}$$

where $\boldsymbol{\kappa}^{(w)}$ has converged to $\boldsymbol{\kappa}^*$. We define the rate matrices

$$\begin{aligned}\mathcal{J}^{(U)}(\boldsymbol{\kappa}^*) &= \mathbf{I}_d - \mathcal{I}_c^{(U)-1} \mathbf{I}_o, \\ \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*) &= \mathbf{I}_d - \mathcal{I}_c^{(P)-1} \mathbf{I}_o,\end{aligned}$$

which can be written as

$$\begin{aligned}\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*) &= \mathcal{I}_c^{(U)-1} \mathbf{I}_o, \\ \mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*) &= \mathcal{I}_c^{(P)-1} \mathbf{I}_o.\end{aligned}$$

The rate of convergence associated with the complete data specification $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ is $\lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)\}$ and $\lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)\}$ respectively (McLachlan and Krishnan, 1997, Section 3.9.1). The term $\lambda_{\min}(\cdot)$ refers to the minimum eigenvalue. We aim to show

$$\lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)\} \geq \lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)\}, \quad (4.14)$$

i.e., the rate of convergence of a REML EM algorithm using the complete data specification $\mathbf{y}_c^{(U)}$ is faster than the rate of convergence of a REML EM algorithm using the complete data specification $\mathbf{y}_c^{(P)}$. We have not been able to prove that this inequality always holds. However, in our experience the rate of convergence when using $\mathbf{y}_c^{(U)}$ has always been faster than that when using $\mathbf{y}_c^{(P)}$. By considering a Loewner partial ordering (see Appendix C) of the complete information matrices

associated with $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ we are able to present the conditions under which the inequality in (4.14) will hold.

If we consider the difference between $\mathcal{I}_c^{(P)}$ and $\mathcal{I}_c^{(U)}$ to be positive definite then we can define the following Loewner partial ordering

$$\mathcal{I}_c^{(U)} \preceq \mathcal{I}_c^{(P)}.$$

Using Corollary C.1 (a) we can write

$$\mathcal{I}_c^{(U)^{-1}} \succeq \mathcal{I}_c^{(P)^{-1}}.$$

As a consequence of Corollary C.1 (b) we have the following inequalities between the complete data information matrices associated with $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$

$$\det(\mathcal{I}_c^{(U)^{-1}}) \geq \det(\mathcal{I}_c^{(P)^{-1}}), \quad (4.15)$$

$$\frac{1}{\det(\mathcal{I}_c^{(U)})} \geq \frac{1}{\det(\mathcal{I}_c^{(P)})}, \quad (4.16)$$

$$\det(\mathcal{I}_c^{(U)}) \leq \det(\mathcal{I}_c^{(P)}), \quad (4.17)$$

where going from (4.15) to (4.16) uses the second relation in Result A.3.3. Multiplying both sides of (4.17) by $\det(\mathbf{I}_o)$ and using the third relation in Result A.3.3 we have

$$\begin{aligned} \det(\mathcal{I}_c^{(U)^{-1}}) \det(\mathbf{I}_o) &\geq \det(\mathcal{I}_c^{(P)^{-1}}) \det(\mathbf{I}_o), \\ \det(\mathcal{I}_c^{(U)^{-1}} \mathbf{I}_o) &\geq \det(\mathcal{I}_c^{(P)^{-1}} \mathbf{I}_o), \\ \det\{\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)\} &\geq \det\{\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)\}. \end{aligned}$$

As a consequence of Corollary C.1 (b) and (c) it follows that

$$\boldsymbol{\lambda}_k\{\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)\} \geq \boldsymbol{\lambda}_k\{\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)\},$$

where $\boldsymbol{\lambda}_k(\cdot)$ is a vector of k eigenvalues. If the k eigenvalues of $\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)$ and $\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)$ are arranged in descending order then

$$\lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(U)}(\boldsymbol{\kappa}^*)\} \geq \lambda_{\min}\{\mathbf{I}_d - \mathcal{J}^{(P)}(\boldsymbol{\kappa}^*)\},$$

which concludes the proof. The key part of the proof is the inequality in (4.17). Elements of the complete data information matrices for both $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ can be used to check if the difference between $\mathbf{I}_c^{(P)}$ and $\mathbf{I}_c^{(U)}$ is positive definite for particular cases.

4.7 Example: lamb weight data

We apply a REML EM algorithm using the complete data specifications $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ to the lamb weight data presented in Callanan and Harville (1991) and provided in Table 4.1.

The data consists of birth weights (in pounds) of single-birth male lambs that are the progeny of 62 ewes. The age of the ewes are categorised into 3 groups (1 = 1 – 2 years; 2 = 2 – 3 years; 3 = over 3 years). The lambs were sired by one of 23 rams which belonged to 5 different population lines. A possible model for lamb birth weight involves fitting sire as a random effect and age and line as fixed effects (along with the overall mean). Using symbolic notation this model can be written as

$$\text{weight} \sim \text{mean} + \text{age} + \text{line} + \mathbf{sire},$$

where bold type is used to indicate that a term is being fitted as a random effect. Writing this model in matrix notation we have

$$\mathbf{y}_o = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u} + \mathbf{e},$$

where $\mathbf{y}_o$ is a 62×1 vector of observed birth weights, $\boldsymbol{\beta}$ is a 7×1 vector of fixed effects comprising the overall mean and the main effects of age and line. The 62×7 design matrix $\mathbf{X}$ is of full-column rank. The vector $\mathbf{u}$ is a 23×1 vector of random sire effects and $\mathbf{Z}$ is its 62×23 design matrix. The vector $\mathbf{e}$ is a 62×1 vector of residuals. We assume that $\mathbf{u}$ and $\mathbf{e}$ are independent of each other and have the following distributions

Table 4.1: Lamb weight data.

Sire	DAMage	Line	weight	Sire	DAMage	Line	weight
11	1	1	6.2	33	3	3	13.0
12	1	1	13.0	34	2	3	12.0
13	1	1	9.5	41	1	4	9.2
13	1	1	10.1	41	1	4	10.6
13	1	1	11.4	41	1	4	10.6
13	2	1	11.8	41	3	4	7.7
13	3	1	12.9	41	3	4	10.0
13	3	1	13.1	41	3	4	11.2
14	1	1	10.4	42	1	4	10.2
14	2	1	8.5	42	1	4	10.9
21	3	2	13.5	43	1	4	11.7
22	2	2	10.1	43	3	4	9.9
22	3	2	11.0	51	1	5	11.7
22	3	2	14.0	51	1	5	12.6
22	3	2	15.5	52	1	5	9.0
23	1	2	12.0	52	3	5	11.0
24	1	2	11.5	53	3	5	9.0
24	3	2	10.8	53	3	5	12.0
31	2	3	9.0	54	3	5	9.9
31	3	3	9.5	55	2	5	13.5
31	3	3	12.6	56	2	5	10.9
32	1	3	11.0	56	3	5	5.9
32	2	3	10.1	57	2	5	10.0
32	2	3	11.7	57	2	5	12.7
32	3	3	8.5	57	3	5	13.2
32	3	3	8.8	57	3	5	13.3
32	3	3	9.9	58	1	5	10.7
32	3	3	10.9	58	1	5	11.0
32	3	3	11.0	58	1	5	12.5
32	3	3	13.9	58	3	5	9.0
33	1	3	11.6	58	3	5	10.2

$$\begin{aligned} \mathbf{u} &\sim N(\mathbf{0}, \sigma_s^2 \mathbf{I}_{23}), \\ \mathbf{e} &\sim N(\mathbf{0}, \sigma^2 \mathbf{I}_{62}), \end{aligned}$$

so that $\mathbf{R} = \sigma^2 \mathbf{I}_{62}$ and $\mathbf{G} = \sigma_s^2 \mathbf{I}_{23}$. The updating equation for σ^2 is provided in (3.14) and (4.5) for $\mathbf{y}_c^{(P)}$ and $\mathbf{y}_c^{(U)}$ respectively. The updating equation for σ_s^2 is provided in (4.6) and applies to both complete data specifications. We define $\boldsymbol{\kappa} = (\sigma^2, \sigma_s^2)^T$ to be the vector of variance parameters to be estimated using a REML EM algorithm. The convergence criterion of Foulley and Van Dyk (2000) was used so convergence is defined as being achieved on the w -th iteration if

$$\sqrt{\frac{(\boldsymbol{\kappa}^{(w+1)} - \boldsymbol{\kappa}^{(w)})^T (\boldsymbol{\kappa}^{(w+1)} - \boldsymbol{\kappa}^{(w)})}{\boldsymbol{\kappa}^{T(w)} \boldsymbol{\kappa}^{(w)}}} < \varepsilon, \quad (4.18)$$

where $\varepsilon = 10^{-8}$. Two sets of starting values were considered for both complete data specifications. The first set of starting values were $\boldsymbol{\kappa}^{(0)} = (\sigma^{2(0)}, \sigma_s^{2(0)})^T = (1, 0.01)^T$ and the second set of starting values were $\boldsymbol{\kappa}^{(0)} = (\sigma^{2(0)}, \sigma_s^{2(0)})^T = (1, 5)^T$. The statistical software package R (R Development Core Team, 2011) was used to implement both REML EM algorithms. REML estimates of σ^2 and σ_s^2 , denoted σ^{2*} and σ_s^{2*} , the number of iterations until convergence and the rate of convergence for both complete data specifications and both sets of starting values are provided in Table 4.2.

Table 4.2: Results for two REML EM algorithms applied to the lamb weight data.

Complete data	σ^{2*}	σ_s^{2*}	Iterations	Rate of convergence
$\sigma^{2(0)} = 1, \sigma_s^{2(0)} = 0.01$				
$\mathbf{y}_c^{(P)}$	2.9615972	0.5170759	1296	0.9630649
$\mathbf{y}_c^{(U)}$	2.9615972	0.5170759	1296	0.9629925
$\sigma^{2(0)} = 1, \sigma_s^{2(0)} = 5$				
$\mathbf{y}_c^{(P)}$	2.9615966	0.5170773	342	0.9630648
$\mathbf{y}_c^{(U)}$	2.9615966	0.5170773	341	0.9629924

For both sets of starting values and for both complete data specifications the observed information matrix was found to be (to 4 decimal places)

$$\mathbf{I}_o(\boldsymbol{\kappa}^*; \mathbf{y}_o) = \begin{pmatrix} 2.6598 & 1.1753 \\ 1.1753 & 2.13448 \end{pmatrix},$$

which can be used to calculate the standard errors associated with the REML estimates of σ^2 and σ_s^2 . It is worth noting that the empirical rate of convergence can be used to check the rate of convergence and the observed information matrix.

Based on the rates of convergence in Table 4.2 there appears to be little to gain in using $\mathbf{y}_c^{(U)}$ instead of $\mathbf{y}_c^{(P)}$. Our experience with a number of data sets is that the difference in the rate of convergence between $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ when using a REML EM algorithm is only minor. However, this difference can be significant when a parameter expanded version of the REML EM algorithm, referred to as a REML PX EM algorithm, is implemented. We consider a REML PX EM algorithm for both complete data specifications in the next chapter.

Chapter 5

Two REML PX EM algorithms for the linear mixed model

A common criticism of the EM algorithm is that it can be slow to converge. Since the Dempster et al. (1977) paper there have been a number of schemes suggested for speeding up the convergence of the EM algorithm. In this chapter we will consider one such scheme, known as the parameter expanded (PX) EM algorithm (Liu et al., 1998). From a practical implementation point of view a PX EM algorithm is preferred to an EM algorithm as there is nothing to lose and often a lot to gain. Nothing to lose in the sense that the PX EM algorithm, like the EM algorithm, has monotonic convergence. Furthermore, the E-step of the PX EM algorithm is exactly the same as the E-step of the EM algorithm. The gains in using the PX EM algorithm compared to the EM algorithm are in rates of convergence. Liu et al. (1998) show that the PX EM algorithm has a rate of convergence that is no slower than the EM algorithm but is often much faster. In this chapter we will consider REML PX EM algorithms for the two complete data specifications considered in the previous chapter.

5.1 Linear mixed model specification

For the PX EM algorithm the linear mixed model is formulated by introducing the auxillary parameter λ (say) so that the linear mixed model in (2.1) is expanded to

$$\mathbf{y}_o = \mathbf{X}\boldsymbol{\beta}_* + \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f} + \boldsymbol{\epsilon}, \quad (5.1)$$

where $\mathbf{X}$ and $\mathbf{Z}$ are $n \times t$ and $n \times b$ design matrices respectively, and $\boldsymbol{\Lambda} = \boldsymbol{\Lambda}(\boldsymbol{\lambda})$ is a $b \times b$ real invertible matrix which is a function of the $v \times 1$ auxillary parameter vector $\boldsymbol{\lambda}$. It is assumed that the joint distribution of $\mathbf{f}$ and $\boldsymbol{\epsilon}$ is

$$\begin{pmatrix} \mathbf{f} \\ \boldsymbol{\epsilon} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{D} & \mathbf{0} \\ \mathbf{0} & \mathbf{R}_* \end{pmatrix} \right\},$$

where $\mathbf{D} = \mathbf{D}(\mathbf{d})$ and $\mathbf{R}_* = \sigma_*^2 \boldsymbol{\Sigma}_*(\boldsymbol{\phi}_*)$ are symmetric positive definite matrices. The marginal distribution of $\mathbf{y}_o$ under the parameter expanded model is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}_*, \mathbf{H}_*)$$

where $\mathbf{H}_* = \mathbf{Z}\boldsymbol{\Lambda}\mathbf{D}\boldsymbol{\Lambda}^T\mathbf{Z}^T + \mathbf{R}_*$. The expanded parameter set for the model in (5.1) is $\boldsymbol{\mathcal{K}} = (\boldsymbol{\kappa}_*, \boldsymbol{\lambda}^T)^T$ where $\boldsymbol{\kappa}_* = (\sigma_*^2, \mathbf{d}^T, \boldsymbol{\phi}_*^T)^T$. The subscript $*$ is used to distinguish parameters associated with the linear mixed model specified in (5.1) and the linear mixed model specified in (2.1). The expanded parameter vector $\boldsymbol{\mathcal{K}}$ must satisfy two conditions:

1. $\boldsymbol{\mathcal{K}}$ can be reduced to the original parameter vector $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$ by a many-to-one reduction function. For the model specified in (5.1) this reduction function, labelled $\mathcal{R}$ is $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T = \mathcal{R}(\boldsymbol{\mathcal{K}}) = (\sigma_*^2, \text{vech}(\boldsymbol{\Lambda}\mathbf{D}\boldsymbol{\Lambda}^T), \boldsymbol{\phi}_*^T)^T$.
2. When the auxillary parameter $\boldsymbol{\lambda}$ is set to its null value $\boldsymbol{\lambda}_0$ (say) the expanded complete data model is reduced to the original complete data model. If we define $\mathbf{u} = \boldsymbol{\Lambda}\mathbf{f}$ then $\mathbf{G} = \boldsymbol{\Lambda}\mathbf{D}\boldsymbol{\Lambda}^T$ and when $\boldsymbol{\Lambda}(\boldsymbol{\lambda}_0) = \mathbf{I}_b$ the expanded model in (5.1) is reduced to the original model in (2.1) with $\mathbf{u} = \mathbf{f}$ and $\mathbf{G} = \mathbf{D}$.

Liu et al. (1998) state that

The underlying statistical principle of PX-EM is to perform a ‘covariance adjustment’ to correct the M step, capitalising on extra information captured in the imputed complete data...

An intuitively appealing way to consider this ‘covariance adjustment’ is to compare the joint distribution of the incomplete and missing data under the parameter expanded model, i.e.,

$$\begin{pmatrix} \mathbf{y}_o \\ \mathbf{f} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{X}\boldsymbol{\beta}_* \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{H}_* & \mathbf{Z}\boldsymbol{\Lambda}\mathbf{D} \\ \mathbf{D}\boldsymbol{\Lambda}^T\mathbf{Z}^T & \mathbf{D} \end{pmatrix} \right\}$$

with the joint distribution of the incomplete and missing data under the original model, i.e.,

$$\begin{pmatrix} \mathbf{y}_o \\ \mathbf{u} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{X}\boldsymbol{\beta} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{H} & \mathbf{Z}\mathbf{G} \\ \mathbf{G}\mathbf{Z}^T & \mathbf{D} \end{pmatrix} \right\}.$$

The idea of parameter expansion as a ‘covariance adjustment’ is easily seen by the introduction of $\boldsymbol{\Lambda}$ into the covariance of the former. As Liu et al. (1998) note, the covariance of the later does not generally reflect the “true” relationship between the imputed missing data $\tilde{\mathbf{u}}^{(w+1)}$ and the known incomplete data vector $\mathbf{y}_o$, particularly in the initial iterations of an EM algorithm. The auxiliary parameter $\boldsymbol{\Lambda}$ is used to “correct” the estimate of $\mathbf{G}$ to produce an adjusted estimate. At convergence there is no longer any need to adjust the estimate of $\mathbf{G}$ and $\boldsymbol{\Lambda}(\boldsymbol{\lambda}) = \mathbf{I}_b$.

5.2 REML PX EM algorithm for $\mathbf{y}_c^{(P)}$

A REML PX EM algorithm for the linear mixed model using the complete data specification $\mathbf{y}_c^{(P)} = (\mathbf{y}_o^T, \mathbf{f}^T)^T$ was considered by Foulley and Van Dyk (2000) who used the random effects approach, and Cullis et al. (2004) and Knight (2008) who used the fixed effects approach. Our exposition of a REML PX EM algorithm using $\mathbf{y}_c^{(P)}$ is brief as the detail can be found in the aforementioned publications.

5.2.1 Complete data log density function

The complete data log-density function for $\mathbf{y}_c^{(P)}$ is

$$\begin{aligned} \log\{f_X(\mathbf{y}_c^{(P)}; \boldsymbol{\Theta})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R}_*)\} + \log\{\det(\mathbf{D})\} \right. \\ &\quad \left. + \mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} + \boldsymbol{\epsilon}^T \mathbf{R}_*^{-1} \boldsymbol{\epsilon} \right], \end{aligned} \tag{5.2}$$

where $\boldsymbol{\epsilon} = \mathbf{y}_o - \mathbf{X}\boldsymbol{\beta}_* - \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f}$ and $\boldsymbol{\Theta} = (\boldsymbol{\beta}_*^T, \boldsymbol{\kappa}_*^T, \boldsymbol{\lambda}^T)^T$. The subscript X will be used to distinguish between the complete data log-density function and the $Q(\cdot)$ function associated with the model in (5.1) and the model in (2.1).

5.2.2 REML PX EM algorithm E-step

The complete data log-density function in (5.2) can be partitioned

$$\log\{f_X(\mathbf{y}_c^{(P)}; \boldsymbol{\Theta})\} = \log\{f_X(\boldsymbol{\epsilon}; \boldsymbol{\Theta})\} + \log\{f_X(\mathbf{f}; \mathbf{d})\},$$

where

$$\begin{aligned}\log\{f_X(\boldsymbol{\epsilon}; \boldsymbol{\Theta})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{R}_*)\} + \boldsymbol{\epsilon}^T \mathbf{R}_*^{-1} \boldsymbol{\epsilon} \right], \\ \log\{f_X(\mathbf{f}; \mathbf{d})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D})\} + \mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} \right].\end{aligned}$$

The E-step of the REML PX EM algorithm involves taking the expectation of (5.2) over $k(\mathbf{f}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)})$ where $\boldsymbol{\kappa}^{(w)} = \{\boldsymbol{\kappa}_*^{(w)T}, (\boldsymbol{\lambda} = \boldsymbol{\lambda}_0)^T\}^T$. We can write

$$\begin{aligned}Q_X^{(P)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)}) &= E[\log\{f_X(\mathbf{y}_c^{(P)}; \boldsymbol{\Theta})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= E[\log\{f_X(\boldsymbol{\epsilon}; \boldsymbol{\Theta})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] + E[\log\{f_X(\mathbf{f}; \mathbf{d})\}|\mathbf{y}_2; \boldsymbol{\kappa}^{(w)}] \\ &= Q_X^{(P)}(\sigma_*^2, \phi_*, \boldsymbol{\lambda}; \boldsymbol{\kappa}^{(w)}) + Q_X^{(P)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)}),\end{aligned}$$

where

$$\begin{aligned}Q_X^{(P)}(\sigma_*^2, \phi_*, \boldsymbol{\lambda}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{R}_*)\} + E(\boldsymbol{\epsilon}^T \mathbf{R}_*^{-1} \boldsymbol{\epsilon} | \mathbf{y}_2; \boldsymbol{\kappa}^{(w)}) \right], \\ Q_X^{(P)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{D})\} + E(\mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} | \mathbf{y}_2; \boldsymbol{\kappa}^{(w)}) \right].\end{aligned}$$

When $\boldsymbol{\Lambda}(\boldsymbol{\lambda}_0) = \mathbf{I}_b$ the REML PX EM algorithm E-step is the same as the REML EM algorithm E-step.

5.2.3 REML PX EM algorithm M-step

The M-step of the REML PX EM algorithm involves maximising $Q_X^{(P)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ with respect to $\boldsymbol{\kappa} = (\sigma_*^2, \mathbf{d}^T, \boldsymbol{\phi}_*^T, \boldsymbol{\lambda}^T)^T$ and then applying the reduction function $\mathcal{R}(\boldsymbol{\kappa}) = (\sigma_*^2, \text{vech}(\boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T), \boldsymbol{\phi}_*^T)^T$ to obtain estimates of $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$.

Update for σ_*^2

The update for σ_*^2 is the same as that presented in (3.14), i.e.,

$$\sigma_*^{2(w+1)} = \frac{1}{n} \left\{ \tilde{\mathbf{e}}^{(w)T} \boldsymbol{\Sigma}^{-1(w)} \tilde{\mathbf{e}}^{(w)} + \text{tr}(\boldsymbol{\Sigma}^{-1(w)} \mathbf{W} \mathbf{C}^{-1(w)} \mathbf{W}^T) \right\}. \quad (5.3)$$

Applying the reduction function we have $\sigma^{2(w+1)} = \sigma_*^{2(w+1)}$.

Update for d_{gh}

The update for d_{gh} is the same as that for γ_{gh} when using the REML EM algorithm except we substitute $\mathbf{D}$ for $\mathbf{G}$, i.e., differentiating $Q_X^{(P)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)})$ with respect to d_{gh} we have

$$\frac{\partial Q_X^{(P)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)})}{\partial d_{gh}} = -\frac{1}{2} \left\{ \text{tr} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{gh}} \right) - E \left(\mathbf{f}^T \mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{gh}} \mathbf{D}^{-1} \mathbf{f} \middle| \mathbf{y}_2; \boldsymbol{\kappa}^{(w)} \right) \right\}.$$

The updating formula for d_{gh} depends on the specific form of $\mathbf{D}$. If we define the update for $\mathbf{D}$ as $\mathbf{D}^{(w+1)}$ then applying the reduction function we have $\boldsymbol{\gamma}^{(w+1)} = \text{vech}(\boldsymbol{\Lambda}^{(w+1)} \mathbf{D}^{(w+1)} \boldsymbol{\Lambda}^{(w+1)T})$. For a one-way random effects model $\mathbf{D}$ has the form $\mathbf{D} = d \mathbf{I}_b$. In this case the updating equation is

$$d^{(w+1)} = \frac{1}{b} \left\{ \tilde{\mathbf{u}}^{(w)T} \tilde{\mathbf{u}}^{(w)} + \text{tr}(\mathbf{C}^{ZZ(w)}) \right\}, \quad (5.4)$$

and the reduction function is $\sigma_u^{2(w+1)} = \lambda^{2(w+1)} d^{(w+1)}$.

Update for ϕ_{*ij}

The update for ϕ_{*ij} is the same as that for the REML EM algorithm. Differentiating $Q_X^{(P)}(\sigma_*^2, \phi_*, \lambda; \mathcal{K}_*^{(w)})$ with respect to ϕ_{*ij} we have

$$\frac{\partial Q_X^{(P)}(\sigma_*^2, \phi_*, \lambda; \mathcal{K}_*^{(w)})}{\partial \phi_{*ij}} = -\frac{1}{2} \left\{ \text{tr} \left(\Sigma_*^{-1} \frac{\partial \Sigma_*}{\partial \phi_{*ij}} \right) - \frac{1}{\sigma_*^2} E \left(\epsilon^T \Sigma_*^{-1} \frac{\partial \Sigma_*}{\partial \phi_{*ij}} \Sigma_*^{-1} \epsilon \middle| \mathbf{y}_2; \mathcal{K}_*^{(w)} \right) \right\}.$$

The updating equation for ϕ_{*ij} depends on the specific form of Σ_* . If we define the update for ϕ_{*ij} as $\phi_{*ij}^{(w+1)}$ then applying the reduction function we have $\phi_{ij}^{(w+1)} = \phi_{*ij}^{(w+1)}$.

Update for λ

To form an update for λ we re-write $Q_X^{(P)}(\sigma_*^2, \phi_*, \lambda; \mathcal{K}_*^{(w)})$ using the identity

$$\begin{aligned} \epsilon^T R_*^{-1} \epsilon &= (\mathbf{y}_o - \mathbf{X}\beta_* - \mathbf{Z}\Lambda\mathbf{f})^T R_*^{-1} (\mathbf{y}_o - \mathbf{X}\beta_* - \mathbf{Z}\Lambda\mathbf{f}) \\ &= (\mathbf{y}_o - \mathbf{X}\beta_*)^T R_*^{-1} (\mathbf{y}_o - \mathbf{X}\beta_*) - 2(\mathbf{y}_o - \mathbf{X}\beta_*)^T R_*^{-1} \mathbf{Z}\Lambda\mathbf{f} \\ &\quad + \mathbf{f}^T \Lambda^T \mathbf{Z}^T R_*^{-1} \mathbf{Z}\Lambda\mathbf{f}. \end{aligned} \tag{5.5}$$

Performing the REML PX EM E-step on the right hand side of (5.5) requires the joint conditional distribution

$$\begin{pmatrix} \mathbf{y}_o - \mathbf{X}\beta \\ \mathbf{u} \end{pmatrix} \middle| \mathbf{y}_2 \sim N \left\{ \begin{pmatrix} \mathbf{H} \mathbf{P} \mathbf{y}_o \\ \mathbf{G} \mathbf{Z}^T \mathbf{P} \mathbf{y}_o \end{pmatrix}, \begin{pmatrix} \mathbf{X} \mathbf{C}^{XX} \mathbf{X}^T & -\mathbf{X} \mathbf{C}^{XZ} \\ -\mathbf{C}^{ZX} \mathbf{X}^T & \mathbf{C}^{ZZ} \end{pmatrix} \right\}.$$

The update for λ can be expressed as

$$\lambda^{(w+1)} = \mathbf{A}^{-1(w)} \mathbf{b}^{(w)},$$

where the elements of the $b^2 \times b^2$ matrix $\mathbf{A}^{(w)}$ are

$$a_{gh,qr}^{(w)} = \tilde{\mathbf{u}}^{(w)T} \frac{\partial \Lambda^{(w)T}}{\partial \lambda_{gh}^{(w)}} \mathbf{Z}^T \mathbf{R}^{-1(w)} \mathbf{Z} \frac{\partial \Lambda^{(w)}}{\partial \lambda_{qr}^{(w)}} \tilde{\mathbf{u}}^{(w)} + \text{tr} \left(\frac{\partial \Lambda^{(w)T}}{\partial \lambda_{gh}^{(w)}} \mathbf{Z}^T \mathbf{R}^{-1(w)} \mathbf{Z} \frac{\partial \Lambda^{(w)}}{\partial \lambda_{qr}^{(w)}} \mathbf{C}^{ZZ(w)} \right), \tag{5.6}$$

and the elements of the $b^2 \times 1$ vector $\mathbf{b}^{(w)}$ are

$$b_{gh}^{(w)} = \tilde{\mathbf{u}}^{(w)} \frac{\partial \Lambda^{(w)T}}{\partial \lambda_{gh}^{(w)}} \mathbf{Z}^T \mathbf{R}^{-1(w)} (\mathbf{y}_o - \mathbf{X} \hat{\boldsymbol{\beta}}^{(w)}) - \text{tr} \left(\frac{\partial \Lambda^{(w)T}}{\partial \lambda_{gh}^{(w)}} \mathbf{Z}^T \mathbf{R}^{-1(w)} \mathbf{X} \mathbf{C}^{\mathbf{XZ}^{(w)}} \right), \quad (5.7)$$

where $\hat{\boldsymbol{\beta}}^{(w)}$ is the generalised least squares estimate of $\boldsymbol{\beta}$ at the w -th iterate. This completes the M-step.

5.2.4 Rate matrix

Liu et al. (1998) note that the convergence of $\boldsymbol{\kappa}^{(w)}$ determines the rate of convergence of the the PX EM algorithm. The rate matrix for the REML PX EM algorithm using the complete data specification $\boldsymbol{y}_c^{(P)}$ is

$$\mathcal{J}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0) = \mathbf{I}_{d+v} - \{\mathcal{I}_c^{(P)}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0)\}^{-1} \mathbf{I}_o(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0), \quad (5.8)$$

where $\boldsymbol{\kappa}^{(w)}$ has converged to $\boldsymbol{\kappa}^*$, and where

$$\mathbf{I}_o(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0) = \begin{pmatrix} \mathbf{I}_o(\boldsymbol{\kappa}^*) & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix},$$

and

$$\mathcal{I}_c^{(P)}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0) = \begin{pmatrix} \mathcal{I}_{c_{\boldsymbol{\kappa}, \boldsymbol{\kappa}}}^{(P)} & \mathcal{I}_{c_{\boldsymbol{\kappa}, \boldsymbol{\lambda}}}^{(P)} \\ & \mathcal{I}_{c_{\boldsymbol{\lambda}, \boldsymbol{\lambda}}}^{(P)} \end{pmatrix}_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^*, \boldsymbol{\lambda} = \boldsymbol{\lambda}_0}.$$

The terms $\mathcal{I}_{c_{\boldsymbol{\kappa}, \boldsymbol{\kappa}}}^{(P)}$ and $\mathbf{I}_o(\boldsymbol{\kappa}^*)$ can be calculated using the complete and missing data information matrices presented in Section 3.7 and Section 3.8. The elements of $\mathcal{I}_{c_{\boldsymbol{\kappa}, \boldsymbol{\lambda}}}^{(P)}$, which comprise $\mathcal{I}_{c_{\sigma^2, \boldsymbol{\lambda}}}^{(P)}$, $\mathcal{I}_{c_{\gamma, \boldsymbol{\lambda}}}^{(P)}$, and $\mathcal{I}_{c_{\phi, \boldsymbol{\lambda}}}^{(P)}$, and the elements of $\mathcal{I}_{c_{\boldsymbol{\lambda}, \boldsymbol{\lambda}}}^{(P)}$ are presented below.

Elements of $\mathcal{I}_{c_{\sigma^2, \boldsymbol{\lambda}}}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\sigma^2}, \lambda_{gh}}^{(P)} &= \frac{1}{(\sigma^{2*})^2} \left[\text{tr} \left\{ \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^{*T} (\mathbf{y}_o - \mathbf{X} \hat{\beta}) - \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZX*} \mathbf{X}^T \right\} \right. \\ &\quad \left. - \tilde{\mathbf{u}}^{*T} \mathbf{Z}^T \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* - \text{tr} \left(\mathbf{Z}^T \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \right]. \end{aligned}$$

Elements of $\mathcal{I}_{c_{\gamma}, \lambda}^{(P)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$ it can be shown that

$$\mathcal{I}_{c_{\gamma_{gh}}, \lambda_{gh}}^{(P)} = \tilde{\mathbf{u}}^{*T} \mathbf{G}^{-1*} \frac{\partial \mathbf{G}^*}{\partial \gamma_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* + \text{tr} \left(\mathbf{G}^{-1*} \frac{\partial \mathbf{G}^*}{\partial \gamma_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right).$$

Elements of $\mathcal{I}_{c_{\phi}, \lambda}^{(P)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $g = 1, \dots, b$, $h = 1, \dots, b$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\phi_{ij}}, \lambda_{gh}}^{(P)} &= \frac{1}{\sigma^{2*}} \left[\text{tr} \left\{ \Sigma^{-1*} \frac{\partial \Sigma^*}{\partial \phi_{ij}^*} \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^{*T} (\mathbf{y}_o - \mathbf{X} \hat{\beta}) \right. \right. \\ &\quad \left. - \Sigma^{-1*} \frac{\partial \Sigma^*}{\partial \phi_{ij}^*} \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZX*} \mathbf{X}^T \right\} \\ &\quad - \tilde{\mathbf{u}}^{*T} \mathbf{Z}^T \Sigma^{-1*} \frac{\partial \Sigma^*}{\partial \phi_{ij}^*} \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* \\ &\quad \left. - \text{tr} \left(\mathbf{Z}^T \Sigma^{-1*} \frac{\partial \Sigma^*}{\partial \phi_{ij}^*} \Sigma^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \right]. \end{aligned}$$

Elements of $\mathcal{I}_{c_{\lambda}, \lambda}^{(P)}$

We assume that the elements of the real invertible matrix Λ are λ_{gh} , $g = 1, \dots, b$, $h = 1, \dots, b$ and that the second order partial derivative of Λ can be expressed as

$$\frac{\partial^2 \Lambda}{\partial \lambda_{gh} \partial \lambda_{qr}},$$

where $q = 1, \dots, b$, $r = 1, \dots, b$. It can be shown that

$$\begin{aligned}
\mathcal{I}_{c_{\lambda_{gh}, \lambda_{qr}}}^{(P)} &= \tilde{\mathbf{u}}^{*T} \frac{\partial \Lambda^{*T}}{\partial \lambda_{qr}^*} \mathbf{Z}^T \mathbf{R}^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* + \text{tr} \left(\frac{\partial \Lambda^{*T}}{\partial \lambda_{qr}^*} \mathbf{Z}^T \mathbf{R}^{-1*} \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \\
&+ \text{tr} \left(\frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \Lambda^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \Lambda^{-1*} \right) + \tilde{\mathbf{u}}^{*T} \frac{\partial \Lambda^{*T}}{\partial \lambda_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \tilde{\mathbf{u}}^* \\
&+ \text{tr} \left(\frac{\partial \Lambda^{*T}}{\partial \lambda_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \mathbf{C}^{ZZ*} \right).
\end{aligned}$$

5.3 REML PX EM algorithm for $\mathbf{y}_c^{(U)}$

A REML PX EM algorithm for the linear mixed model using the complete data specification $\mathbf{y}_c^{(U)} = (\mathbf{y}_2^T, \mathbf{f}^T)^T$ has not been considered previously. We present the joint distribution of the incomplete and missing data for the parameter expanded model, the complete data log-density function, the E and M-steps, and the rate matrix for the REML PX EM algorithm using $\mathbf{y}_c^{(U)}$.

For the parameter expanded model based on $\mathbf{y}_c^{(U)}$ the joint distribution of the incomplete and missing data is

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{f} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_2^T \mathbf{H}_* \mathbf{L}_2 & \mathbf{L}_2^T \mathbf{Z} \Lambda \mathbf{D} \\ \mathbf{D} \Lambda^T \mathbf{Z}^T \mathbf{L}_2 & \mathbf{D} \end{pmatrix} \right\} \quad (5.9)$$

where $\mathbf{H}_* = \mathbf{Z} \Lambda \mathbf{D} \Lambda^T \mathbf{Z}^T + \mathbf{R}_*$. This can be compared to the joint distribution of the incomplete and missing data for the original model presented in (4.1).

5.3.1 Complete data log density function

The complete data log density function for $\mathbf{y}_c^{(U)}$ is

$$\begin{aligned}
\log\{f_X(\mathbf{y}_c^{(U)}; \mathbf{K})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R}_* \mathbf{L}_2)\} + \log\{\det(\mathbf{D})\} \right. \\
&\quad \left. + \mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} + (\mathbf{y}_o - \mathbf{Z} \Lambda \mathbf{f})^T \mathbf{S}_* (\mathbf{y}_o - \mathbf{Z} \Lambda \mathbf{f}) \right],
\end{aligned}$$

where $\mathbf{S}_* = \mathbf{R}_*^{-1} - \mathbf{R}_*^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{R}_*^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}_*^{-1}$. Like the REML EM, a major difference between (5.2) and (5.10) is that the REML PX EM algorithm E-step using

the former requires at least the joint conditional distribution of $\mathbf{u}$ and $\mathbf{y}_o - \mathbf{X}\boldsymbol{\beta}$ given $\mathbf{y}_2$ whereas the later only requires $\mathbf{u}$ given $\mathbf{y}_2$.

5.3.2 REML PX EM algorithm E-step

The complete data log-density function in (5.10) can be partitioned

$$\log\{f_X(\mathbf{y}_c^{(U)}; \mathcal{K})\} = \log\{f_X(\mathbf{L}_2^T \boldsymbol{\epsilon}; \mathcal{K})\} + \log\{f_X(\mathbf{f}; \mathbf{d})\},$$

where

$$\begin{aligned} \log\{f_X(\mathbf{L}_2^T \boldsymbol{\epsilon}; \mathcal{K})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R}_* \mathbf{L}_2)\} + (\mathbf{y}_o - \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f})^T \mathbf{S}_* (\mathbf{y}_o - \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f}) \right], \\ \log\{f_X(\mathbf{f}; \mathbf{d})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D})\} + \mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} \right]. \end{aligned}$$

The E-step of the REML PX EM algorithm involves taking the expectation of (5.10) over $k(\mathbf{f}|\mathbf{y}_2; \mathcal{K}^{(w)})$ where $\mathcal{K}^{(w)} = \{\boldsymbol{\kappa}_*^{(w)T}, (\boldsymbol{\lambda} = \boldsymbol{\lambda}_0)^T\}^T$. We can write

$$\begin{aligned} Q_X^{(U)}(\mathcal{K}; \mathcal{K}^{(w)}) &= E[\log\{f_X(\mathbf{y}_c^{(U)}; \mathcal{K})\} | \mathbf{y}_2; \mathcal{K}^{(w)}] \\ &= E[\log\{f_X(\mathbf{L}_2^T \boldsymbol{\epsilon}; \mathcal{K})\} | \mathbf{y}_2; \mathcal{K}^{(w)}] + E[\log\{f_X(\mathbf{f}; \mathbf{d})\} | \mathbf{y}_2; \mathcal{K}^{(w)}] \\ &= Q_X^{(U)}(\sigma_*^2, \boldsymbol{\phi}_*, \boldsymbol{\lambda}; \mathcal{K}^{(w)}) + Q_X^{(U)}(\mathbf{d}; \mathcal{K}^{(w)}), \end{aligned}$$

where

$$\begin{aligned} Q_X^{(U)}(\sigma_*^2, \boldsymbol{\phi}_*, \boldsymbol{\lambda}; \mathcal{K}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R}_* \mathbf{L}_2)\} \right. \\ &\quad \left. + E\left\{(\mathbf{y}_o - \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f})^T \mathbf{S}_* (\mathbf{y}_o - \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f}) | \mathbf{y}_2; \mathcal{K}^{(w)}\right\} \right], \\ Q_X^{(U)}(\mathbf{d}; \mathcal{K}^{(w)}) &= -\frac{1}{2} \left[\log\{\det(\mathbf{D})\} + E(\mathbf{f}^T \mathbf{D}^{-1} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)}) \right]. \end{aligned}$$

Note that $Q_X^{(U)}(\mathbf{d}; \mathcal{K}^{(w)})$ is equivalent to $Q_X^{(P)}(\mathbf{d}; \mathcal{K}^{(w)})$. When $\boldsymbol{\Lambda}(\boldsymbol{\lambda}_0) = \mathbf{I}_b$ the REML PX EM algorithm E-step is the same as the REML EM algorithm E-step.

5.3.3 REML PX EM algorithm M-step

The M-step of the REML PX EM algorithm involves maximising $Q_X^{(U)}(\boldsymbol{\kappa}; \boldsymbol{\kappa}^{(w)})$ with respect to $\boldsymbol{\kappa} = (\sigma_*^2, \mathbf{d}^T, \boldsymbol{\phi}_*^T, \boldsymbol{\lambda}^T)^T$ and then applying the reduction function $\mathcal{R}(\boldsymbol{\kappa}) = (\sigma_*^2, \text{vech}(\boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T), \boldsymbol{\phi}_*^T)^T$ to obtain estimates of $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\gamma}^T, \boldsymbol{\phi}^T)^T$.

Update for σ_*^2

The update for σ_*^2 is the same as that presented in (4.5), i.e.,

$$\sigma_*^{2(w+1)} = \frac{1}{n-t} \left\{ (\mathbf{y}_o - \mathbf{Z} \tilde{\mathbf{u}}^{(w)})^T \mathbf{U} (\mathbf{y}_o - \mathbf{Z} \tilde{\mathbf{u}}^{(w)}) + \text{tr}(\mathbf{Z}^T \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ(w)}) \right\}, \quad (5.10)$$

where $\mathbf{U} = \boldsymbol{\Sigma}^{-1} - \boldsymbol{\Sigma}^{-1} \mathbf{X} (\mathbf{X}^T \boldsymbol{\Sigma}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \boldsymbol{\Sigma}^{-1}$. Applying the reduction function we have $\sigma^{2(w+1)} = \sigma_*^{2(w+1)}$.

Update for d_{gh}

Since $Q_X^{(U)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)})$ is equivalent to $Q_X^{(P)}(\mathbf{d}; \boldsymbol{\kappa}^{(w)})$ the update for d_{gh} is the same as that presented in (5.4).

Update for ϕ_{*ij}

The update for ϕ_{*ij} is the same as that for the REML EM algorithm. Differentiating $Q_X^{(U)}(\sigma_*^2, \boldsymbol{\phi}_*, \boldsymbol{\lambda}; \boldsymbol{\kappa}^{(w)})$ with respect to ϕ_{*ij} we have

$$\begin{aligned} \frac{\partial Q_X^{(U)}(\sigma_*^2, \boldsymbol{\phi}_*, \boldsymbol{\lambda}; \boldsymbol{\kappa}^{(w)})}{\partial \phi_{*ij}} &= -\frac{1}{2} \left[\text{tr} \left(\mathbf{U}_* \frac{\partial \boldsymbol{\Sigma}_*}{\partial \phi_{*ij}} \right) \right. \\ &\quad \left. - \frac{1}{\sigma_*^2} E \left\{ (\mathbf{y}_o - \mathbf{Z} \boldsymbol{\Lambda} \mathbf{f})^T \mathbf{U}_* \frac{\partial \boldsymbol{\Sigma}_*}{\partial \phi_{*ij}} \mathbf{U}_* (\mathbf{y}_o - \mathbf{Z} \boldsymbol{\Lambda} \mathbf{f}) \middle| \mathbf{y}_2; \boldsymbol{\kappa}^{(w)} \right\} \right]. \end{aligned}$$

The updating equation for ϕ_{*ij} depends on the specific form of $\boldsymbol{\Sigma}_*$. If we define the update for ϕ_{*ij} as $\phi_{*ij}^{(w+1)}$ then applying the reduction function we have $\phi_{ij}^{(w+1)} = \phi_{*ij}^{(w+1)}$.

Update for λ

To derive an update for the auxillary parameter vector λ we consider the $b \times b$ real invertible matrix Λ to have elements λ_{gh} and $\lambda = \text{vec}(\Lambda)$. To derive an update for λ we consider the partial derivative of $Q_X^{(U)}(\sigma_*^2, \phi_*, \lambda; \mathcal{K}^{(w)})$ with respect to λ_{gh} , i.e.,

$$\begin{aligned} \frac{\partial Q_X^{(U)}(\sigma_*^2, \phi_*, \lambda; \mathcal{K}^{(w)})}{\partial \lambda_{gh}} &= -\frac{1}{2} \left\{ -2E \left(\mathbf{y}_o^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) \right. \\ &\quad \left. + 2E \left(\mathbf{f}^T \Lambda^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) \right\}. \end{aligned}$$

Equating to zero we have

$$E \left(\mathbf{f}^T \Lambda^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) = E \left(\mathbf{y}_o^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right). \quad (5.11)$$

We can express Λ as

$$\Lambda = \sum_{q=1}^b \sum_{r=1}^b \lambda_{qr} \mathbf{J}_{qr},$$

where $\mathbf{J}_{qr}$ is a $b \times b$ indicator matrix with a unit entry in row q , column r and zeros elsewhere. We can re-write the left hand side of (5.11) as

$$\begin{aligned} E \left(\mathbf{f}^T \Lambda^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) &= E \left(\mathbf{f}^T \sum_{k=1}^b \sum_{l=1}^b \lambda_{kl} \mathbf{J}_{kl}^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) \\ &= \sum_{k=1}^b \sum_{l=1}^b \lambda_{kl} E \left(\mathbf{f}^T \mathbf{J}_{qr}^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right). \end{aligned} \quad (5.12)$$

Substituting (5.12) into (5.11) we have

$$\sum_{q=1}^b \sum_{r=1}^b \lambda_{qr} E \left(\mathbf{f}^T \mathbf{J}_{qr}^T \mathbf{Z}^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right) = E \left(\mathbf{y}_o^T \mathbf{S}_* \mathbf{Z} \frac{\partial \Lambda}{\partial \lambda_{gh}} \mathbf{f} | \mathbf{y}_2; \mathcal{K}^{(w)} \right).$$

The update for λ can be expressed as

$$\boldsymbol{\lambda}^{(w+1)} = \mathbf{A}^{-1(w)} \mathbf{b}^{(w)},$$

where the elements of the $b^2 \times b^2$ matrix $\mathbf{A}^{(w)}$ are

$$a_{gh,qr}^{(w)} = \tilde{\mathbf{u}}^{(w)T} \mathbf{J}_{qr}^T \mathbf{Z}^T \mathbf{S}^{(w)} \mathbf{Z} \frac{\partial \boldsymbol{\Lambda}^{(w)}}{\partial \lambda_{gh}^{(w)}} \tilde{\mathbf{u}}^{(w)} + \text{tr} \left(\mathbf{J}_{qr}^T \mathbf{Z}^T \mathbf{S}^{(w)} \mathbf{Z} \frac{\partial \boldsymbol{\Lambda}^{(w)}}{\partial \lambda_{gh}^{(w)}} \mathbf{C}^{ZZ(w)} \right) \quad (5.13)$$

and the elements of the $b^2 \times 1$ vector $\mathbf{b}^{(w)}$ are

$$b_{gh}^{(w)} = \mathbf{y}_o^T \mathbf{S}^{(w)} \mathbf{Z} \frac{\partial \boldsymbol{\Lambda}^{(w)}}{\partial \lambda_{gh}^{(w)}} \tilde{\mathbf{u}}^{(w)}. \quad (5.14)$$

Note the difference between (5.13), (5.14) and (5.6), (5.7). The latter, based on $\mathbf{y}_c^{(P)}$ requires $\mathbf{C}^{XZ}$ and $\mathbf{C}^{ZZ}$ whereas the former, based on $\mathbf{y}_c^{(U)}$ only requires $\mathbf{C}^{ZZ}$.

We have derived an update for $\boldsymbol{\lambda}$ where the variance-covariance matrix $\mathbf{D}$ is considered to be completely unstructured. When $\mathbf{D}$ is structured the update for $\boldsymbol{\lambda}$ will have a simpler form. For the one-way random effects model for example, $\mathbf{D} = d\mathbf{I}_b$ and therefore $\boldsymbol{\Lambda} = \lambda\mathbf{I}_b$. In this case the update for λ has a relatively simple form.

This completes the M-step.

5.3.4 Rate matrix

The rate matrix for the REML PX EM algorithm using the complete data specification $\mathbf{y}_c^{(U)}$ is the same as that presented in (5.8) except we substitute $\boldsymbol{\mathcal{I}}_c^{(U)}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0)$ for $\boldsymbol{\mathcal{I}}_c^{(P)}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0)$. We define

$$\boldsymbol{\mathcal{I}}_c^{(U)}(\boldsymbol{\kappa}^*, \boldsymbol{\lambda}_0) = \begin{pmatrix} \boldsymbol{\mathcal{I}}_{c\boldsymbol{\kappa},\boldsymbol{\kappa}}^{(U)} & \boldsymbol{\mathcal{I}}_{c\boldsymbol{\kappa},\boldsymbol{\lambda}}^{(U)} \\ \boldsymbol{\mathcal{I}}_{c\boldsymbol{\lambda},\boldsymbol{\kappa}}^{(U)} & \boldsymbol{\mathcal{I}}_{c\boldsymbol{\lambda},\boldsymbol{\lambda}}^{(U)} \end{pmatrix}_{\boldsymbol{\kappa} = \boldsymbol{\kappa}^*, \boldsymbol{\lambda} = \boldsymbol{\lambda}_0}.$$

The elements of $\boldsymbol{\mathcal{I}}_c^{(U)}$, which comprise $\boldsymbol{\mathcal{I}}_{c\sigma^2,\boldsymbol{\lambda}}^{(U)}$, $\boldsymbol{\mathcal{I}}_{c\boldsymbol{\gamma},\boldsymbol{\lambda}}^{(U)}$, and $\boldsymbol{\mathcal{I}}_{c\boldsymbol{\phi},\boldsymbol{\lambda}}^{(U)}$, and the elements of $\boldsymbol{\mathcal{I}}_{c\boldsymbol{\lambda},\boldsymbol{\lambda}}^{(U)}$ are presented below.

Elements of $\mathcal{I}_{c_{\sigma^2, \lambda}}^{(U)}$

For $g = 1, \dots, b$, $h = 1, \dots, b$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\sigma^2, \lambda_{gh}}}^{(U)} = & \frac{1}{(\sigma^{2*})^2} \left\{ \mathbf{y}_o^T \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* - \tilde{\mathbf{u}}^{*T} \mathbf{Z}^T \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* \right. \\ & \left. - \text{tr} \left(\mathbf{Z}^T \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \right\}. \end{aligned}$$

Elements of $\mathcal{I}_{c_{\gamma, \lambda}}^{(U)}$

The elements for $\mathcal{I}_{c_{\gamma, \lambda}}^{(U)}$ are equivalent to $\mathcal{I}_{c_{\gamma, \lambda}}^{(P)}$, i.e., for $g = 1, \dots, b$, $h = 1, \dots, b$ we have

$$\mathcal{I}_{c_{\gamma_{gh}, \lambda_{gh}}}^{(U)} = \tilde{\mathbf{u}}^{*T} \mathbf{G}^{-1*} \frac{\partial \mathbf{G}^*}{\partial \gamma_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* + \text{tr} \left(\mathbf{G}^{-1*} \frac{\partial \mathbf{G}^*}{\partial \gamma_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right).$$

Elements of $\mathcal{I}_{c_{\phi, \lambda}}^{(U)}$

For $i = 1, \dots, n$, $j = 1, \dots, n$, $g = 1, \dots, b$, $h = 1, \dots, b$ it can be shown that

$$\begin{aligned} \mathcal{I}_{c_{\phi_{ij}, \lambda_{gh}}}^{(U)} = & \frac{1}{\sigma^{2*}} \left\{ \mathbf{y}_o^T \mathbf{U}^* \frac{\partial \Sigma^*}{\partial \phi_{ij}} \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* - \tilde{\mathbf{u}}^{*T} \mathbf{Z}^T \mathbf{U}^* \frac{\partial \Sigma^*}{\partial \phi_{ij}} \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* \right. \\ & \left. - \text{tr} \left(\mathbf{Z}^T \mathbf{U}^* \frac{\partial \Sigma^*}{\partial \phi_{ij}} \mathbf{U}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \right\}. \end{aligned}$$

Elements of $\mathcal{I}_{c_{\lambda, \lambda}}^{(U)}$

We make the same assumptions as those made for the elements of $\mathcal{I}_{c_{\lambda, \lambda}}^{(P)}$. For $g = 1, \dots, b$, $h = 1, \dots, b$, $q = 1, \dots, b$, $r = 1, \dots, b$ it can be shown that

$$\begin{aligned}
\mathcal{I}_{c_{\lambda_{gh}, \lambda_{qr}}}^{(U)} &= \tilde{\mathbf{u}}^{\star T} \frac{\partial \Lambda^{\star T}}{\partial \lambda_{qr}^*} \mathbf{Z}^T \mathbf{S}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \tilde{\mathbf{u}}^* + \text{tr} \left(\frac{\partial \Lambda^{\star T}}{\partial \lambda_{qr}^*} \mathbf{Z}^T \mathbf{S}^* \mathbf{Z} \frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \mathbf{C}^{ZZ*} \right) \\
&+ \text{tr} \left(\frac{\partial \Lambda^*}{\partial \lambda_{gh}^*} \Lambda^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \Lambda^{-1*} \right) + \tilde{\mathbf{u}}^{\star T} \frac{\partial \Lambda^{\star T}}{\partial \lambda_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \tilde{\mathbf{u}}^* \\
&+ \text{tr} \left(\frac{\partial \Lambda^{\star T}}{\partial \lambda_{gh}^*} \mathbf{G}^{-1*} \frac{\partial \Lambda^*}{\partial \lambda_{qr}^*} \mathbf{C}^{ZZ*} \right).
\end{aligned}$$

5.4 Example: lamb weight data

We apply a REML PX EM algorithm using the complete data specifications $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ to the lamb weight data presented in Table 4.1. We fit the same model as that fitted in the previous chapter for the REML EM algorithm, i.e., the overall mean, age, and line are fitted as fixed effects, and sire is fitted as a random effect. For the REML PX EM algorithm this model can be written in matrix notation as

$$\mathbf{y}_o = \mathbf{X}\boldsymbol{\beta}_* + \mathbf{Z}\boldsymbol{\Lambda}\mathbf{f} + \boldsymbol{\epsilon},$$

where $\mathbf{y}_o$, $\mathbf{X}$, and $\mathbf{Z}$ are defined in the same way as they were in the previous chapter for the REML EM algorithm. The vector $\boldsymbol{\beta}_*$ is defined in the same way as $\boldsymbol{\beta}$. We assume that the 23×1 vector of rescaled random effects $\mathbf{f}$, and the 62×1 vector of residuals $\boldsymbol{\epsilon}$, are independent of each other and have the following distributions

$$\begin{aligned}
\mathbf{f} &\sim N(\mathbf{0}, d\mathbf{I}_{23}), \\
\boldsymbol{\epsilon} &\sim N(\mathbf{0}, \sigma_*^2 \mathbf{I}_{62}).
\end{aligned}$$

The updating equation for σ_*^2 is provided in (5.3) and (5.10) for $\mathbf{y}_c^{(P)}$ and $\mathbf{y}_c^{(U)}$ respectively. The updating equation for d is provided in (5.4) and applies to both complete data specifications. We define $\boldsymbol{\kappa} = (\sigma^2, \sigma_s^2)^T$ as the parameter vector of interest and $\boldsymbol{\mathcal{K}} = (\sigma_*^2, d, \lambda)^T$ as the expanded parameter vector. The reduction function is $\boldsymbol{\kappa} = \mathcal{R}(\boldsymbol{\mathcal{K}}) = (\sigma_*^2, \lambda^2 d)^T$. We use an equivalent set of starting values to those in the previous chapter, i.e., the first set of starting values are $\boldsymbol{\mathcal{K}}^{(0)} = (\sigma_*^{2(0)}, d^{(0)}, \lambda^{(0)})^T = (1, 0.01, 1)^T$ and the second set of starting values are $\boldsymbol{\mathcal{K}}^{(0)} = (\sigma_*^{2(0)}, d^{(0)}, \lambda^{(0)})^T = (1, 5, 1)^T$. The convergence criteria presented in (4.18)

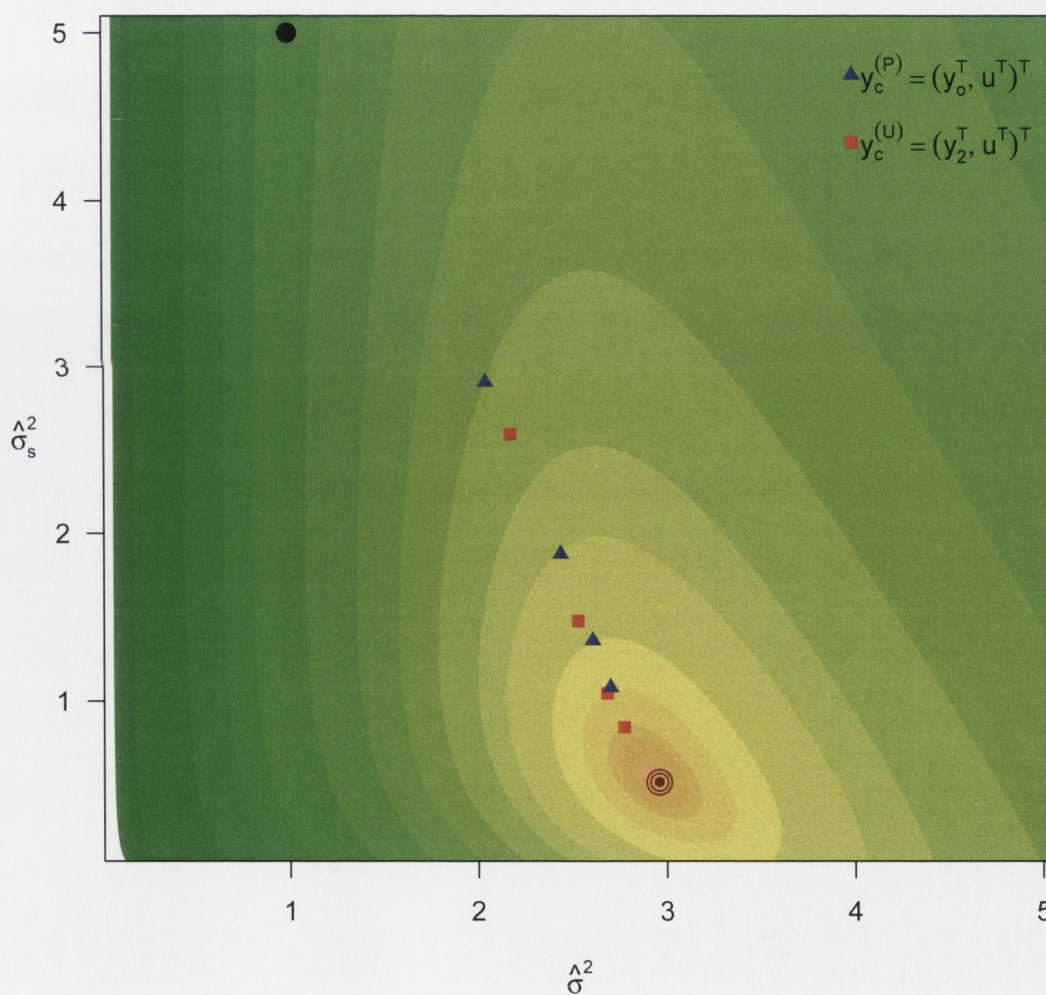
was used to determine if the sequence of iterates $\boldsymbol{\kappa}^{(0)}, \boldsymbol{\kappa}^{(1)}, \dots$ converged to a (local) maximum of the REML log-likelihood function. The statistical software package R (R Development Core Team, 2011) was used to implement both REML PX EM algorithms. REML estimates of σ^2 and σ_s^2 , denoted σ^{2*} and σ_s^{2*} , the number of iterations until convergence and the rate of convergence for both complete data specifications and both sets of starting values are provided in Table 5.1.

Table 5.1: Results for two REML PX EM algorithms applied to the lamb weight data.

Complete data	σ^{2*}	σ_s^{2*}	Iterations	Rate of convergence
$\sigma_*^{2(0)} = 1, d^{(0)} = 0.01, \lambda^{(0)} = 1$				
$\mathbf{y}_c^{(P)}$	2.961597	0.5170765	83	0.8167880
$\mathbf{y}_c^{(U)}$	2.961597	0.5170765	57	0.7434961
$\sigma_*^{2(0)} = 1, d^{(0)} = 5, \lambda^{(0)} = 1$				
$\mathbf{y}_c^{(P)}$	2.961597	0.5170767	78	0.8167879
$\mathbf{y}_c^{(U)}$	2.961597	0.5170767	55	0.7434961

The observed information matrix is exactly the same as that obtained using a REML EM algorithm. A comparison of the results in Table 4.2 and Table 5.1 show that the REML PX EM converges in far fewer iterations than the REML EM algorithm. Table 5.1 highlights the benefit in using the complete data specification $\mathbf{y}_c^{(U)}$ compared to $\mathbf{y}_c^{(P)}$. The rate of convergence is noticeably less and convergence is achieved in approximately 30% fewer iterations for both sets of starting values. A contour plot of the REML log-likelihood surface and the first four iterations of the REML PX EM algorithms using the complete data specification $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$ is provided in Figure 5.1. This figure illustrates the faster rate of convergence associated with using the complete data specification $\mathbf{y}_c^{(U)}$. The complete data specification $\mathbf{y}_c^{(U)}$ makes a larger initial step towards the REML solution and its' third iterate is approximately equivalent to the fourth iterate of the REML PX EM algorithm using the complete data specification $\mathbf{y}_c^{(P)}$. Figure 5.1 also highlights another feature of EM algorithms in general, which is that they often get close to a solution quite quickly but then take a long time to converge.

Figure 5.1: Contour plot of the residual log-likelihood surface for the lamb weight data and the first 4 iterations of the REML PX EM algorithms using the complete data specification $\mathbf{y}_c^{(U)}$ and $\mathbf{y}_c^{(P)}$. The starting values of $\sigma^{2(0)} = 1$ and $\sigma_s^{2(0)} = 5$ are represented by the black circle. The REML estimates of $\sigma^{2*} = 2.962$ and $\sigma_s^{2*} = 0.517$ are represented by the “target”.



Chapter 6

Two model specifications for a factor analytic mixed model

In this and the next two chapters we consider REML EM and REML PX EM algorithms for factor analytic mixed models. The difference between linear mixed models and factor analytic mixed models is that in the former the vector of random effects is assumed to have a Gaussian distribution with a mean of zero and a variance defined by a symmetric positive definite variance-covariance matrix. In the later the vector of random effects is modelled using factor analysis. Our aim in the next three chapters is to consider ways in which to reduce the number of iterations until convergence for REML EM (and REML PX EM) algorithms applied to factor analytic mixed models. In all three chapters we consider the efficacy of using a less strict convergence criteria. In this chapter we consider two different model specifications, and in the next two chapters we consider a new parameter expansion and an alternative missing data specification respectively.

The motivation for the work presented in this chapter and the next two chapters is the analysis of multi-environment trial data produced by plant breeding programmes. Factor analytic mixed models were recommended by Smith et al. (2001) for the analysis of this type of data. In Australia, a major source of multi-environment trial data is the National Variety Trial programme which aims to deliver robust, independent results on the performance of recently released cereal grain varieties. The measure most often used to rank performance is grain yield. This programme comprises large numbers of varieties grown over large numbers of trials. As an example of the numbers involved, the 2005 to 2010 non-Durum wheat data

comprises 538 varieties and 928 trials where an “average” trial contains 162 plots. This data set therefore contains approximately 150000 measurements of grain yield. Fitting a factor analytic mixed model of the type advocated by Smith et al. (2001) to a data set of this size and structure is currently infeasible.

Although not applied to multi-environment trial data comprising tens of thousands of records, factor analytic mixed models are routinely applied to multi-environment trial data which contain thousands of records. For example, Thompson et al. (2003) apply a factor analytic mixed model to multi-environment trial data comprising 7613 records, 62 trials, and 3300 variety by trial combinations. One of the challenges in fitting a factor analytic mixed model to data of this size and structure is that modern statistical software packages, such as ASReml-R (Butler et al., 2007), typically implement a Newton-Raphson type algorithm, such as the average information algorithm (Gilmour et al., 1995), and successive iterations of these algorithms don’t always increase the residual log-likelihood function, parameter updates don’t always remain in their parameter space, and the algorithm can fail to converge. A potential solution to these problems is to use a REML EM or REML PX EM algorithm.

For a REML EM or REML PX EM algorithm to be a practical option for estimating the variance parameters associated with a factor analytic mixed model applied to multi-environment trial data the number of iterations until convergence has to be at a manageable level as each iteration can be computationally expensive in terms of computing time. Depending on the size and structure of the data this might mean having an algorithm that converges in no more than 50 iterations and preferably less than 25 iterations.

In this chapter we consider two different model specifications for a factor analytic mixed model. The first is based on the most commonly employed model specification, where the k factors are assumed to have unit variance. The second model specification is closely related to the PX EM algorithm model for a factor analytic mixed model presented by Liu et al. (1998). In this specification the variance-covariance matrix associated with the k factors is assumed to be unstructured. The main aim of this chapter is to show, using an example, that large gains in the number of iterations until convergence using a REML EM algorithm for factor analytic mixed models can be made by choosing a different model specification to that most commonly employed.

6.1 Model specification

We consider a simplified plant breeding case where we have m varieties grown in p trials. Every variety is grown in every trial and there is r replicates of every variety within a trial, i.e., the total number of observations $n = rmp$. Consider the model

$$\text{Model } \mathcal{O} \quad \mathbf{y}_o = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}_g\mathbf{u}_g + \mathbf{e},$$

where $\boldsymbol{\beta}$ is a $p \times 1$ vector of fixed effects representing trial means and $\mathbf{X}$ is its associated $n \times p$ design matrix of full column rank. It is assumed that the $n \times 1$ vector of residuals $\mathbf{e}$ is distributed $N(\mathbf{0}, \mathbf{R})$ where $\mathbf{R} = \sigma^2 \mathbf{I}_n$. The $n \times mp$ design matrix $\mathbf{Z}_g$ is associated with the $mp \times 1$ vector $\mathbf{u}_g$ of variety by trial effects. These effects are assumed to be correlated within a trial but are themselves independent, i.e.,

$$V(\mathbf{u}_g) = \mathbf{G}_e \otimes \mathbf{I}_m,$$

where $\mathbf{G}_e$ is a $p \times p$ symmetric positive definite matrix and is commonly referred to as the genetic variance-covariance matrix. The diagonal elements of $\mathbf{G}_e$ are the genetic variances for each trial and off-diagonal elements are genetic covariances between pairs of trials. A completely general or unstructured form of $\mathbf{G}_e$ involves $p(p+1)/2$ parameters. For even moderately large values of p it is computationally difficult to estimate unstructured $\mathbf{G}_e$. A more parsimonious representation of $\mathbf{G}_e$ can be achieved using factor analysis to model the variety effects within each trial, i.e.,

$$\begin{aligned} \mathbf{u}_g &= (\boldsymbol{\lambda}_1 \otimes \mathbf{I}_m)\mathbf{f}_1 + \dots + (\boldsymbol{\lambda}_k \otimes \mathbf{I}_m)\mathbf{f}_k + \boldsymbol{\delta} \\ &= (\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} + \boldsymbol{\delta}, \end{aligned}$$

where $\boldsymbol{\Lambda}$ is a $p \times k$ matrix of loadings and $\mathbf{f}$ is a $mk \times 1$ vector of varietal scores that can be partitioned as $\mathbf{f} = (\mathbf{f}_1^T, \dots, \mathbf{f}_k^T)^T$. The matrix of factor loadings $\boldsymbol{\Lambda}$ has elements λ_{jr} , $j = 1, \dots, p$, $r = 1, \dots, k$. To ensure the uniqueness of the estimates of factor loadings when $k > 1$, $k(k-1)/2$ independent constraints on the elements of $\boldsymbol{\Lambda}$ are necessary. The constraints to be used are arbitrary. However, Smith et al. (2001) recommend the constraint $\lambda_{jr} = 0$ for $j < r = 2, \dots, k$.

The vector $\boldsymbol{\delta}$ is a $mp \times 1$ vector of residuals that can be partitioned with respect to trials, i.e., $\boldsymbol{\delta} = (\boldsymbol{\delta}_1^T, \dots, \boldsymbol{\delta}_p^T)^T$. It is assumed that $\mathbf{f}$ and $\boldsymbol{\delta}$ are independent of each other and that

$$\boldsymbol{\delta} \sim N(\mathbf{0}, \boldsymbol{\Psi} \otimes \mathbf{I}_m),$$

where $\boldsymbol{\Psi} = \boldsymbol{\Psi}(\boldsymbol{\psi})$ is a $p \times p$ diagonal matrix. The genetic variance-covariance matrix is therefore

$$\mathbf{G}_e = \boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T + \boldsymbol{\Psi}.$$

The vector of variance parameters $\boldsymbol{\psi}$ is commonly referred to as the specific variances. For simplicity we will assume that $\boldsymbol{\Psi} = \psi \mathbf{I}_p$, i.e., there is a common specific variance.

The two model specifications we consider differ in their assumption regarding the variance of $\mathbf{f}$. Using the same notation as Liu et al. (1998) we refer to the first model as Model $\mathcal{O}$ where it is assumed that

$$\mathbf{f} \sim N(\mathbf{0}, \mathbf{I}_{mk}).$$

We refer to the second model as Model $\mathcal{M}$ where it is assumed that

$$\mathbf{f} \sim N(\mathbf{0}, \mathbf{D} \otimes \mathbf{I}_m),$$

where $\mathbf{D} = \mathbf{D}(\mathbf{d})$ is a $k \times k$ symmetric positive definite matrix and $\mathbf{d}$ is a vector of variance parameters.

Model $\mathcal{O}$ is a special case of Model $\mathcal{M}$, namely, when $\mathbf{D} = \mathbf{I}_k$. Therefore, from here on we only present the development of a REML EM algorithm for Model $\mathcal{M}$ as the complete data density function and the E and M-steps for Model $\mathcal{O}$ can be obtained by substituting an identity matrix of order k for $\mathbf{D}$.

For Model $\mathcal{M}$ the distribution of the observed data is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}, \mathbf{H}),$$

where $\mathbf{H} = \mathbf{Z}_g(\mathbf{G}_e \otimes \mathbf{I}_m)\mathbf{Z}_g^T + \mathbf{R}$ where $\mathbf{R} = \sigma^2\mathbf{I}_n$ and the vector of variance parameters to be estimated is $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \psi, \mathbf{d}^T)^T$. The estimate of the genetic variance-covariance matrix, denoted $\hat{\mathbf{G}}_e$, and the best linear unbiased predictor of variety by trial effects, denoted $\tilde{\mathbf{u}}_g$, are the quantities most often reported to plant breeders. The former for providing the genetic correlation between trials and the later for determining which varieties are broadly adaptable across environments.

It should be noted that $\boldsymbol{\Lambda}$ and $\mathbf{D}$ are not unique for Model $\mathcal{M}$. We don't consider this an impediment to the implementation of Model $\mathcal{M}$ for the analysis of plant breeding data as $\hat{\mathbf{G}}_e$ and $\tilde{\mathbf{u}}_g$ are unique and are the quantities of primary interest.

6.2 Complete data density function

We define the complete data to be $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$ where the missing data vector is $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$. Model $\mathcal{O}$ and Model $\mathcal{M}$ could be considered new in the sense that we consider the incomplete data to be the transformed observed data vector associated with the marginal log-density function in the Verbyla (1990) conditional derivation of REML. This has not been considered previously for REML EM or REML PX EM algorithms applied to factor analytic mixed models.

To form the complete data log-density function requires the joint distribution

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{f} \\ \boldsymbol{\delta} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2 & \mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Lambda} \mathbf{D} \otimes \mathbf{I}_m) & \mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Psi} \otimes \mathbf{I}_m) \\ (\mathbf{D} \boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{L}_2 & \mathbf{D} \otimes \mathbf{I}_m & \mathbf{0} \\ (\boldsymbol{\Psi} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{L}_2 & \mathbf{0} & \boldsymbol{\Psi} \otimes \mathbf{I}_m \end{pmatrix} \right\} \quad (6.1)$$

The complete data log-density function can be expressed as

$$\log\{f(\mathbf{y}_c; \boldsymbol{\Gamma})\} = \log\{f(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\} + \log\{f(\mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\}.$$

Consider $\log\{f(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\}$. Using (6.1) and Result A.1 it can be shown that

$$\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta} \sim N\{\mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{f} + \mathbf{L}_2^T \mathbf{Z}_g \boldsymbol{\delta}, \mathbf{L}_2^T \mathbf{R} \mathbf{L}_2\},$$

and the log-density function is

$$\begin{aligned} \log\{f(\mathbf{y}_2|\mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} \right. \\ &\quad \left. + \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} - \mathbf{Z}_g\boldsymbol{\delta}\}^T \mathbf{S} \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} - \mathbf{Z}_g\boldsymbol{\delta}\} \right], \end{aligned}$$

where

$$\begin{aligned} \mathbf{S} &= \mathbf{L}_2(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \\ &= \mathbf{R}^{-1} - \mathbf{R}^{-1} \mathbf{X}(\mathbf{X}^T \mathbf{R}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}^{-1}. \end{aligned}$$

The marginal log-density function for the missing data is

$$\begin{aligned} \log\{f(\mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\} &= \log\{f(\mathbf{f}; \mathbf{d})\} + \log\{f(\boldsymbol{\delta}; \boldsymbol{\psi})\} \\ &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\boldsymbol{\Psi} \otimes \mathbf{I}_m)\} \right. \\ &\quad \left. + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \boldsymbol{\delta}^T (\boldsymbol{\Psi}^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right]. \end{aligned}$$

The complete data log-density function is therefore

$$\begin{aligned} \log\{f(\mathbf{y}_c; \boldsymbol{\Gamma})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} + \log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\boldsymbol{\Psi} \otimes \mathbf{I}_m)\} \right. \\ &\quad + \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} - \mathbf{Z}_g\boldsymbol{\delta}\}^T \mathbf{S} \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} - \mathbf{Z}_g\boldsymbol{\delta}\} \\ &\quad \left. + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \boldsymbol{\delta}^T (\boldsymbol{\Psi}^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right]. \end{aligned} \tag{6.2}$$

6.3 REML EM algorithm E-step for Model $\mathcal{M}$

The E-step involves taking the conditional expectation of (6.2) over $k(\mathbf{f}, \boldsymbol{\delta}|\mathbf{y}_2; \boldsymbol{\Gamma}^{(w)})$ where $\boldsymbol{\Gamma}^{(w)}$ is the w -th iterate of $\boldsymbol{\Gamma}$. We can write

$$\begin{aligned} Q(\boldsymbol{\Gamma}; \boldsymbol{\Gamma}^{(w)}) &= E[\log\{f(\mathbf{y}_c; \boldsymbol{\Gamma})\}|\mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] \\ &= E[\log\{f(\mathbf{y}_2|\mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\Gamma})\}|\mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] + E[\log\{f(\mathbf{f}; \mathbf{d})\}|\mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] \\ &\quad + E[\log\{f(\boldsymbol{\delta}; \boldsymbol{\psi})\}|\mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] \\ &= Q(\sigma^2, \boldsymbol{\lambda}; \boldsymbol{\Gamma}^{(w)}) + Q(\mathbf{d}; \boldsymbol{\Gamma}^{(w)}) + Q(\boldsymbol{\psi}; \boldsymbol{\Gamma}^{(w)}). \end{aligned} \tag{6.3}$$

To compute the $Q(\cdot)$ functions in (6.3) requires the joint conditional distribution of $\mathbf{f}$ and $\boldsymbol{\delta}$ given $\mathbf{y}_2$. Using (6.1) and Result A.1 we have

$$\begin{pmatrix} \mathbf{f} \\ \boldsymbol{\delta} \end{pmatrix} \middle| \mathbf{y}_2 \sim N \left\{ \begin{pmatrix} (\mathbf{D}\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{y}_o \\ (\boldsymbol{\Psi} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{y}_o \end{pmatrix}, \begin{pmatrix} \mathbf{T}_{ff} & \mathbf{T}_{f\delta} \\ \mathbf{T}_{\delta f} & \mathbf{T}_{\delta\delta} \end{pmatrix} \right\},$$

where

$$\begin{aligned} \mathbf{T}_{ff} &= (\mathbf{D} \otimes \mathbf{I}_m) - (\mathbf{D}\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\boldsymbol{\Lambda} \mathbf{D} \otimes \mathbf{I}_m), \\ \mathbf{T}_{f\delta} &= \mathbf{T}_{\delta f}^T = -(\boldsymbol{\Psi} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\boldsymbol{\Lambda} \mathbf{D} \otimes \mathbf{I}_m), \\ \mathbf{T}_{\delta\delta} &= (\boldsymbol{\Psi} \otimes \mathbf{I}_m) - (\boldsymbol{\Psi} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\boldsymbol{\Psi} \otimes \mathbf{I}_m) \end{aligned}$$

and

$$\begin{aligned} \mathbf{P} &= \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \\ &= \mathbf{H}^{-1} - \mathbf{H}^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{H}^{-1}. \end{aligned}$$

We can write

$$\begin{aligned} Q(\sigma^2, \boldsymbol{\lambda}; \boldsymbol{\Gamma}^{(w)}) &= -\frac{1}{2} \left[\log \{ \det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2) \} \right. \\ &\quad + \{ \mathbf{y}_o - \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \}^T \mathbf{S} \{ \mathbf{y}_o - \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \} \\ &\quad + \text{tr} \left\{ (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \right\} + \text{tr} \left\{ \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)} \right\} \\ &\quad \left. - 2 \text{tr} \left\{ \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}) \right\} \right], \\ Q(\mathbf{d}; \boldsymbol{\Gamma}^{(w)}) &= -\frac{1}{2} \left[\log \{ \det(\mathbf{D} \otimes \mathbf{I}_m) \} + \tilde{\mathbf{f}}^{(w)T} (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr}(\mathbf{T}_{ff}^{(w)}) \right], \\ Q(\boldsymbol{\psi}; \boldsymbol{\Gamma}^{(w)}) &= -\frac{1}{2} \left[\log \{ \det(\boldsymbol{\Psi} \otimes \mathbf{I}_m) \} + \tilde{\boldsymbol{\delta}}^{(w)T} (\boldsymbol{\Psi}^{-1} \otimes \mathbf{I}_m) \tilde{\boldsymbol{\delta}}^{(w)} + \text{tr}(\mathbf{T}_{\delta\delta}^{(w)}) \right], \end{aligned}$$

where

$$\begin{aligned} \tilde{\mathbf{f}}^{(w)} &= (\mathbf{D}^{(w)} \boldsymbol{\Lambda}^{(w)T} \otimes \mathbf{I}_m) \mathbf{P}^{(w)} \mathbf{y}_o, \\ E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}) &= \tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)} \quad \text{using Result A.2.} \end{aligned}$$

This completes the E-step.

6.4 REML EM algorithm M-step for Model $\mathcal{M}$

The M-step involves maximising $Q(\mathbf{\Gamma}; \mathbf{\Gamma}^{(w)})$ with respect to $\mathbf{\Gamma}$ to find an update for $\mathbf{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \boldsymbol{\psi}, \mathbf{d}^T)^T$. The partitioning of $Q(\mathbf{\Gamma}; \mathbf{\Gamma}^{(w)})$ into $Q(\sigma^2, \boldsymbol{\lambda}; \mathbf{\Gamma}^{(w)})$, $Q(\mathbf{d}; \mathbf{\Gamma}^{(w)})$ and $Q(\boldsymbol{\psi}; \mathbf{\Gamma}^{(w)})$ simplifies the process.

6.4.1 Update for σ^2

Differentiating $Q(\sigma^2, \boldsymbol{\lambda}; \mathbf{\Gamma}^{(w)})$ with respect to σ^2 we have

$$\begin{aligned} \frac{\partial Q(\sigma^2, \boldsymbol{\lambda}; \mathbf{\Gamma}^{(w)})}{\partial \sigma^2} = & -\frac{1}{2} \left[\frac{n-p}{\sigma^2} \right. \\ & - \frac{1}{(\sigma^2)^2} \{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \}^T \mathbf{K} \{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \} \\ & - \frac{1}{(\sigma^2)^2} \text{tr} \left\{ (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \right\} \\ & \left. - \frac{1}{(\sigma^2)^2} \text{tr} \left(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)} \right) - \frac{2}{(\sigma^2)^2} \text{tr} \left\{ \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}) \right\} \right], \end{aligned}$$

where $\mathbf{K} = \mathbf{I}_n - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$. Equating to zero and solving for σ^2 we have

$$\begin{aligned} \sigma^{2(w+1)} = & \frac{1}{n-p} \left[\{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \}^T \mathbf{K} \{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \} \right. \\ & + \text{tr} \left\{ (\boldsymbol{\Lambda}^{(w)T} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \right\} \\ & \left. + \text{tr} \left(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)} \right) + 2 \text{tr} \left\{ \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}) \right\} \right]. \end{aligned}$$

6.4.2 Update for λ_{jr}

We denote the elements of the $p \times k$ matrix of factor loadings $\boldsymbol{\Lambda}$ as λ_{jr} , $j = 1, \dots, p$, $r = 1, \dots, k$ and $\boldsymbol{\lambda} = \text{vec}(\boldsymbol{\Lambda})$. Finding an update for $\boldsymbol{\lambda}$ is similar to finding an update for the auxillary parameter in the model presented in (5.1).

We can write $Q(\sigma^2, \boldsymbol{\lambda}; \Gamma^{(w)})$ as

$$\begin{aligned} Q(\sigma^2, \boldsymbol{\lambda}; \Gamma^{(w)}) = & -\frac{1}{2} \left[\log \{ \det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2) \} + \mathbf{y}_o^T \mathbf{S} \mathbf{y}_o \right. \\ & - 2E \{ \mathbf{y}_o^T \mathbf{S} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \} \\ & - 2E \{ \boldsymbol{\delta}^T \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \} \\ & \left. + E \{ \mathbf{f}^T (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \} \right]. \end{aligned}$$

Differentiating $Q(\sigma^2, \boldsymbol{\lambda}; \Gamma^{(w)})$ with respect to λ_{jr} we have

$$\begin{aligned} \frac{\partial Q(\sigma^2, \boldsymbol{\lambda}; \Gamma^{(w)})}{\partial \lambda_{jr}} = & -\frac{1}{2} \left[-2E \left\{ \mathbf{y}_o^T \mathbf{S} \mathbf{Z}_g \left(\frac{\partial \boldsymbol{\Lambda}}{\partial \lambda_{jr}} \otimes \mathbf{I}_m \right) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \right\} \right. \\ & + 2\text{tr} \left\{ \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g \left(\frac{\partial \boldsymbol{\Lambda}}{\partial \lambda_{jr}} \otimes \mathbf{I}_m \right) E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \Gamma^{(w)}) \right\} \\ & \left. + 2E \left\{ \mathbf{f}^T (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g \left(\frac{\partial \boldsymbol{\Lambda}}{\partial \lambda_{jr}} \otimes \mathbf{I}_m \right) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \right\} \right]. \end{aligned}$$

Equating to zero and using the identity

$$\boldsymbol{\Lambda}^T = \sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} \mathbf{J}_{lq}^T$$

where $\mathbf{J}_{lq}$ is a $p \times k$ indicator matrix having a 1 in row l , column q and zeroes elsewhere we obtain

$$\begin{aligned} & \sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} E \{ \mathbf{f}^T (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \} \\ & = \sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} \left[E \{ \mathbf{y}_o^T \mathbf{S} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)} \} \right. \\ & \quad \left. - \text{tr} \{ \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) E(\mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \Gamma^{(w)}) \} \right], \end{aligned}$$

where $\mathbf{J}_{jr}$ is a $p \times k$ indicator matrix having a 1 in row j , column r , and zeroes elsewhere. This equation can be written as

$$\sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} a_{jr,lq}^{(w)} = b_{jr}^{(w)}$$

and in matrix notation as

$$\mathbf{A}^{(w)} \boldsymbol{\lambda} = \mathbf{b}^{(w)},$$

where $\mathbf{A}^{(w)}$ is a $pk \times pk$ matrix with elements

$$\begin{aligned} a_{jr,lq}^{(w)} &= E\{\mathbf{f}^T(\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)}\} \\ &= \tilde{\mathbf{f}}^{(w)T} (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} \\ &\quad + \text{tr}\{(\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)}\} \end{aligned}$$

and $\mathbf{b}^{(w)}$ is a $pk \times 1$ vector with elements

$$\begin{aligned} b_{jr}^{(w)} &= E\{\mathbf{y}_o^T \mathbf{S} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma^{(w)}\} - \text{tr}\{\mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) E(\mathbf{f} \mathbf{f}^T | \mathbf{y}_2; \Gamma^{(w)})\} \\ &= \mathbf{y}_o^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \text{tr}\{\mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)})\}. \end{aligned}$$

The update for $\boldsymbol{\lambda}$ is therefore

$$\boldsymbol{\lambda}^{(w+1)} = \mathbf{A}^{(w)-1} \mathbf{b}^{(w)}.$$

6.4.3 Update for ψ

Assuming $\boldsymbol{\Psi} = \psi \mathbf{I}_p$ we can write $Q(\boldsymbol{\psi}; \Gamma^{(w)})$ as

$$Q(\boldsymbol{\psi}; \Gamma^{(w)}) = -\frac{1}{2} \left\{ mp \log(\psi) + \frac{1}{\psi} \tilde{\boldsymbol{\delta}}^{(w)T} \tilde{\boldsymbol{\delta}}^{(w)} + \frac{1}{\psi} \text{tr}(\mathbf{T}_{\delta\delta}^{(w)}) \right\}.$$

Differentiating with respect to ψ , equating to zero, and solving for ψ we have

$$\psi^{(w+1)} = \frac{1}{mp} \left\{ \tilde{\boldsymbol{\delta}}^{(w)T} \tilde{\boldsymbol{\delta}}^{(w)} + \text{tr}(\mathbf{T}_{\delta\delta}^{(w)}) \right\}.$$

6.4.4 Update for d_{qr}

The $k \times k$ matrix $\mathbf{D}$ is symmetric with elements d_{qr} , $q = 1, \dots, k$, $r = 1, \dots, k$, $q \geq r$, i.e., due to the symmetry of $\mathbf{D}$ we only consider the lower left triangle of $\mathbf{D}$. We define the elements of $\mathbf{D}^{-1}$ to be d^{qr} , $q = 1, \dots, k$, $r = 1, \dots, k$, $q \geq r$. Differentiating $Q(\mathbf{d}; \Gamma^{(w)})$ with respect to d_{qr} we have

$$\begin{aligned} \frac{\partial Q(\mathbf{d}; \Gamma^{(w)})}{\partial d_{qr}} &= -\frac{1}{2} \left[\text{tr} \left\{ \left(\mathbf{D}^{-1} \otimes \mathbf{I}_m \right) \left(\frac{\partial \mathbf{D}}{\partial d_{qr}} \otimes \mathbf{I}_m \right) \right\} \right. \\ &\quad \left. - \tilde{\mathbf{f}}^{(w)T} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \mathbf{D}^{-1} \otimes \mathbf{I}_m \right) \tilde{\mathbf{f}}^{(w)} \right. \\ &\quad \left. - \text{tr} \left\{ \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \mathbf{D}^{-1} \otimes \mathbf{I}_m \right) \mathbf{T}_{ff}^{(w)} \right\} \right] \\ &= -\frac{1}{2} \left[\text{tr} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \otimes \mathbf{I}_m \right) - \text{tr} \left\{ \mathbf{V}^{(w)} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \mathbf{D}^{-1} \otimes \mathbf{I}_m \right) \right\} \right. \\ &\quad \left. - \text{tr} \left\{ \mathbf{T}_{ff}^{(w)} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \mathbf{D}^{-1} \otimes \mathbf{I}_m \right) \right\} \right], \end{aligned} \quad (6.4)$$

where $\mathbf{V}^{(w)} = \tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{f}}^{(w)T}$ is symmetric and can be partitioned

$$\mathbf{V}^{(w)} = \begin{pmatrix} \mathbf{V}_{11}^{(w)} & & \\ \vdots & \ddots & \\ \mathbf{V}_{k1}^{(w)} & \dots & \mathbf{V}_{kk}^{(w)} \end{pmatrix} = \begin{pmatrix} \tilde{\mathbf{f}}_1^{(w)} \tilde{\mathbf{f}}_1^{(w)T} & & \\ \vdots & \ddots & \\ \tilde{\mathbf{f}}_k^{(w)} \tilde{\mathbf{f}}_1^{(w)T} & \dots & \tilde{\mathbf{f}}_k^{(w)} \tilde{\mathbf{f}}_k^{(w)T} \end{pmatrix},$$

where each $\mathbf{V}_{gh}^{(w)}$ is a $m \times m$ matrix. Working with the first term in (6.4) we have

$$\begin{aligned} \text{tr} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \otimes \mathbf{I}_m \right) &= m \text{tr} \left(\mathbf{D}^{-1} \frac{\partial \mathbf{D}}{\partial d_{qr}} \right) \\ &= m d^{qr}. \end{aligned}$$

Working with the second term in (6.4) it can be shown that

$$\begin{aligned}
 \text{tr}\left\{\mathbf{V}^{(w)}\left(\mathbf{D}^{-1}\frac{\partial\mathbf{D}}{\partial d_{qr}}\mathbf{D}^{-1}\otimes\mathbf{I}_m\right)\right\} &= \sum_{g=1}^k\sum_{h=1}^k\text{tr}(\mathbf{V}_{gh}^{(w)}d^{gq}d^{rh}) \\
 &= \sum_{g=1}^k\sum_{h=1}^kd^{gq}\text{tr}(\mathbf{V}_{gh}^{(w)})d^{rh} \\
 &= \sum_{g=1}^k\sum_{h=1}^kd^{gq}\delta_{gh}^{(w)}d^{rh},
 \end{aligned}$$

where $\delta_{gh}^{(w)} = \text{tr}(\mathbf{V}_{gh}^{(w)})$ are the elements of the $k \times k$ matrix $\mathbf{\Delta}^{(w)}$. Working with the third term and noting that $\mathbf{T}_{ff}^{(w)}$ can be partitioned

$$\mathbf{T}_{ff}^{(w)} = \begin{pmatrix} \mathbf{T}_{ff_{11}}^{(w)} & & \\ \vdots & \ddots & \\ \mathbf{T}_{ff_{k1}}^{(w)} & \dots & \mathbf{T}_{ff_{kk}}^{(w)} \end{pmatrix},$$

where each $\mathbf{T}_{ff_{gh}}^{(w)}$ is a $m \times m$ matrix, we have

$$\begin{aligned}
 \text{tr}\left\{\mathbf{T}_{ff}^{(w)}\left(\mathbf{D}^{-1}\frac{\partial\mathbf{D}}{\partial d_{qr}}\mathbf{D}^{-1}\otimes\mathbf{I}_m\right)\right\} &= \sum_{g=1}^k\sum_{h=1}^k\text{tr}(\mathbf{T}_{ff_{gh}}^{(w)}d^{gq}d^{rh}) \\
 &= \sum_{g=1}^k\sum_{h=1}^kd^{gq}\text{tr}(\mathbf{T}_{ff_{gh}}^{(w)})d^{rh} \\
 &= \sum_{g=1}^k\sum_{h=1}^kd^{gq}\xi_{gh}^{(w)}d^{rh},
 \end{aligned}$$

where $\xi_{gh}^{(w)} = \text{tr}(\mathbf{T}_{ff_{gh}}^{(w)})$ are the elements of the $k \times k$ matrix $\mathbf{\Xi}^{(w)}$. Equating (6.4) to zero we have

$$m d^{qr} = \sum_{g=1}^k\sum_{h=1}^kd^{gq}\delta_{gh}^{(w)}d^{rh} + \sum_{g=1}^k\sum_{h=1}^kd^{gq}\xi_{gh}^{(w)}d^{rh},$$

which can be written in matrix notation as

$$m\mathbf{D}^{-1} = \mathbf{D}^{-1}\mathbf{\Delta}^{(w)}\mathbf{D}^{-1} + \mathbf{D}^{-1}\mathbf{\Xi}^{(w)}\mathbf{D}^{-1}.$$

Pre and post multiplying by $\mathbf{D}$ we have

$$m\mathbf{D} = \mathbf{\Delta}^{(w)} + \mathbf{\Xi}^{(w)}.$$

The update for d_{qr} is therefore

$$md_{qr}^{(w+1)} = \delta_{qr}^{(w)} + \xi_{qr}^{(w)}.$$

Using the identity

$$\delta_{qr}^{(w)} = \text{tr}(\mathbf{V}_{qr}^{(w)}) = \text{tr}(\tilde{\mathbf{f}}_r^{(w)} \tilde{\mathbf{f}}_q^{(w)T}) = \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)},$$

the update for d_{qr} can also be expressed as

$$d_{qr}^{(w+1)} = \frac{1}{m} \left\{ \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)} + \text{tr}(\mathbf{T}_{ff_{qr}}^{(w)}) \right\}.$$

This completes the M-step.

6.5 Example: National Variety Trial data

The purpose of this example is to investigate algorithm performance. We do not provide an interpretation of the results from the models fitted. The interpretation of factor analytic mixed models applied to plant breeding data can be found in Smith et al. (2001). In the following when we refer to Model $\mathcal{O}$ or Model $\mathcal{M}$ it should be read as “the REML EM algorithm used to estimate the variance parameters associated with Model A ” where Model A is Model $\mathcal{O}$ or Model $\mathcal{M}$ as the case may be.

Model $\mathcal{O}$ and Model $\mathcal{M}$ with $k = 1$ factors were applied to the small subset of National Variety Trial grain yield data presented in Table 6.2. The subset of data we consider comprises 588 observations made up of 42 varieties grown in 7 trials with 2 replicates of every variety in each trial. In this and later chapters we do not consider models with $k > 1$ factors for two reasons. Firstly, we will not have to consider the effect of placing arbitrary constraints on factor loadings on algorithm performance. This would have been a concern when we consider a REML PX EM algorithm in the next chapter where the parameter expansion is associated with the

factor loadings. Secondly, fitting factor analytic mixed models where $k > 1$ involves a model building process. The most important stage of this process is fitting a factor analytic mixed model with $k = 1$ factors as the parameter estimates obtained from this model can be used to inform the starting values for subsequent models. We do not believe that restricting ourselves to factor analytic mixed models with $k = 1$ factors diminishes the results and conclusions presented in this and the next two chapters.

For both models we define the starting values $\boldsymbol{\lambda}^{(0)} = (0.1, 0.1, 0.1, 0.1, 0.1, 0.1, 0.1)^T$, $\psi^{(0)} = 0.1$, and $\sigma^{2(0)} = 0.1$. For Model $\mathcal{M}$ we define $d^{(0)} = 1$. The criterion presented in (4.18), where $\boldsymbol{\kappa}$ is replaced by $\boldsymbol{\Gamma}$, was used to determine if the sequence $\boldsymbol{\Gamma}^{(0)}, \boldsymbol{\Gamma}^{(1)}, \dots$ converged to a (local) maximum of the residual log-likelihood function. We assess whether or not the stopping rule with $\varepsilon = 10^{-8}$ is too stringent by considering different levels of ε . In Table 6.1 we present the different levels of ε , iterations until convergence, the final residual log-likelihood (ℓ_R), and the Frobenius matrix norm of the estimated genetic variance-covariance matrix when using Model $\mathcal{O}$ and Model $\mathcal{M}$ with $k = 1$ factors.

Model $\mathcal{M}$ converges in fewer iterations than Model $\mathcal{O}$ for all levels of ε . If we consider an accurate estimate of the genetic variance-covariance matrix to be one that has the same Frobenius matrix norm to 5 decimal places as that calculated when $\varepsilon = 10^{-8}$ then, based on the information presented in Table 6.1, it is possible to consider a less stringent convergence criterion for both models. For Model $\mathcal{M}$ there is no cost in terms of the final residual log-likelihood and very little cost in terms of an accurate estimate of genetic variance-covariance matrix by increasing ε to 10^{-5} . Likewise, for Model $\mathcal{O}$ there is little cost in increasing ε to 10^{-6} . At these levels of ε Model $\mathcal{O}$ and Model $\mathcal{M}$ converged in 76 and 38 iterations respectively. Further reductions in the number of iterations until convergence are possible by considering parameter expanded versions of Model $\mathcal{O}$ and Model $\mathcal{M}$. These are considered in the next chapter.

Table 6.1: Iterations until convergence, final residual log-likelihood (ℓ_R), and the Frobenius matrix norm of the estimated genetic variance-covariance matrix when using a REML EM algorithm to fit Model $\mathcal{O}$ and Model $\mathcal{M}$ with $k = 1$ factors and using a stopping rule with different levels of ε . For Model $\mathcal{O}$ and Model $\mathcal{M}$ with $k = 1$ factors there are 9 and 10 variance parameters to be estimated respectively.

ε	Model $\mathcal{O}$			Model $\mathcal{M}$		
	Iter.	ℓ_R	$\ \widehat{\mathbf{G}}_e\ _F$	Iter.	ℓ_R	$\ \widehat{\mathbf{G}}_e\ _F$
10^{-8}	104	300.9849	0.3180123	66	300.9849	0.3180123
10^{-7}	90	300.9849	0.3180121	56	300.9849	0.3180123
10^{-6}	76	300.9849	0.3180101	47	300.9849	0.3180125
10^{-5}	63	300.9849	0.3179936	38	300.9849	0.3180146
10^{-4}	49	300.9849	0.3178231	29	300.9848	0.3180426
10^{-3}	35	300.9842	0.3161148	21	300.9778	0.3182776
10^{-2}	21	300.8976	0.3002162	15	300.8208	0.3185271

CHAPTER 6. TWO MODEL SPECIFICATIONS FOR A FACTOR ANALYTIC MIXED MODEL

Table 6.2: National Variety Trial grain yield data (tonnes per hectare) for 7 trials comprising 2 replicates, labelled “I” and “II”, of 42 varieties.

	Trial 1		Trial 2		Trial 3		Trial 4		Trial 5		Trial 6		Trial 7	
	I	II	I	II	I	II	I	II	I	II	I	II	I	II
V1	5.14	5.27	2.63	0.90	3.18	3.24	2.08	2.18	4.83	5.31	2.85	3.41	3.51	3.26
V2	5.60	5.28	1.91	1.89	3.03	2.66	2.23	2.00	4.13	4.32	2.75	2.52	2.80	2.47
V3	5.22	4.69	2.46	1.90	2.86	2.76	2.09	1.86	4.51	4.40	2.58	2.70	2.80	2.85
V4	4.73	5.23	2.62	1.33	3.08	2.60	1.93	1.73	4.62	4.65	2.69	3.06	2.64	2.13
V5	4.89	5.51	3.24	2.18	2.97	2.87	2.49	2.20	4.67	4.65	2.78	3.18	3.35	3.09
V6	5.21	4.82	2.29	1.99	3.19	3.13	1.81	1.90	4.26	4.26	2.51	2.84	3.38	3.53
V7	5.36	4.93	2.06	1.98	3.18	3.04	2.26	2.34	4.54	5.09	2.62	3.09	3.10	2.83
V8	5.52	5.46	2.32	2.49	3.15	3.11	2.09	2.49	5.01	5.08	2.79	3.28	3.16	3.24
V9	4.93	4.97	2.62	2.82	2.86	3.01	2.46	2.11	4.51	4.67	2.72	2.68	2.81	2.97
V10	5.51	5.22	1.89	2.18	3.03	2.97	1.99	2.04	5.04	4.54	2.35	2.49	2.67	2.60
V11	6.03	5.51	2.89	2.63	2.92	3.27	2.14	2.34	4.52	5.06	3.26	3.34	3.45	3.23
V12	3.83	4.46	0.79	0.70	1.79	1.74	1.86	2.08	2.97	3.48	1.31	1.62	2.09	2.21
V13	5.63	5.07	1.54	1.45	2.51	2.31	1.80	1.81	4.35	4.62	2.84	3.03	2.57	2.76
V14	6.21	5.79	3.37	1.20	2.86	2.86	2.03	2.46	5.11	5.20	2.99	3.19	3.43	2.97
V15	6.09	6.11	2.66	1.65	2.64	2.80	2.16	2.33	4.86	4.95	3.14	3.23	3.29	3.39
V16	6.36	6.08	2.26	2.24	2.73	2.50	2.27	2.52	5.21	5.01	2.99	2.83	2.86	2.88
V17	5.84	5.21	2.13	2.35	2.91	3.20	2.38	2.29	4.60	4.87	3.18	2.90	3.39	3.14
V18	5.56	5.52	2.69	2.05	2.95	2.97	2.06	2.12	4.95	4.91	2.97	2.77	3.31	3.06
V19	5.32	5.50	2.68	2.16	3.33	3.17	2.24	2.08	5.21	5.11	3.01	3.41	3.43	3.25
V20	5.45	5.72	2.51	2.73	3.20	2.89	2.20	2.32	5.05	5.43	2.91	3.13	3.48	3.55
V21	6.03	5.87	2.27	1.89	2.75	2.52	3.05	2.28	4.71	5.01	2.99	3.30	3.39	2.79
V22	5.65	5.55	2.01	1.77	2.85	2.94	2.05	2.26	4.20	4.55	2.73	3.19	3.48	2.82
V23	5.29	5.28	0.69	1.71	2.39	2.55	2.01	2.25	3.38	3.23	2.26	2.26	2.89	2.55
V24	5.92	5.82	1.89	2.07	3.12	2.92	2.14	2.27	4.36	4.64	2.74	3.06	2.96	3.03
V25	5.36	5.74	2.59	2.11	3.08	2.68	2.16	2.61	4.01	4.56	2.76	2.57	3.21	2.98
V26	5.24	5.35	2.32	2.12	2.49	2.73	2.07	2.09	4.13	4.00	2.58	2.92	2.63	2.39
V27	5.27	5.24	1.56	1.99	2.60	2.57	2.16	2.47	5.20	4.88	2.58	3.31	3.40	2.79
V28	5.19	5.34	3.14	1.88	3.23	3.11	2.32	2.17	4.97	4.88	2.99	2.81	2.37	2.37
V29	5.19	4.99	1.25	1.56	2.93	2.94	2.13	1.87	4.36	4.11	2.78	2.39	3.41	3.13
V30	5.33	5.08	2.03	2.04	3.30	2.93	2.36	2.36	5.34	5.08	3.14	3.46	3.17	3.20
V31	5.76	5.58	2.01	1.83	2.98	2.92	2.08	2.22	4.85	4.71	3.03	3.40	2.68	2.86
V32	5.67	5.51	1.90	2.38	2.76	3.45	2.30	2.28	4.89	5.22	2.99	3.07	3.47	3.09
V33	5.17	5.00	2.25	2.15	2.52	2.91	2.54	2.18	4.16	4.39	2.58	2.76	2.81	2.68
V34	5.20	5.06	2.14	2.09	2.61	3.10	1.99	2.39	4.69	5.01	2.42	2.67	2.90	2.87
V35	5.44	5.01	2.39	2.38	3.14	2.99	2.05	2.09	4.45	4.54	2.51	3.40	3.27	3.15
V36	5.68	5.74	2.17	2.49	3.21	3.29	2.17	2.05	4.37	5.06	2.85	3.06	3.28	3.46
V37	6.18	5.86	1.24	1.36	2.30	2.14	2.26	2.17	4.45	4.48	2.86	2.36	2.50	2.51
V38	5.23	5.37	2.04	2.36	3.09	2.69	2.12	2.45	4.30	4.86	3.29	3.16	2.85	2.45
V39	5.76	5.46	2.16	1.26	2.98	2.87	2.08	2.51	5.05	5.56	3.16	2.96	3.44	3.40
V40	5.46	5.17	1.71	1.41	2.68	2.49	2.03	2.43	4.41	4.85	3.04	2.82	3.15	2.97
V41	5.25	4.84	3.02	2.12	2.95	3.17	2.28	2.32	4.46	4.68	2.59	2.42	3.62	3.07
V42	5.24	5.16	2.85	2.17	3.09	2.95	1.95	2.00	4.44	3.89	3.01	3.04	2.21	2.71

Chapter 7

REML PX EM algorithms for two factor analytic mixed model specifications

In this chapter we consider REML PX EM algorithms based on the two factor analytic mixed model specifications considered in the previous chapter, which were referred to as Model $\mathcal{O}$ and Model $\mathcal{M}$. The difference between the two being that for Model $\mathcal{O}$ the variance of the factor scores is the identity matrix whereas for Model $\mathcal{M}$ it comprises an unstructured matrix.

The parameter expansion of Model $\mathcal{O}$ will be referred to as Model $\mathcal{X}$, this is the same notation used by Liu et al. (1998). Model $\mathcal{M}$ and Model $\mathcal{X}$ are very similar and it is of interest to compare the two. However, the main purpose of this chapter is to present a parameter expansion of Model $\mathcal{M}$, which will be referred to as Model $\mathcal{Z}$. The specification of Model $\mathcal{Z}$ is new. In this model the parameter expansion is associated with the specific variances rather than the factor loadings as in Model $\mathcal{X}$. Using the example data considered in the previous chapter it is shown that this model converges in less iterations than either Model $\mathcal{X}$ or Model $\mathcal{M}$. We consider Model $\mathcal{X}$ and Model $\mathcal{Z}$ in turn.

7.1 REML PX EM algorithm - Model $\mathcal{X}$

Using the parameter expansion considered by Liu et al. (1998) we define the parameter expanded model

$$\text{Model } \mathcal{X} \quad \mathbf{y}_o = \mathbf{X}\boldsymbol{\beta}_* + \mathbf{Z}_g\mathbf{u}_g + \mathbf{e},$$

where $\mathbf{u}_g = (\boldsymbol{\Lambda}_* \otimes \mathbf{I}_m)\mathbf{f} + \boldsymbol{\delta}$. We assume $\boldsymbol{\delta} \sim N(\mathbf{0}, \boldsymbol{\Psi}_* \otimes \mathbf{I}_m)$ where $\boldsymbol{\Psi}_* = \psi_*\mathbf{I}_p$ and $\mathbf{f} \sim N(\mathbf{0}, \mathbf{D} \otimes \mathbf{I}_m)$. We assume that the $n \times 1$ vector of residuals is distributed $\mathbf{e} \sim N(\mathbf{0}, \mathbf{R}_*)$ where $\mathbf{R}_* = \sigma_*^2\mathbf{I}_n$. The difference between Model $\mathcal{O}$ and Model $\mathcal{X}$ is the assumed distribution of the $mk \times 1$ vector $\mathbf{f}$. In Model $\mathcal{X}$ it is assumed that $\mathbf{f} \sim N(\mathbf{0}, \mathbf{D} \otimes \mathbf{I}_m)$ whereas in Model $\mathcal{O}$ it is assumed that $\mathbf{f} \sim N(\mathbf{0}, \mathbf{I}_{mk})$. We define $\mathbf{D} = \mathbf{D}(\mathbf{d})$ as a $k \times k$ symmetric positive definite matrix where $\mathbf{d}$ is a vector of variance parameters.

We follow the convention established by Liu et al. (1998) and define the expanded parameter set associated with Model $\mathcal{X}$ as $\boldsymbol{\Gamma}_* = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \psi_*, \mathbf{d}^T)^T$ where the subscript $*$ is used to distinguish parameters associated with Model $\mathcal{O}$ and Model $\mathcal{X}$. For the parameter expanded model the distribution of the observed data is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}_*, \mathbf{H}_*),$$

where $\mathbf{H}_* = (\boldsymbol{\Lambda}_*\mathbf{D}\boldsymbol{\Lambda}_*^T + \boldsymbol{\Psi}_*) \otimes \mathbf{I}_m$. Using the parameter expanded model we identify the parameters associated with Model $\mathcal{O}$ as $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \psi)^T = \mathcal{R}(\boldsymbol{\Gamma}_*) = [\sigma_*^2, \text{vec}\{\boldsymbol{\Lambda}_*\text{Chol}(\mathbf{D})\}, \psi_*]^T$ where Chol denotes the Cholesky decomposition and $\mathcal{R}$ is the reduction function from the expanded parameter set to the original parameter set.

7.1.1 Complete data density function

For Model $\mathcal{X}$ the complete data specification remains the same as that for Model $\mathcal{O}$, i.e., $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$ where the missing data vector is define as $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$. However, the complete data log-density function is different in that it contains the so-called auxillary parameter $\mathbf{d}$. We use the subscript $\mathcal{X}$ to distinguish between the complete data log-density function and $Q(\cdot)$ function for Model $\mathcal{O}$ and Model $\mathcal{X}$. The complete data log-density function for the latter is

$$\log\{f_{\mathcal{X}}(\mathbf{y}_c; \mathbf{\Gamma}_*)\} = \log\{f_{\mathcal{X}}(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \sigma_*^2, \boldsymbol{\lambda}_*)\} + \log\{f_{\mathcal{X}}(\mathbf{f}; \mathbf{d})\} + \log\{f_{\mathcal{X}}(\boldsymbol{\delta}; \psi_*)\}, \quad (7.1)$$

where

$$\begin{aligned} \log\{f_{\mathcal{X}}(\mathbf{f}; \mathbf{d})\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} \right], \\ \log\{f_{\mathcal{X}}(\boldsymbol{\delta}; \psi_*)\} &= -\frac{1}{2} \left[\log\{\det(\boldsymbol{\Psi}_* \otimes \mathbf{I}_m)\} + \boldsymbol{\delta}^T (\boldsymbol{\Psi}_*^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right] \end{aligned}$$

and

$$\begin{aligned} &\log\{f_{\mathcal{X}}(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \sigma_*^2, \boldsymbol{\lambda}_*)\} \\ &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R}_* \mathbf{L}_2)\} \right. \\ &\quad \left. + \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} - \mathbf{Z}_g \boldsymbol{\delta}\}^T \mathbf{S}_* \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} - \mathbf{Z}_g \boldsymbol{\delta}\} \right], \end{aligned}$$

where $\mathbf{S}_* = \mathbf{R}_*^{-1} - \mathbf{R}_*^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{R}_*^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}_*^{-1}$. This is the same as the complete data density function for Model $\mathcal{M}$ where the variance parameters σ^2 , $\boldsymbol{\lambda}$, and ψ have been substituted with σ_*^2 , $\boldsymbol{\lambda}_*$, and ψ_* .

7.2 REML PX EM algorithm E-step for Model $\mathcal{X}$

For Model $\mathcal{X}$ we define $\mathbf{\Gamma}_*^{(w)} = (\sigma_*^2, \boldsymbol{\lambda}_*^{(w)T}, \psi_*^{(w)}, \mathbf{d}_0^T)^T$ where $\mathbf{D}(\mathbf{d}_0) = \mathbf{I}_k$. The E-step of the REML PX EM algorithm involves calculating $Q_{\mathcal{X}}(\mathbf{\Gamma}_*; \mathbf{\Gamma}_*^{(w)})$ which is equivalent to the E-step of the REML PX EM algorithm for Model $\mathcal{O}$. Note that for Model $\mathcal{O}$ and Model $\mathcal{X}$

$$\begin{aligned} \tilde{\mathbf{f}}^{(w)} &= (\boldsymbol{\Lambda}^{(w)T} \otimes \mathbf{I}_m) \mathbf{P}^{(w)} \mathbf{y}_o, \\ \mathbf{T}_{ff}^{(w)} &= \mathbf{I}_{mk} - (\boldsymbol{\Lambda}^{(w)T} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P}^{(w)} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m), \\ \mathbf{T}_{f\delta}^{(w)} &= -(\boldsymbol{\Psi}^{(w)} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P}^{(w)} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m). \end{aligned}$$

7.3 REML PX EM algorithm M-step for Model $\mathcal{X}$

The M-step of the REML PX-EM algorithm involves maximising $Q_{\mathcal{X}}(\Gamma_*; \Gamma_*^{(w)})$ with respect to $\Gamma_* = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \psi_*, \mathbf{d}^T)^T$ and then applying the reduction function $\mathcal{R}(\Gamma_*) = [\sigma_*^2, \text{vec}\{\boldsymbol{\Lambda}_* \text{Chol}(\mathbf{D})\}, \psi_*]^T$ to obtain estimates of $\Gamma^{(w+1)} = (\sigma^{2(w+1)}, \boldsymbol{\lambda}^{(w+1)T}, \psi^{(w+1)})^T$.

7.3.1 Update for σ_*^2

The update for σ_*^2 is the same as the update for σ^2 under Model $\mathcal{O}$, i.e.,

$$\begin{aligned} \sigma_*^{2(w+1)} &= \frac{1}{n-p} \left[\{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \}^T \mathbf{K} \{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \} \right. \\ &\quad + \text{tr} \{ (\boldsymbol{\Lambda}^{(w)})^T \otimes \mathbf{I}_m \} \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \} + \text{tr} \{ \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)} \} \\ &\quad \left. + 2 \text{tr} \{ \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)}) \} \right]. \end{aligned}$$

Applying the reduction function we have $\sigma^{2(w+1)} = \sigma_*^{2(w+1)}$.

7.3.2 Update for λ_{*jr}

The update for λ_{*jr} is the same as that presented for λ_{jr} under Model $\mathcal{O}$ except we substitute $\boldsymbol{\Lambda}_*$ and $\mathbf{S}_*$ for $\boldsymbol{\Lambda}$ and $\mathbf{S}$ respectively, i.e.,

$$\boldsymbol{\lambda}_*^{(w+1)} = \mathbf{A}_*^{(w)-1} \mathbf{b}_*^{(w)},$$

where the elements of $\mathbf{A}_*^{(w)}$ are

$$\begin{aligned} a_{*jr,lq}^{(w)} &= E \{ \mathbf{f}^T (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}_* \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma_*^{(w)} \} \\ &= \tilde{\mathbf{f}}^{(w)T} (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} \\ &\quad + \text{tr} \{ (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \} \end{aligned}$$

and the elements of $\mathbf{b}_*^{(w)}$ are

$$\begin{aligned}
 b_{*jr}^{(w)} &= E\{\mathbf{y}_o^T \mathbf{S}_* \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \Gamma_*^{(w)}\} \\
 &\quad - \text{tr}\{E(\mathbf{Z}_g^T \mathbf{S}_*^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \Gamma_*^{(w)})\} \\
 &= \mathbf{y}_o^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \text{tr}\{\mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g(\mathbf{J}_{jr} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)})\}.
 \end{aligned}$$

Applying the reduction function we have $\boldsymbol{\lambda}^{(w+1)} = \text{vec}\{\boldsymbol{\Lambda}_*^{(w+1)} \text{Chol}(\mathbf{D}^{(w+1)})\}$.

7.3.3 Update for ψ_*

The update for ψ_* is the same as the update for ψ under Model $\mathcal{O}$ where ψ_* replaces ψ , i.e.,

$$\psi_*^{(w+1)} = \frac{1}{mp} \left\{ \tilde{\boldsymbol{\delta}}^{(w)T} \tilde{\boldsymbol{\delta}}^{(w)} + \text{tr}(\mathbf{T}_{\delta\delta}^{(w)}) \right\}$$

Applying the reduction function we have $\psi_*^{(w+1)} = \psi_*^{(w+1)}$.

7.3.4 Update for d_{qr}

The update for d_{qr} is

$$d_{qr}^{(w+1)} = \frac{1}{m} \left\{ \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)} + \text{tr}(\mathbf{T}_{ff_{qr}}^{(w)}) \right\},$$

which is similar to the update for d_{qr} under Model $\mathcal{M}$ presented in the previous chapter. The only difference being the form of $\tilde{\mathbf{f}}^{(w)}$ and $\mathbf{T}_{ff}^{(w)}$. Under Model $\mathcal{M}$ these contain the unstructured matrix $\mathbf{D}$ whereas for the model we are considering, Model $\mathcal{X}$, it is assumed that $\mathbf{D} = \mathbf{I}_k$.

7.4 REML PX EM algorithm - Model $\mathcal{Z}$

We now consider a new specification of a REML PX EM algorithm for a factor analytic mixed model. Consider the parameter expanded model

$$\text{Model } \mathcal{Z} \quad \mathbf{y}_o = \mathbf{X}\boldsymbol{\beta}_* + \mathbf{Z}_g \mathbf{u}_g + \mathbf{e}$$

where $\mathbf{e} \sim N(\mathbf{0}, \mathbf{R}_*)$, $\mathbf{R}_* = \sigma_*^2 \mathbf{I}_n$ and $\mathbf{u}_g = (\mathbf{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} + (\mathbf{\Omega} \otimes \mathbf{I}_m) \boldsymbol{\delta}$, i.e., we introduce the $p \times p$ positive definite diagonal matrix $\mathbf{\Omega} = \mathbf{\Omega}(\boldsymbol{\omega})$. When we assume a common specific variance for $\boldsymbol{\delta}$, i.e., $\boldsymbol{\delta} \sim N(\mathbf{0}, \mathbf{\Psi}_* \otimes \mathbf{I}_m)$ where $\mathbf{\Psi}_* = \psi_* \mathbf{I}_{mp}$ then $\mathbf{\Omega}$ has the form $\mathbf{\Omega} = \omega \mathbf{I}_p$. We assume that $\boldsymbol{\delta}$ and the $mk \times 1$ vector $\mathbf{f} \sim N(\mathbf{0}, \mathbf{D}_* \otimes \mathbf{I}_m)$ are independent of each other. We define $\mathbf{D}_* = \mathbf{D}_*(\mathbf{d}_*)$ as a $k \times k$ symmetric positive definite matrix where $\mathbf{d}_*$ is a vector of variance parameters.

We define the expanded parameter set associated with Model $\mathcal{Z}$ as

$\boldsymbol{\tau} = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \psi_*, \mathbf{d}_*^T, \boldsymbol{\omega}^T)^T$ where the subscript $*$ is used to distinguish parameters associated with Model $\mathcal{M}$ and Model $\mathcal{Z}$. For the parameter expanded model the distribution of the observed data is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}_*, \mathbf{H}_*),$$

where $\mathbf{H}_* = \mathbf{Z}_g \{ (\mathbf{\Lambda}_* \mathbf{D}_* \mathbf{\Lambda}_*^T + \mathbf{\Omega} \mathbf{\Psi}_* \mathbf{\Omega}) \otimes \mathbf{I}_m \} \mathbf{Z}_g^T + \mathbf{R}_*$. Using the parameter expanded model we identify the parameters associated with Model $\mathcal{M}$ as $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \psi, \mathbf{d}^T)^T = \mathcal{R}(\boldsymbol{\tau}) = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \omega^2 \psi_*, \mathbf{d}_*^T)^T$ where $\mathcal{R}$ is the reduction function from the expanded parameter set to the original parameter set.

The difference between the parameter expanded Model $\mathcal{X}$ and Model $\mathcal{Z}$ is that in the former the auxillary parameter $\mathbf{d}$ is associated with the factor loadings in the reduction function, whereas in the later the auxillary parameter $\boldsymbol{\omega}$ is associated with the specific variances.

An important implication of the non-uniqueness of $\mathbf{\Lambda}_*$ and $\mathbf{D}_*$ for Model $\mathcal{Z}$ is that the REML PX EM algorithm applied to this model is not guaranteed to have a rate of convergence at least as fast as the REML EM algorithm applied to Model $\mathcal{M}$. We don't believe this is an impediment to the implementation of Model $\mathcal{Z}$ for the analysis of plant breeding data for two reasons. Firstly, $\hat{\mathbf{G}}_e$ and $\tilde{\mathbf{u}}_g$ are unique; and secondly, the cases in which a REML EM algorithm applied to Model $\mathcal{M}$ would converge faster than a REML PX EM algorithm applied to Model $\mathcal{Z}$ are atypical.

7.4.1 Complete data density function

For Model $\mathcal{Z}$ the complete data specification is the same as that for Model $\mathcal{M}$, i.e., $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$. We use the subscript $\mathcal{Z}$ to distinguish between the complete data log-density function and $Q(\cdot)$ function for Model $\mathcal{M}$ and Model $\mathcal{Z}$. To form the complete data density function for Model $\mathcal{Z}$ requires the joint distribution

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{f} \\ \boldsymbol{\delta} \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{L}_2^T \mathbf{H}_* \mathbf{L}_2 & \mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Lambda}_* \mathbf{D}_* \otimes \mathbf{I}_m) & \mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Omega} \boldsymbol{\Psi}_* \otimes \mathbf{I}_m) \\ (\mathbf{D}_* \boldsymbol{\Lambda}_*^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{L}_2 & \mathbf{D}_* \otimes \mathbf{I}_m & \mathbf{0} \\ (\boldsymbol{\Psi}_* \boldsymbol{\Omega} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{L}_2 & \mathbf{0} & \boldsymbol{\Psi}_* \otimes \mathbf{I}_m \end{pmatrix} \right\}.$$

The complete data log-density function is

$$\log\{f_{\mathcal{Z}}(\mathbf{y}_c; \boldsymbol{\mathcal{T}})\} = \log\{f_{\mathcal{Z}}(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \sigma_*^2, \boldsymbol{\lambda}_*, \omega)\} + \log\{f_{\mathcal{Z}}(\mathbf{f}; \mathbf{d}_*)\} + \log\{f_{\mathcal{Z}}(\boldsymbol{\delta}; \boldsymbol{\psi}_*)\},$$

where

$$\begin{aligned} \log\{f_{\mathcal{Z}}(\mathbf{f}; \mathbf{d}_*)\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D}_* \otimes \mathbf{I}_m)\} + \mathbf{f}^T (\mathbf{D}_*^{-1} \otimes \mathbf{I}_m) \mathbf{f} \right], \\ \log\{f_{\mathcal{Z}}(\boldsymbol{\delta}; \boldsymbol{\psi}_*, \omega)\} &= -\frac{1}{2} \left[\log\{\det(\boldsymbol{\Psi}_* \otimes \mathbf{I}_m)\} + \boldsymbol{\delta}^T (\boldsymbol{\Psi}_*^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right] \end{aligned}$$

and

$$\begin{aligned} \log\{f_{\mathcal{Z}}(\mathbf{y}_2 | \mathbf{f}, \boldsymbol{\delta}; \boldsymbol{\lambda}_*, \boldsymbol{\psi}_*, \omega)\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R}_* \mathbf{L}_2)\} \right. \\ &\quad \left. + \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} - \mathbf{Z}_g(\boldsymbol{\Omega} \otimes \mathbf{I}_m) \boldsymbol{\delta}\}^T \mathbf{S}_* \{\mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} - \mathbf{Z}_g(\boldsymbol{\Omega} \otimes \mathbf{I}_m) \boldsymbol{\delta}\} \right]. \end{aligned}$$

7.5 REML PX EM algorithm E-step for Model $\mathcal{Z}$

For Model $\mathcal{Z}$ we define $\boldsymbol{\mathcal{T}}^{(w)} = (\sigma_*^2, \boldsymbol{\lambda}_*^{(w)T}, \boldsymbol{\psi}_*^{(w)}, \mathbf{d}_*^{(w)T}, \omega_0)^T$ where $\boldsymbol{\Omega}(\omega_0) = \mathbf{I}_p$. The E-step of the REML PX EM algorithm involves calculating $Q_{\mathcal{Z}}(\boldsymbol{\mathcal{T}}; \boldsymbol{\mathcal{T}}^{(w)})$ which is equivalent to $Q(\boldsymbol{\Gamma}; \boldsymbol{\Gamma}^{(w)})$, i.e., the E-step of the REML PX EM algorithm for Model $\mathcal{Z}$ and the E-step of the REML EM algorithm for Model $\mathcal{M}$ are the same.

7.6 REML PX EM algorithm M-step for Model $\mathcal{Z}$

The M-step of the REML PX EM algorithm involves maximising $Q_{\mathcal{Z}}(\boldsymbol{\mathcal{T}}; \boldsymbol{\mathcal{T}}^{(w)})$ with respect to $\boldsymbol{\mathcal{T}} = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \boldsymbol{\psi}_*, \mathbf{d}_*^T, \omega)^T$ and then applying the reduction function $\mathcal{R}(\boldsymbol{\mathcal{T}})$ to obtain $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \boldsymbol{\psi}, \mathbf{d}^T)^T$.

7.6.1 Update for σ_*^2

The update for σ_*^2 is the same as the update for σ^2 under Model $\mathcal{M}$, i.e.,

$$\begin{aligned} \sigma_*^{2(w+1)} &= \frac{1}{n-p} \left[\{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \}^T \mathbf{K} \{ \mathbf{y}_o - \mathbf{Z}_g(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \} \right. \\ &\quad + \text{tr} \left\{ (\boldsymbol{\Lambda}^{(w)})^T \otimes \mathbf{I}_m \right\} \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \Big\} \\ &\quad \left. + \text{tr} \left(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)} \right) + 2 \text{tr} \left\{ \mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g (\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)}) \right\} \right]. \end{aligned}$$

Applying the reduction function we have $\sigma^{2(w+1)} = \sigma_*^{2(w+1)}$.

7.6.2 Update for λ_{*jr}

The derivation of the update for λ_{*jr} is the same as that for λ_{jr} under Model $\mathcal{M}$, i.e.,

$$\boldsymbol{\lambda}_*^{(w+1)} = \mathbf{A}_*^{(w)-1} \mathbf{b}_*^{(w)},$$

where the elements of the $mk \times mk$ matrix $\mathbf{A}_*^{(w)}$ are

$$\begin{aligned} a_{*jr,lq}^{(w)} &= E \{ \mathbf{f}^T (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}_* \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \mathcal{T}^{(w)} \} \\ &= \tilde{\mathbf{f}}^{(w)T} (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} \\ &\quad + \text{tr} \{ (\mathbf{J}_{lq}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \} \end{aligned}$$

and $\mathbf{b}_*^{(w)}$ is a $mk \times 1$ vector with elements

$$\begin{aligned} b_{*jr}^{(w)} &= E \{ \mathbf{y}_o^T \mathbf{S}_* \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \mathcal{T}^{(w)} \} - \text{tr} \{ E (\mathbf{Z}_g^T \mathbf{S}_*^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} \boldsymbol{\delta}^T | \mathbf{y}_2; \mathcal{T}^{(w)}) \} \\ &= \mathbf{y}_o^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} - \text{tr} \{ \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\mathbf{J}_{jr} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)}) \}. \end{aligned}$$

Applying the reduction function we have $\boldsymbol{\lambda}^{(w+1)} = \boldsymbol{\lambda}_*^{(w+1)}$.

7.6.3 Update for ψ_*

The update for ψ_* is the same as that presented for ψ under Model $\mathcal{M}$, i.e.,

$$\psi_*^{(w+1)} = \frac{1}{mp} \left\{ \tilde{\boldsymbol{\delta}}^{(w)T} \tilde{\boldsymbol{\delta}}^{(w)} + \text{tr}(\mathbf{T}_{\delta\delta}^{(w)}) \right\}.$$

Applying the reduction function we have $\psi_*^{(w+1)} = \omega^{2(w+1)} \psi_*^{(w+1)}$.

7.6.4 Update for d_{*qr}

The update for d_{*qr} is the same as that for d_{qr} under Model $\mathcal{M}$, i.e.,

$$d_{*qr}^{(w+1)} = \frac{1}{m} \left\{ \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)} + \text{tr}(\mathbf{T}_{ffqr}^{(w)}) \right\}.$$

Applying the reduction function we have $d_{qr}^{(w+1)} = d_{*qr}^{(w+1)}$.

7.6.5 Update for ω

Differentiating $Q_{\mathcal{Z}}(\boldsymbol{\tau}; \boldsymbol{\tau}^{(w)})$ with respect to ω we have

$$\begin{aligned} \frac{\partial Q_{\mathcal{Z}}(\boldsymbol{\tau}; \boldsymbol{\tau}^{(w)})}{\partial \omega} = & -\frac{1}{2} \left\{ -2\mathbf{y}_o^T \mathbf{S} \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} \right. \\ & + 2\text{tr} \left\{ (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)}) \right\} \\ & \left. + 2\omega \tilde{\boldsymbol{\delta}}^{(w)T} \mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} + 2\omega \text{tr}(\mathbf{Z}_g^T \mathbf{S} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)}) \right\} \end{aligned}$$

Equating to zero and solving for ω we have

$$\omega^{(w+1)} = \frac{\mathbf{y}_o^T \mathbf{S}^{(w)} \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} - \text{tr} \left\{ (\boldsymbol{\Lambda}^{(w)T} \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g (\tilde{\mathbf{f}}^{(w)} \tilde{\boldsymbol{\delta}}^{(w)T} + \mathbf{T}_{f\delta}^{(w)}) \right\}}{\tilde{\boldsymbol{\delta}}^{(w)T} \mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g \tilde{\boldsymbol{\delta}}^{(w)} + \text{tr}(\mathbf{Z}_g^T \mathbf{S}^{(w)} \mathbf{Z}_g \mathbf{T}_{\delta\delta}^{(w)})}$$

7.7 Example: National Variety Trial data

In the following when we refer to Model $\mathcal{O}$ or Model $\mathcal{M}$ it should be read as “the REML EM algorithm used to estimate the variance parameters associated with

Model A ” where Model A is Model $\mathcal{O}$ or Model $\mathcal{M}$ as the case may be. When we refer to Model $\mathcal{X}$ or Model $\mathcal{Z}$ it should be read as “the REML PX EM algorithm used to estimate the variance parameters associated with Model B ” where Model B is Model $\mathcal{X}$ or Model $\mathcal{Z}$ as the case may be. We consider the small subset of National Variety Trial data used in the previous two chapters and presented in Table 6.2.

For Model $\mathcal{X}$ and Model $\mathcal{Z}$ with $k = 1$ factors we define the starting values $\boldsymbol{\lambda}^{(0)} = (0.1, 0.1, 0.1, 0.1, 0.1, 0.1, 0.1)^T$, $\psi^{(0)} = 0.1$, $\sigma^{2(0)} = 0.1$, and $d^{(0)} = 1$. For Model $\mathcal{Z}$ we define $\omega^{(0)} = 1$. In Table 7.1 iterations until convergence and the final residual log-likelihood are considered for different levels of ε for both models.

If we compare the Table 7.1 and Table 6.1 we see that Model $\mathcal{X}$ converges in significantly fewer iterations than Model $\mathcal{O}$ and Model $\mathcal{Z}$ converges in 1 or 2 iterations less than Model $\mathcal{M}$. Of particular interest is the comparison between Model $\mathcal{X}$ and Model $\mathcal{M}$ as they differ in only one major respect, i.e., the later incorporates the variance parameter d in the E-step. For this particular set of data Model $\mathcal{M}$ converges in significantly fewer iterations than Model $\mathcal{X}$ which suggests that having the unstructured positive definite matrix $\mathbf{D}$ as part of the E-step can improve algorithm performance in terms of the number of iterations until convergence.

In regard to a less stringent convergence criteria, there is little or no cost in increasing ε to 10^{-5} for both models. Our preferred model and stopping rule is Model $\mathcal{Z}$ using $\varepsilon = 10^{-5}$. With this combination convergence is obtained in 36 iterations compared to the 45 iterations it takes Model $\mathcal{X}$ using $\varepsilon = 10^{-5}$ to converge. For smaller data sets such as the one considered here having an algorithm converge in 36 iterations might be acceptable. However, for larger data sets this number of iterations until convergence would be unsatisfactory in terms of elapsed time. In the next chapter we consider a new missing data specification for factor analytic mixed models with the aim of making further gains in REML PX EM algorithm performance.

Table 7.1: Iterations until convergence, final residual log-likelihood (ℓ_R), and the Frobenius matrix norm of the estimated genetic variance-covariance matrix when using a REML PX EM algorithm to fit Model $\mathcal{X}$ and Model $\mathcal{Z}$ with $k = 1$ factors and using a stopping rule with different levels of ε . For Model $\mathcal{X}$ and Model $\mathcal{Z}$ where $k = 1$ there are 9 and 10 variance parameters to be estimated respectively.

ε	Model $\mathcal{X}$			Model $\mathcal{Z}$		
	Iter.	ℓ_R	$\ \widehat{\mathbf{G}}_e\ _F$	Iter.	ℓ_R	$\ \widehat{\mathbf{G}}_e\ _F$
10^{-8}	74	300.9849	0.3180123	65	300.9849	0.3180123
10^{-7}	64	300.9849	0.3180123	55	300.9849	0.3180123
10^{-6}	54	300.9849	0.3180123	46	300.9849	0.3180123
10^{-5}	45	300.9849	0.3180126	36	300.9849	0.3180126
10^{-4}	35	300.9849	0.3180178	27	300.9847	0.3180158
10^{-3}	26	300.9844	0.3180826	19	300.9741	0.3179943
10^{-2}	17	300.9276	0.3186011	14	300.8342	0.3169530

Chapter 8

A new missing data specification for factor analytic mixed models

In Chapter 6 we considered two model specifications for a factor analytic mixed model and showed that one of these specifications, Model $\mathcal{M}$, could greatly reduce the number of iterations until convergence when using a REML EM algorithm. In Chapter 7 we considered a new REML PX EM algorithm for factor analytic mixed models where the parameter expansion is associated with the specific variances. This model was referred to as Model $\mathcal{Z}$. In both these chapters the distribution of the missing data vector was considered to be the joint distribution of the vector of factor scores and the vector of residuals associated with the factor analytic model.

In this chapter we consider an alternative missing data specification and hence complete data specification for factor analytic mixed models. The new missing data specification we consider is based on the joint distribution of the vector of factor scores and the vector of variety by trial effects. This new missing data specification is motivated by the “dependent formulation” of factor analytic variance models presented in Thompson et al. (2003).

This chapter proceeds by showing that the complete data log-density function associated with the new complete data specification for Model $\mathcal{M}$ is the same as the complete data density function for this model presented in Chapter 6. We consider Model $\mathcal{M}$ and Model $\mathcal{Z}$ using the new missing data specification and demonstrate, using the example considered in the previous two chapters and a simulation study, that this new missing data specification can further reduce the number of iterations until convergence.

8.1 REML EM algorithm - Model $\mathcal{M}$

We considered Model $\mathcal{M}$ in Chapter 6. We provide a succinct overview here for convenience. We define

$$\text{Model } \mathcal{M} \quad \mathbf{y}_o = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}_g\mathbf{u}_g + \mathbf{e},$$

where $\boldsymbol{\beta}$ is a $p \times 1$ vector of fixed effects representing trial means and $\mathbf{X}$ is its associated $n \times p$ design matrix of full column rank. It is assumed that the $n \times 1$ vector of residuals $\mathbf{e}$ is distributed $N(\mathbf{0}, \mathbf{R})$ where $\mathbf{R} = \sigma^2 \mathbf{I}_n$. The $n \times mp$ design matrix $\mathbf{Z}_g$ is associated with the $mp \times 1$ vector $\mathbf{u}_g$ of variety by trial effects. Factor analysis is used to model the variety effects within each trial, i.e.,

$$\mathbf{u}_g = (\boldsymbol{\Lambda} \otimes \mathbf{I}_m)\mathbf{f} + \boldsymbol{\delta},$$

where $\boldsymbol{\Lambda}$ is a $p \times k$ matrix of loadings and $\mathbf{f}$ is a $mk \times 1$ vector of varietal scores that can be partitioned as $\mathbf{f} = (\mathbf{f}_1^T, \dots, \mathbf{f}_k^T)^T$. It is assumed that $\mathbf{f} \sim N(\mathbf{0}, \mathbf{D} \otimes \mathbf{I}_m)$ where $\mathbf{D} = \mathbf{D}(\mathbf{d})$ is a $k \times k$ symmetric positive definite matrix and $\mathbf{d}$ is a vector of variance parameters. The vector $\boldsymbol{\delta}$ is a $mp \times 1$ vector of residuals that can be partitioned with respect to trials, i.e., $\boldsymbol{\delta} = (\boldsymbol{\delta}_1^T, \dots, \boldsymbol{\delta}_p^T)^T$ and is independent of $\mathbf{f}$. It is assumed that $\boldsymbol{\delta} \sim N(\mathbf{0}, \boldsymbol{\Psi} \otimes \mathbf{I}_m)$ where $\boldsymbol{\Psi} = \boldsymbol{\Psi}(\boldsymbol{\psi})$ is a $p \times p$ diagonal matrix. The vector of variance parameters $\boldsymbol{\psi}$ is commonly referred to as the specific variances. For simplicity we will assume that $\boldsymbol{\Psi} = \psi \mathbf{I}_p$, i.e., there is a common specific variance. It is assumed that the distribution of the vector of variety effects within trials is

$$\mathbf{u}_g \sim N(\mathbf{0}, \mathbf{G}_e \otimes \mathbf{I}_m),$$

where $\mathbf{G}_e = \boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T + \boldsymbol{\Psi}$ and is referred to as the genetic variance-covariance matrix. For Model $\mathcal{M}$ the distribution of the observed data is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}, \mathbf{H}),$$

where $\mathbf{H} = \mathbf{Z}_g(\mathbf{G}_e \otimes \mathbf{I}_m)\mathbf{Z}_g^T + \mathbf{R}$ and the vector of variance parameters to be estimated is $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \psi, \mathbf{d}^T)^T$.

8.1.1 Complete data density function for $\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$

We define the new missing data specification as

$$\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T,$$

where the superscript D is used to distinguish between this missing data specification and $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$. Likewise, the new complete data specification is defined as

$$\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T.$$

To form the complete data log-likelihood function for $\mathbf{y}_c^{(D)}$ requires the joint distribution

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{f} \\ \mathbf{u}_g \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{V}_{11} & \mathbf{V}_{12} & \mathbf{V}_{13} \\ \mathbf{V}_{21} & \mathbf{V}_{22} & \mathbf{V}_{23} \\ \mathbf{V}_{31} & \mathbf{V}_{32} & \mathbf{V}_{33} \end{pmatrix} \right\}, \quad (8.1)$$

where

$$\begin{aligned} \mathbf{V}_{11} &= \mathbf{L}_2^T \mathbf{H} \mathbf{L}_2, \\ \mathbf{V}_{22} &= \mathbf{D} \otimes \mathbf{I}_m, \\ \mathbf{V}_{33} &= (\boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T + \boldsymbol{\Psi}) \otimes \mathbf{I}_m, \\ \mathbf{V}_{12} &= \mathbf{V}_{21}^T = \mathbf{L}_2^T \mathbf{Z}_g (\boldsymbol{\Lambda} \mathbf{D} \otimes \mathbf{I}_m), \\ \mathbf{V}_{13} &= \mathbf{V}_{31}^T = \mathbf{L}_2^T \mathbf{Z}_g \{(\boldsymbol{\Lambda} \mathbf{D} \boldsymbol{\Lambda}^T + \boldsymbol{\Psi}) \otimes \mathbf{I}_m\}, \\ \mathbf{V}_{23} &= \mathbf{V}_{32}^T = \mathbf{D} \boldsymbol{\Lambda}^T \otimes \mathbf{I}_m. \end{aligned}$$

The complete data log-density function can be expressed as

$$\log\{f(\mathbf{y}_c^{(D)}; \boldsymbol{\Gamma})\} = \log\{f(\mathbf{y}_2 | \mathbf{y}_m^{(D)}; \boldsymbol{\Gamma})\} + \log\{f(\mathbf{y}_m^{(D)}; \boldsymbol{\Gamma})\}.$$

Consider $\log\{f(\mathbf{y}_2 | \mathbf{y}_m^{(D)}; \boldsymbol{\Gamma})\}$. Using (8.6) and Result A.1 it can be shown that

$$\mathbf{y}_2 | \mathbf{y}_m^{(D)} \sim N(\mathbf{L}_2^T \mathbf{Z}_g \mathbf{u}_g, \mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)$$

and the log-density function for $\mathbf{y}_2$ given $\mathbf{y}_m^{(D)}$ is (excluding constants)

$$\begin{aligned} \log\{f(\mathbf{y}_2|\mathbf{y}_m^{(D)}; \Gamma)\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} \right. \\ &\quad \left. + (\mathbf{y}_o - \mathbf{Z}_g \mathbf{u}_g)^T \mathbf{S} (\mathbf{y}_o - \mathbf{Z}_g \mathbf{u}_g) \right], \end{aligned} \quad (8.2)$$

where $\mathbf{S} = \mathbf{R}^{-1} - \mathbf{R}^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{R}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{R}^{-1}$. Now consider the log-density function $\log\{f(\mathbf{y}_m^{(D)}; \Gamma)\}$. We define

$$\mathcal{V} = \begin{Bmatrix} \mathbf{D} \otimes \mathbf{I}_m & \mathbf{D} \mathbf{\Lambda}^T \otimes \mathbf{I}_m \\ \mathbf{\Lambda} \mathbf{D} \otimes \mathbf{I}_m & (\mathbf{\Lambda} \mathbf{D} \mathbf{\Lambda}^T + \mathbf{\Psi}) \otimes \mathbf{I}_m \end{Bmatrix}.$$

The log-density function for $\mathbf{y}_m^{(D)}$ is

$$\log\{f(\mathbf{y}_m^{(D)}; \Gamma)\} = -\frac{1}{2} \left[\log\{\det(\mathcal{V})\} + \mathbf{y}_m^{(D)T} \mathcal{V}^{-1} \mathbf{y}_m^{(D)} \right]. \quad (8.3)$$

To express (8.3) in terms of $\mathbf{f}$ and $\mathbf{u}_g$ and their associated variance parameters we require $\det(\mathcal{V})$ and $\mathcal{V}^{-1}$. Using the fourth identity in Result A.3.3 we have

$$\begin{aligned} \det(\mathcal{V}) &= \det(\mathbf{D} \otimes \mathbf{I}_m) \det[\{(\mathbf{\Lambda} \mathbf{D} \mathbf{\Lambda}^T + \mathbf{\Psi}) \otimes \mathbf{I}_m\} \\ &\quad - \{(\mathbf{\Lambda} \mathbf{D} \otimes \mathbf{I}_m)(\mathbf{D}^{-1} \otimes \mathbf{I}_m)(\mathbf{D} \mathbf{\Lambda}^T \otimes \mathbf{I}_m)\}] \\ &= \det(\mathbf{D} \otimes \mathbf{I}_m) \det[\{(\mathbf{\Lambda} \mathbf{D} \mathbf{\Lambda}^T + \mathbf{\Psi}) \otimes \mathbf{I}_m\} - (\mathbf{\Lambda} \mathbf{D} \mathbf{\Lambda}^T \otimes \mathbf{I}_m)] \\ &= \det(\mathbf{D} \otimes \mathbf{I}_m) \det(\mathbf{\Psi} \otimes \mathbf{I}_m). \end{aligned}$$

Using Result A.3.6 we can express $\mathcal{V}^{-1}$ as

$$\mathcal{V}^{-1} = \begin{pmatrix} \mathcal{V}^{11} & \mathcal{V}^{12} \\ \mathcal{V}^{21} & \mathcal{V}^{22} \end{pmatrix},$$

where

$$\begin{aligned} \mathcal{V}^{11} &= (\mathbf{D}^{-1} + \mathbf{\Lambda}^T \mathbf{\Psi}^{-1} \mathbf{\Lambda}) \otimes \mathbf{I}_m, \\ \mathcal{V}^{12} &= \mathcal{V}^{21T} \\ &= -(\mathbf{\Lambda}^T \mathbf{\Psi}^{-1} \otimes \mathbf{I}_m), \\ \mathcal{V}^{22} &= \mathbf{\Psi}^{-1} \otimes \mathbf{I}_m. \end{aligned}$$

We can express (8.3) as

$$\begin{aligned}
 \log\{f(\mathbf{y}_m^{(D)}; \Gamma)\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\Psi \otimes \mathbf{I}_m)\} \right. \\
 &\quad + \mathbf{f}^T \{(\mathbf{D}^{-1} + \Lambda^T \Psi^{-1} \Lambda) \otimes \mathbf{I}_m\} \mathbf{f} - 2\mathbf{u}_g^T (\Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad \left. + \mathbf{u}_g^T (\Psi^{-1} \otimes \mathbf{I}_m) \mathbf{u}_g \right] \\
 &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\Psi \otimes \mathbf{I}_m)\} \right. \\
 &\quad + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \mathbf{f}^T (\Lambda^T \Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad \left. - 2\mathbf{u}_g^T (\Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} + \mathbf{u}_g^T (\Psi^{-1} \otimes \mathbf{I}_m) \mathbf{u}_g \right]. \tag{8.4}
 \end{aligned}$$

8.1.2 Equivalence of $\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$ and $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$

If we substitute $\mathbf{u}_g = (\Lambda \otimes \mathbf{I}_m) \mathbf{f} + \boldsymbol{\delta}$ into (8.2) then it is easily seen that $\log\{f(\mathbf{y}_2 | \mathbf{y}_m^{(D)}; \Gamma)\} = \log\{f(\mathbf{y}_2 | \mathbf{y}_m; \Gamma)\}$. Making the same substitution into (8.4) we have

$$\begin{aligned}
 \log\{f(\mathbf{y}_m^{(D)}; \Gamma)\} &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\Psi \otimes \mathbf{I}_m)\} \right. \\
 &\quad + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \mathbf{f}^T (\Lambda^T \Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad - 2\{\mathbf{f}^T (\Lambda^T \otimes \mathbf{I}_m) + \boldsymbol{\delta}^T\} (\Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad \left. + \{\mathbf{f}^T (\Lambda^T \otimes \mathbf{I}_m) + \boldsymbol{\delta}^T\} (\Psi^{-1} \otimes \mathbf{I}_m) \{(\Lambda \otimes \mathbf{I}_m) \mathbf{f} + \boldsymbol{\delta}\} \mathbf{f} \right] \\
 &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\Psi \otimes \mathbf{I}_m)\} \right. \\
 &\quad + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \mathbf{f}^T (\Lambda^T \Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad - 2\mathbf{f}^T (\Lambda^T \Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} - 2\boldsymbol{\delta}^T (\Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad + \mathbf{f}^T (\Lambda^T \Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} + 2\boldsymbol{\delta}^T (\Psi^{-1} \Lambda \otimes \mathbf{I}_m) \mathbf{f} \\
 &\quad \left. + \boldsymbol{\delta}^T (\Psi^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right] \\
 &= -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + \log\{\det(\Psi \otimes \mathbf{I}_m)\} \right. \\
 &\quad \left. + \mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} + \boldsymbol{\delta}^T (\Psi^{-1} \otimes \mathbf{I}_m) \boldsymbol{\delta} \right] \\
 &= \log\{f(\mathbf{y}_m; \Gamma)\}.
 \end{aligned}$$

Since $\log\{f(\mathbf{y}_m^{(D)}; \Gamma)\} = \log\{f(\mathbf{y}_m; \Gamma)\}$ and $\log\{f(\mathbf{y}_2 | \mathbf{y}_m^{(D)}; \Gamma)\} = \log\{f(\mathbf{y}_2 | \mathbf{y}_m; \Gamma)\}$ then it follows that $\log\{f(\mathbf{y}_c^{(D)}; \Gamma)\} = \log\{f(\mathbf{y}_c; \Gamma)\}$.

8.2 REML EM algorithm E-step for Model $\mathcal{M}$ using $\mathbf{y}_c^{(D)}$

The E-step involves taking the conditional expectation of $\log\{f(\mathbf{y}_c^{(D)}; \mathbf{\Gamma})\}$ over $k(\mathbf{f}, \mathbf{u}_g | \mathbf{y}_2; \mathbf{\Gamma}^{(w)})$ where $\mathbf{\Gamma}^{(w)}$ is the w -th iterate of $\mathbf{\Gamma}$. We can write

$$\begin{aligned} Q^{(D)}(\mathbf{\Gamma}; \mathbf{\Gamma}^{(w)}) &= E[\log\{f(\mathbf{y}_c^{(D)}; \mathbf{\Gamma})\} | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}] \\ &= E[\log\{f(\mathbf{y}_2 | \mathbf{y}_m^{(D)}; \mathbf{\Gamma})\} | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}] + E[\log\{f(\mathbf{y}_m^{(D)}; \mathbf{\Gamma})\} | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}] \\ &= Q^{(D)}(\sigma^2; \mathbf{\Gamma}^{(w)}) + Q^{(D)}(\boldsymbol{\lambda}, \psi; \mathbf{\Gamma}^{(w)}) + Q^{(D)}(\mathbf{d}; \mathbf{\Gamma}^{(w)}). \end{aligned} \quad (8.5)$$

We use the superscript D on the $Q(\cdot)$ function to distinguish between the $Q(\cdot)$ function using $\mathbf{y}_c^{(D)} = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$ and the $Q(\cdot)$ function using $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \boldsymbol{\delta}^T)^T$. To compute the $Q(\cdot)$ functions in (8.5) requires the joint conditional distribution of $\mathbf{f}$ and $\mathbf{u}_g$ given $\mathbf{y}_2$. Using (8.6) and Result A.1 we have

$$\begin{pmatrix} \mathbf{f} \\ \mathbf{u}_g \end{pmatrix} \Big| \mathbf{y}_2 \sim N \left\{ \begin{pmatrix} (\boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{y}_o \\ (\mathbf{G}_e \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{y}_o \end{pmatrix}, \begin{pmatrix} \mathbf{T}_{ff} & \mathbf{T}_{fu_g} \\ \mathbf{T}_{u_g f} & \mathbf{T}_{u_g u_g} \end{pmatrix} \right\},$$

where

$$\begin{aligned} \mathbf{T}_{ff} &= (\mathbf{D} \otimes \mathbf{I}_m) - (\mathbf{D} \boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\boldsymbol{\Lambda} \mathbf{D} \otimes \mathbf{I}_m), \\ \mathbf{T}_{fu_g} &= \mathbf{T}_{u_g f}^T = (\mathbf{D} \boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) - (\mathbf{D} \boldsymbol{\Lambda}^T \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\mathbf{G}_e \otimes \mathbf{I}_m), \\ \mathbf{T}_{u_g u_g} &= (\mathbf{G}_e \otimes \mathbf{I}_m) - (\mathbf{G}_e \otimes \mathbf{I}_m) \mathbf{Z}_g^T \mathbf{P} \mathbf{Z}_g (\mathbf{G}_e \otimes \mathbf{I}_m) \end{aligned}$$

and

$$\begin{aligned} \mathbf{P} &= \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \\ &= \mathbf{H}^{-1} - \mathbf{H}^{-1} \mathbf{X} (\mathbf{X}^T \mathbf{H}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{H}^{-1}. \end{aligned}$$

We can write

$$Q^{(D)}(\sigma^2; \mathbf{\Gamma}^{(w)}) = -\frac{1}{2} \left[\log\{\det(\mathbf{L}_2^T \mathbf{R} \mathbf{L}_2)\} + E\{(\mathbf{y}_o - \mathbf{Z}_g \mathbf{u}_g)^T \mathbf{S} (\mathbf{y}_o - \mathbf{Z}_g \mathbf{u}_g) | \mathbf{y}_2; \mathbf{\Gamma}^{(w)}\} \right]$$

and

$$\begin{aligned} Q^{(D)}(\boldsymbol{\lambda}, \psi; \boldsymbol{\Gamma}^{(w)}) &= -\frac{1}{2} \left[mp \log(\psi) + E[\mathbf{f}^T \{(\boldsymbol{\Lambda}^T \boldsymbol{\Psi}^{-1} \boldsymbol{\Lambda}) \otimes \mathbf{I}_m\} \mathbf{f} | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] \right. \\ &\quad - 2E\{\mathbf{f}^T (\boldsymbol{\Lambda}^T \boldsymbol{\Psi}^{-1} \otimes \mathbf{I}_m) \mathbf{u}_g | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}\} \\ &\quad \left. + E\{\mathbf{u}_g^T (\boldsymbol{\Psi}^{-1} \otimes \mathbf{I}_m) \mathbf{u}_g | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}\} \right] \end{aligned}$$

and

$$Q^{(D)}(\mathbf{d}; \boldsymbol{\Gamma}^{(w)}) = -\frac{1}{2} \left[\log\{\det(\mathbf{D} \otimes \mathbf{I}_m)\} + E[\mathbf{f}^T (\mathbf{D}^{-1} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}] \right].$$

8.3 REML EM algorithm M-step for Model $\mathcal{M}$ using $\mathbf{y}_c^{(D)}$

8.3.1 Update for σ^2

Differentiating $Q^{(D)}(\sigma^2; \boldsymbol{\Gamma}^{(w)})$ with respect to σ^2 we have

$$\begin{aligned} \frac{\partial Q^{(D)}(\sigma^2; \boldsymbol{\Gamma}^{(w)})}{\partial \sigma^2} &= -\frac{1}{2} \left\{ \frac{n-p}{\sigma^2} - \frac{1}{(\sigma^2)^2} (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)})^T \mathbf{K} (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)}) \right. \\ &\quad \left. - \frac{1}{(\sigma^2)^2} \text{tr}(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{u_g u_g}^{(w)}) \right\}, \end{aligned}$$

where $\mathbf{K} = \mathbf{I}_n - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$. Equating to zero and solving for σ^2 we have

$$\sigma^{2(w+1)} = \frac{1}{n-p} \left\{ (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)})^T \mathbf{K} (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)}) + \text{tr}(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{u_g u_g}^{(w)}) \right\}.$$

8.3.2 Update for λ_{jr}

Differentiating $Q^{(D)}(\boldsymbol{\lambda}, \psi; \boldsymbol{\Gamma}^{(w)})$ with respect to λ_{jr} we have

$$\begin{aligned} \frac{\partial Q^{(D)}(\boldsymbol{\lambda}, \psi; \boldsymbol{\Gamma}^{(w)})}{\partial \lambda_{jr}} &= -\frac{1}{2} \left[2E \left\{ \mathbf{f}^T \left(\boldsymbol{\Lambda}^T \boldsymbol{\Psi}^{-1} \frac{\partial \boldsymbol{\Lambda}}{\partial \lambda_{jr}} \otimes \mathbf{I}_m \right) \mathbf{f} | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)} \right\} \right. \\ &\quad \left. - 2\text{tr} \left\{ \left(\boldsymbol{\Psi}^{-1} \frac{\partial \boldsymbol{\Lambda}}{\partial \lambda_{jr}} \otimes \mathbf{I}_m \right) E \left(\mathbf{f} \mathbf{u}_g^T | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)} \right) \right\} \right]. \end{aligned}$$

Equating to zero and using the identity

$$\boldsymbol{\Lambda}^T = \sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} \mathbf{J}_{lq}^T,$$

where $\mathbf{J}_{lq}$ is a $p \times k$ indicator matrix having a 1 in row l , column q and zeroes elsewhere we obtain

$$\sum_{l=1}^p \sum_{q=1}^k \lambda_{lq} E \{ \mathbf{f}^T (\mathbf{J}_{lq}^T \boldsymbol{\Psi}^{-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)} \} = \text{tr} \{ (\boldsymbol{\Psi}^{-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) E(\mathbf{f} \mathbf{u}_g^T | \mathbf{y}_2; \boldsymbol{\Gamma}^{(w)}) \}.$$

In matrix notation this can be expressed as

$$\mathbf{A}^{(w)} \boldsymbol{\lambda} = \mathbf{b}^{(w)},$$

where $\mathbf{A}^{(w)}$ is a $mk \times mk$ matrix with elements

$$a_{jr,lq}^{(w)} = \tilde{\mathbf{f}}^{(w)T} (\mathbf{J}_{lq}^T \boldsymbol{\Psi}^{(w)-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr} \{ (\mathbf{J}_{lq}^T \boldsymbol{\Psi}^{(w)-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \}$$

and $\mathbf{b}^{(w)}$ is a $mk \times 1$ vector with elements

$$b_{jr}^{(w)} = \text{tr} \{ (\boldsymbol{\Psi}^{(w)-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{fu_g}^{(w)}) \}.$$

The update for $\boldsymbol{\lambda}$ is therefore

$$\boldsymbol{\lambda}^{(w+1)} = \mathbf{A}^{(w)-1} \mathbf{b}^{(w)}.$$

8.3.3 Update for ψ

Differentiating $Q^{(D)}(\boldsymbol{\lambda}, \psi; \boldsymbol{\Gamma}^{(w)})$ with respect to ψ we have

$$\begin{aligned} \frac{\partial Q^{(D)}(\boldsymbol{\lambda}, \psi; \boldsymbol{\Gamma}^{(w)})}{\partial \psi} = & -\frac{1}{2} \left[\frac{mp}{\psi} - \frac{1}{\psi^2} \tilde{\mathbf{f}}^{(w)T} (\boldsymbol{\Lambda}^T \boldsymbol{\Lambda} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} \right. \\ & - \frac{1}{\psi^2} \text{tr}\{(\boldsymbol{\Lambda}^T \boldsymbol{\Lambda} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)}\} \\ & - \frac{1}{\psi^2} \tilde{\mathbf{u}}_g^{(w)T} \tilde{\mathbf{u}}_g^{(w)} - \frac{1}{\psi^2} \text{tr}(\mathbf{T}_{u_g u_g}^{(w)}) \\ & \left. + \frac{2}{\psi^2} \text{tr}\{(\boldsymbol{\Lambda} \otimes \mathbf{I}_m)(\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{f u_g}^{(w)})\} \right]. \end{aligned}$$

Equating to zero and solving for ψ we have

$$\begin{aligned} \psi^{(w+1)} = & \frac{1}{mp} \left[\tilde{\mathbf{f}}^{(w)T} (\boldsymbol{\Lambda}^{(w)T} \boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr}\{(\boldsymbol{\Lambda}^{(w)T} \boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)}\} \right. \\ & \left. + \tilde{\mathbf{u}}_g^{(w)T} \tilde{\mathbf{u}}_g^{(w)} + \text{tr}(\mathbf{T}_{u_g u_g}^{(w)}) - 2 \text{tr}\{(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m)(\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{f u_g}^{(w)})\} \right]. \end{aligned}$$

8.3.4 Update for d_{qr}

Since $Q^{(D)}(\mathbf{d}; \boldsymbol{\Gamma}^{(w)})$ is equivalent to $Q(\mathbf{d}; \boldsymbol{\Gamma}^{(w)})$ the update for d_{qr} , $q = 1, \dots, k$, $r = 1, \dots, j$, $q \geq r$ is the same as that presented in Chapter 6 for Model $\mathcal{M}$, i.e.,

$$d_{qr}^{(w+1)} = \frac{1}{m} \left\{ \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)} + \text{tr}(\mathbf{T}_{ff_{qr}}^{(w)}) \right\}.$$

8.4 REML PX EM algorithm - Model $\mathcal{Z}$

We considered Model $\mathcal{Z}$ in Chapter 7. We provide a succinct overview here for convenience. We define

$$\text{Model } \mathcal{Z} \quad \mathbf{y}_o = \mathbf{X} \boldsymbol{\beta}_* + \mathbf{Z}_g \mathbf{u}_g + \mathbf{e},$$

where $\mathbf{e}$ is distributed $N(\mathbf{0}, \mathbf{R}_*)$ and $\mathbf{R}_* = \sigma_*^2 \mathbf{I}_n$. For Model $\mathcal{Z}$ we consider

$$\mathbf{u}_g = (\mathbf{\Lambda}_* \otimes \mathbf{I}_m) \mathbf{f} + (\mathbf{\Omega} \otimes \mathbf{I}_m) \boldsymbol{\delta},$$

where $\mathbf{\Omega} = \mathbf{\Omega}(\boldsymbol{\omega})$ is a $p \times p$ diagonal positive definite diagonal matrix. When we assume common specific variances $\mathbf{\Omega}$ has the form $\mathbf{\Omega} = \omega \mathbf{I}_p$. We assume $\mathbf{f} \sim N(\mathbf{0}, \mathbf{D}_* \otimes \mathbf{I}_m)$ where $\mathbf{D}_* = \mathbf{D}_*(\mathbf{d}_*)$ and $\boldsymbol{\delta} \sim N(\mathbf{0}, \mathbf{\Psi}_* \otimes \mathbf{I}_m)$ where $\mathbf{\Psi}_* = \psi_* \mathbf{I}_p$. The $mk \times 1$ vector $\mathbf{f}$ and the $mp \times 1$ vector $\boldsymbol{\delta}$ are assumed to be independent of each other. We define the expanded parameter set associated with Model $\mathcal{Z}$ as $\boldsymbol{\mathcal{T}} = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \psi_*, \mathbf{d}_*^T, \boldsymbol{\omega}^T)^T$ where the subscript $*$ is used to distinguish parameters associated with Model $\mathcal{M}$ and Model $\mathcal{Z}$. For Model $\mathcal{Z}$ the distribution of the observed data is

$$\mathbf{y}_o \sim N(\mathbf{X}\boldsymbol{\beta}_*, \mathbf{H}_*),$$

where

$$\mathbf{H}_* = \mathbf{Z}_g \mathbf{G}_{e_*} \mathbf{Z}_g^T + \mathbf{R}_*$$

and

$$\mathbf{G}_{e_*} = \{(\mathbf{\Lambda}_* \mathbf{D}_* \mathbf{\Lambda}_*^T + \omega^2 \psi_* \mathbf{I}_p) \otimes \mathbf{I}_m\}.$$

Using the parameter expanded model we identify the parameters associated with Model $\mathcal{M}$ as $\boldsymbol{\Gamma} = (\sigma^2, \boldsymbol{\lambda}^T, \psi, \mathbf{d}^T)^T = \mathcal{R}(\boldsymbol{\mathcal{T}}) = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \omega^2 \psi_*, \mathbf{d}_*^T)^T$ where $\mathcal{R}$ is the reduction function from the expanded parameter set to the original parameter set.

8.4.1 Complete data density function

For Model $\mathcal{Z}$ the complete data specification is the same as that for Model $\mathcal{M}$, i.e., $\mathbf{y}_c = (\mathbf{y}_2^T, \mathbf{f}^T, \mathbf{u}_g^T)^T$. We use the subscript $\mathcal{Z}$ to distinguish between the complete data log-density function and $Q(\cdot)$ function for Model $\mathcal{M}$ and Model $\mathcal{Z}$. To form the complete data density function for Model $\mathcal{Z}$ requires the joint distribution

$$\begin{pmatrix} \mathbf{y}_2 \\ \mathbf{f} \\ \mathbf{u}_g \end{pmatrix} \sim N \left\{ \begin{pmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{V}_{11*} & \mathbf{V}_{12*} & \mathbf{V}_{13*} \\ \mathbf{V}_{21*} & \mathbf{V}_{22*} & \mathbf{V}_{23*} \\ \mathbf{V}_{31*} & \mathbf{V}_{32*} & \mathbf{V}_{33*} \end{pmatrix} \right\}, \quad (8.6)$$

where

$$\begin{aligned}
 V_{11*} &= L_2^T H_* L_2, \\
 V_{22*} &= D_* \otimes I_m, \\
 V_{33*} &= G_{e*} \otimes I_m, \\
 V_{12*} &= V_{21*}^T = L_2^T Z_g (\Lambda_* D_* \otimes I_m), \\
 V_{13*} &= V_{31*}^T = L_2^T Z_g (G_{e*} \otimes I_m) \\
 V_{23*} &= V_{32*}^T = D_* \Lambda_*^T \otimes I_m.
 \end{aligned}$$

The log-density function for $\mathbf{y}_2$ given $\mathbf{y}_m^{(D)}$ is (excluding constants)

$$\begin{aligned}
 \log\{f_Z(\mathbf{y}_2|\mathbf{y}_m^{(D)}; \Gamma)\} &= -\frac{1}{2} \left[\log\{\det(L_2^T R_* L_2)\} \right. \\
 &\quad \left. + (\mathbf{y}_o - Z_g \mathbf{u}_g)^T S_* (\mathbf{y}_o - Z_g \mathbf{u}_g) \right],
 \end{aligned}$$

where $S_* = R_*^{-1} - R_*^{-1} X (X^T R_*^{-1} X)^{-1} X^T R_*^{-1}$. Now consider the log-density function $\log\{f_Z(\mathbf{y}_m^{(D)}; \Gamma)\}$. We define

$$\mathcal{V}_* = \begin{Bmatrix} D_* \otimes I_m & D_* \Lambda_*^T \otimes I_m \\ \Lambda_* D_* \otimes I_m & G_{e*} \otimes I_m \end{Bmatrix}.$$

The log-density function for $\mathbf{y}_m^{(D)}$ is

$$\log\{f_Z(\mathbf{y}_m^{(D)}; \Gamma)\} = -\frac{1}{2} \left[\log\{\det(\mathcal{V}_*)\} + \mathbf{y}_m^{(D)T} \mathcal{V}_*^{-1} \mathbf{y}_m^{(D)} \right].$$

Using the fourth identity in Result A.3.3 we can write

$$\det(\mathcal{V}_*) = \det(D_* \otimes I_m) \det(\Omega \Psi_* \Omega \otimes I_m).$$

Using Result A.3.6 we can write

$$\mathcal{V}_*^{-1} = \begin{pmatrix} \mathcal{V}_*^{11} & \mathcal{V}_*^{12} \\ \mathcal{V}_*^{21} & \mathcal{V}_*^{22} \end{pmatrix},$$

where

$$\begin{aligned}\mathcal{V}_*^{11} &= \{D_*^{-1} + \Lambda_*^T(\Omega\Psi_*\Omega)^{-1}\Lambda_*\} \otimes I_m, \\ \mathcal{V}_*^{12} &= \mathcal{V}_*^{21^T} = -\{\Lambda_*^T(\Omega\Psi_*\Omega)^{-1} \otimes I_m\}, \\ \mathcal{V}_*^{22} &= (\Omega\Psi_*\Omega)^{-1} \otimes I_m.\end{aligned}$$

Therefore,

$$\begin{aligned}\log\{f_{\mathcal{Z}}(\mathbf{y}_m^{(D)}; \Gamma)\} &= -\frac{1}{2} \left[\log\{\det(D_* \otimes I_m)\} + \log\{\det(\Omega\Psi_*\Omega \otimes I_m)\} \right. \\ &\quad + \mathbf{f}^T [\{D_*^{-1} + \Lambda_*^T(\Omega\Psi_*\Omega)^{-1}\Lambda_*\} \otimes I_m] \mathbf{f} - 2\mathbf{u}_g^T \{(\Omega\Psi_*\Omega)^{-1}\Lambda_* \otimes I_m\} \mathbf{f} \\ &\quad \left. + \mathbf{u}_g^T \{(\Omega\Psi_*\Omega)^{-1} \otimes I_m\} \mathbf{u}_g \right].\end{aligned}$$

8.5 REML PX EM algorithm E-step for Model $\mathcal{Z}$ using $\mathbf{y}_c^{(D)}$

For Model $\mathcal{Z}$ we define $\mathcal{T}^{(w)} = (\sigma_*^2, \boldsymbol{\lambda}_*^{(w)T}, \psi_*^{(w)}, \mathbf{d}_*^{(w)T}, \omega_0)^T$ where $\Omega(\omega_0) = I_p$. The E-step REML PX EM algorithm for Model $\mathcal{Z}$ and the E-step of the REML EM algorithm for Model $\mathcal{M}$ are the same.

8.6 REML PX EM algorithm M-step for Model $\mathcal{Z}$ using $\mathbf{y}_c^{(D)}$

The M-step of the REML PX EM algorithm involves maximising $Q_{\mathcal{Z}}(\mathcal{T}; \mathcal{T}^{(w)})$ with respect to $\mathcal{T} = (\sigma_*^2, \boldsymbol{\lambda}_*^T, \psi_*, \mathbf{d}_*^T, \omega)^T$ and then applying the reduction function $\mathcal{R}(\mathcal{T}) = (\sigma_*^2, \boldsymbol{\lambda}_*, \omega^2\psi_*, \mathbf{d}_*^T)^T$ to obtain $\Gamma = (\sigma^2, \boldsymbol{\lambda}^T, \psi, \mathbf{d}^T)^T$.

8.6.1 Update for σ_*^2

The update for σ_*^2 is the same as the update for σ^2 under Model $\mathcal{M}$, i.e.,

$$\sigma_*^{2(w+1)} = -\frac{1}{n-p} \left\{ (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)})^T \mathbf{K} (\mathbf{y}_o - \mathbf{Z}_g \tilde{\mathbf{u}}_g^{(w)}) + \text{tr}(\mathbf{Z}_g^T \mathbf{K} \mathbf{Z}_g \mathbf{T}_{u_g u_g}^{(w)}) \right\}.$$

Applying the reduction function we have $\sigma_*^{2(w+1)} = \sigma_*^{2(w+1)}$.

8.6.2 Update for λ_{*jr}

The update for λ_{*jr} is the same as that for λ_{jr} under Model $\mathcal{M}$, i.e.,

$$\boldsymbol{\lambda}_*^{(w+1)} = \mathbf{A}_*^{(w)-1} \mathbf{b}_*^{(w)},$$

where the elements of the $mk \times mk$ matrix $\mathbf{A}_*^{(w)}$ are

$$a_{jr,lq}^{(w)} = \tilde{\mathbf{f}}^{(w)T} (\mathbf{J}_{lq}^T \boldsymbol{\Psi}^{(w)-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr}\{(\mathbf{J}_{lq}^T \boldsymbol{\Psi}^{(w)-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)}\}$$

and $\mathbf{b}^{(w)}$ is a $mk \times 1$ vector with elements

$$b_{jr}^{(w)} = \text{tr}\{\boldsymbol{\Psi}^{-1} \mathbf{J}_{jr} \otimes \mathbf{I}_m\} (\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{fu_g}^{(w)}).$$

Applying the reduction function we have $\boldsymbol{\lambda}^{(w+1)} = \boldsymbol{\lambda}_*^{(w+1)}$.

8.6.3 Update for ψ_*

The update for ψ_* is the same as that for ψ under Model $\mathcal{M}$, i.e.,

$$\begin{aligned} \psi_*^{(w+1)} = & \frac{1}{mp} \left[\tilde{\mathbf{f}}^{(w)T} (\boldsymbol{\Lambda}^{(w)T} \boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr}\{(\boldsymbol{\Lambda}^{(w)T} \boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)}\} \right. \\ & \left. + \tilde{\mathbf{u}}_g^{(w)T} \tilde{\mathbf{u}}_g^{(w)} + \text{tr}(\mathbf{T}_{u_g u_g}^{(w)}) - 2\text{tr}\{(\boldsymbol{\Lambda}^{(w)} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{fu_g}^{(w)})\} \right]. \end{aligned}$$

Applying the reduction function we have $\psi^{(w+1)} = \omega^{2(w+1)} \psi_*^{(w+1)}$.

8.6.4 Update for d_{*qr}

The update for d_{*qr} is the same as that for d_{qr} under Model $\mathcal{M}$, i.e.,

$$d_{*qr}^{(w+1)} = \frac{1}{m} \left\{ \tilde{\mathbf{f}}_q^{(w)T} \tilde{\mathbf{f}}_r^{(w)} + \text{tr}(\mathbf{T}_{ffqr}^{(w)}) \right\}.$$

Applying the reduction function we have $d_{qr}^{(w+1)} = d_{*qr}^{(w+1)}$.

8.6.5 Update for ω

Differentiating $Q_Z(\mathcal{T}; \mathcal{T}^{(w)})$ with respect to ω we have

$$\begin{aligned} \frac{\partial Q_Z(\mathcal{T}; \mathcal{T}^{(w)})}{\partial \omega} = & -\frac{1}{2} \left[\frac{2mp}{\omega} - \frac{2}{\psi \omega^3} E\{ \mathbf{f}^T (\Lambda_*^T \Lambda_* \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \mathcal{T}^{(w)} \} \right. \\ & \left. - \frac{2}{\psi \omega^3} E(\mathbf{u}_g^T \mathbf{u}_g | \mathbf{y}_2; \mathcal{T}^{(w)}) + \frac{4}{\psi \omega^3} E\{ \mathbf{u}_g^T (\Lambda_* \otimes \mathbf{I}_m) \mathbf{f} | \mathbf{y}_2; \mathcal{T}^{(w)} \} \right]. \end{aligned}$$

Equating to zero and solving for ω^2 we have

$$\begin{aligned} \omega^{2(w+1)} = & \frac{1}{mp\psi^{(w)}} \left[\tilde{\mathbf{f}}^{(w)T} (\Lambda^{(w)T} \Lambda^{(w)} \otimes \mathbf{I}_m) \tilde{\mathbf{f}}^{(w)} + \text{tr}\{ (\Lambda^{(w)T} \Lambda^{(w)} \otimes \mathbf{I}_m) \mathbf{T}_{ff}^{(w)} \} \right. \\ & \left. + \tilde{\mathbf{u}}_g^{(w)T} \tilde{\mathbf{u}}_g^{(w)} + \text{tr}(\mathbf{T}_{u_g u_g}^{(w)}) - 2\text{tr}\{ (\Lambda^{(w)} \otimes \mathbf{I}_m) (\tilde{\mathbf{f}}^{(w)} \tilde{\mathbf{u}}_g^{(w)T} + \mathbf{T}_{fu_g}^{(w)}) \} \right] \end{aligned}$$

8.7 Example: National Variety Trial data

In the following example and simulation study when we refer to Model $\mathcal{O}$ or Model $\mathcal{M}$ it should be read as “the REML EM algorithm used to estimate the variance parameters associated with Model A ” where Model A is Model $\mathcal{O}$ or Model $\mathcal{M}$ as the case may be. When we refer to Model $\mathcal{X}$ or Model $\mathcal{Z}$ it should be read as “the REML PX EM algorithm used to estimate the variance parameters associated with Model B ” where Model B is Model $\mathcal{X}$ or Model $\mathcal{Z}$ as the case may be. We consider the small subset of National Variety Trial data used in the previous two chapters and presented in Table 6.2. For every model fitted to the simulated data we define the starting values $\boldsymbol{\lambda}^{(0)} = (0.1, 0.1, 0.1, 0.1, 0.1, 0.1, 0.1)^T$, $\psi^{(0)} = 0.1$, $\sigma^{2(0)} = 0.1$, and $d^{(0)} = 1$, $\omega^{(0)} = 1$ when required.

Results for Model $\mathcal{O}$, $\mathcal{M}$, $\mathcal{X}$, and Model $\mathcal{Z}$ with $k = 1$ factors and using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ are provided in Table 8.1. In the interests of brevity, parameter updates for Model $\mathcal{O}$ and Model $\mathcal{X}$ using the missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ were omitted from this chapter, however they can be derived using information in this and previous chapters.

Comparing the results in Table 8.1 with the results in Table 6.1 and Table 7.1 we can see that for all models and levels of ε that an algorithm using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ outperforms an algorithm using the missing data specification $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ in terms of iterations until convergence. If we had to choose a model and a value of ε then we would choose Model $\mathcal{Z}$ based on the new missing data specification and $\varepsilon = 10^{-5}$. This model results in significantly fewer iterations to convergence than the next best option. When the stopping rule is based on $\varepsilon = 10^{-5}$ there is no loss in terms of the final residual log-likelihood and the Frobenious matrix norm of the variance-covariance matrix is accurate to 5 decimal places. This model and stopping rule combination converges in 23 iterations which is acceptable from a practical point of view. In terms of iterations until convergence and accuracy of the estimated genetic variance-covariance matrix, Model $\mathcal{Z}$ using the new missing data specification compares favourably to the next best model options, which are Model $\mathcal{M}$ and Model $\mathcal{X}$ using the new missing data specification.

8.8 Simulation study

Model $\mathcal{Z}$ using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ was identified as the best model when applied to the small subset of National Variety Trial data presented in Table 6.2 and a stopping rule with $\varepsilon = 10^{-5}$ was identified as being sufficient in determining convergence. We consider a simulation study to confirm these findings.

The simulation study involved Model $\mathcal{O}$, $\mathcal{M}$, $\mathcal{X}$, and Model $\mathcal{Z}$ using both $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ and $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$. Each model was fitted to 3 lots of 100 randomly generated data sets. Each lot of 100 were generated using $\sigma^2 = 0.05$ and $d = 1$ but differed in the specification of the genetic variance-covariance matrix. The three different genetic variance-covariance matrices used to generate the data can be broadly classified as consisting of high, medium, and low correlations. They were generated by considering the same vector of factor loadings, i.e.,

$\boldsymbol{\lambda} = (0.1, 0.5, -0.05, -0.1, 0.2, 0.05, 0.005)^T$ but three different values for the com-

Table 8.1: Iterations until convergence when using a stopping rule with varying levels of ϵ , final residual log-likelihood (ℓ_R), and the Frobenius matrix norm of the estimated genetic variance-covariance matrix when using a REML EM algorithm to fit Model $\mathcal{O}$ and Model $\mathcal{M}$ and a REML PX EM algorithm to fit Model $\mathcal{X}$ and $\mathcal{Z}$ with $k = 1$ factors to the small subset of National Variety Trial data presented in Table 6.2. All algorithms use the missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$.

ϵ	Model $\mathcal{O}$			Model $\mathcal{X}$			Model $\mathcal{M}$			Model $\mathcal{Z}$		
	Iter.	ℓ_R	$\ \hat{\mathbf{G}}_e\ _F$	Iter.	ℓ_R	$\ \hat{\mathbf{G}}_e\ _F$	Iter.	ℓ_R	$\ \hat{\mathbf{G}}_e\ _F$	Iter.	ℓ_R	$\ \hat{\mathbf{G}}_e\ _F$
10^{-8}	66	300.9849	0.3180123	54	300.9849	0.3180123	54	300.9849	0.3180123	40	300.9849	0.3180123
10^{-7}	57	300.9849	0.3180122	46	300.9849	0.3180124	46	300.9849	0.3180124	35	300.9849	0.3180123
10^{-6}	48	300.9849	0.3180113	38	300.9849	0.3180129	38	300.9849	0.3180129	29	300.9849	0.3180124
10^{-5}	39	300.9849	0.3180019	30	300.9849	0.3180189	30	300.9849	0.3180189	23	300.9849	0.3180132
10^{-4}	30	300.9849	0.3179079	22	300.9849	0.3180807	23	300.9849	0.3180634	18	300.9849	0.3180216
10^{-3}	21	300.9842	0.3169733	15	300.9802	0.3185419	16	300.9824	0.3184078	13	300.9830	0.3180960
10^{-2}	13	300.9369	0.3101104	10	300.8763	0.3200647	10	300.8763	0.3200647	9	300.9220	0.3178401

mon specific variance, i.e., $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$. The three genetic variance-covariance (and correlation) matrices that were used to generate the data are presented in Table 8.2. When using the common specific variance $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$ the correlations in the genetic correlation matrix have ranges of $(-0.69, 0.88)$, $(-0.37, 0.61)$, and $(-0.14, 0.26)$ respectively.

For every model fitted to the simulated data we define the starting values $\boldsymbol{\lambda}^{(0)} = (0.1, 0.1, 0.1, 0.1, 0.1, 0.1, 0.1)^T$, $\psi^{(0)} = 0.1$, $\sigma^{2(0)} = 0.1$, and $d^{(0)} = 1$, $\omega^{(0)} = 1$ when required. The minimum, median, and maximum number of iterations until convergence for all models using the two different missing data specifications and varying levels of ε are provided in Table 8.3. Model $\mathcal{Z}$ is the preferred model for both missing data specifications and at all levels of ε , regardless of the level of ψ . Model $\mathcal{Z}$ using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ is preferred to Model $\mathcal{Z}$ when using the missing data specification $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ at all levels of ε and when $\psi = 0.05$ and $\psi = 0.25$, i.e., when the genetic variance-covariance matrix consists of low to medium correlations.

An interesting outcome of the simulation study is the performance of Model $\mathcal{Z}$ using the missing data specification $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ when $\psi = 0.01$, i.e., when the genetic variance-covariance matrix consists of high correlations. This model and missing data specification significantly outperforms all other models. For levels of ε smaller than 10^{-5} the performance gain, in terms of iterations until convergence, is between $\frac{1}{2}$ to almost $\frac{1}{3}$ the number of iterations until convergence of the next best model, which is Model $\mathcal{Z}$ using the new missing data specification. A caveat on this result is that performance is judged only on the number of iterations until convergence and does not consider computing time. Although we have not considered a comprehensive computing strategy for each model in this thesis, the time per iteration when using the missing data specification $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ is approximately 2.5 times greater than time per iteration when using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$. Therefore, the two have comparable performance in terms of elapsed time to convergence when $\psi = 0.01$. In all other cases Model $\mathcal{Z}$ using the new missing data specification is far superior in terms of elapsed time to convergence than all other models. The reduction in time per iteration for models using the new missing data specification concurs with results presented in Thompson et al. (2003). Thompson et al. (2003) found that time per iteration was substantially reduced when applying the average information algorithm to a “dependent” formulation of the factor analytic mixed model.

In the small subset of National Variety Trial data example it was found that using the less stringent convergence criteria $\varepsilon = 10^{-5}$ had no effect on the final residual log-likelihood and little effect on the estimate of the genetic variance-covariance matrix. In Table 8.4 summary statistics for the absolute difference between the Frobenius matrix norm of the estimated genetic variance-covariance matrix using $\varepsilon = 10^{-8}$ and the Frobenius matrix norm of the estimated genetic variance covariance matrix using another level of ε are presented. If we consider the standard to be the Frobenius matrix norm of the estimated variance-covariance matrix obtained when using $\varepsilon = 10^{-8}$ then the data in Table 8.4 supports an increase in ε to 10^{-5} . For Model $\mathcal{Z}$ and $\varepsilon = 10^{-5}$ the worst case is when using the new missing data specification and when $\psi = 0.01$. In this case the median value of the difference between the Frobenius matrix norm of the estimated variance-covariance matrix using $\varepsilon = 10^{-8}$ and $\varepsilon = 10^{-5}$ is 0.000028. This difference at the other two levels of ψ is less than 0.00001. These differences indicate that from a practical point of view there is little loss in terms of an accurate estimate of the genetic variance-covariance matrix in moving to a less stringent stopping rule based on $\varepsilon = 10^{-5}$.

We recommend Model $\mathcal{Z}$ using the new missing data specification when fitting a factor analytic mixed model. In Figure 8.1 we compare the number of iterations until convergence for Model $\mathcal{Z}$ using $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ and $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$. The convergence criteria for both was based on $\varepsilon = 10^{-5}$. Model $\mathcal{Z}$ using the new missing data specification and applied to data generated using $\psi = 0.05$ and $\psi = 0.25$ converges (on average) in 54% and 16% respectively of the number of iterations it takes Model $\mathcal{Z}$ using $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$ to converge. When $\psi = 0.05$ and $\psi = 0.25$ the median (50th percentile in Figure 8.1) number of iterations to convergence is 15 and 25 iterations respectively for Model $\mathcal{Z}$ using the new missing data specification. Both are acceptable from a practical point of view. When $\psi = 0.01$ the median number of iterations to convergence is 56. This is unacceptably high and greater than the median number of iterations to convergence for Model $\mathcal{Z}$ using $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$. However, when $\psi = 0.01$ the two missing data specifications for Model $\mathcal{Z}$ are comparable in terms of elapsed time to convergence. It is also worth noting that increasing ε to 10^{-4} reduces the median number of iterations until convergence of Model $\mathcal{Z}$ using $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ to 28 iterations. This is a feasible alternative if we are prepared to accept an estimate of the genetic variance-covariance matrix which is slightly less accurate than that obtained when using $\varepsilon = 10^{-5}$.

Table 8.2: The three genetic variance-covariance matrices based on $\lambda = (0.1, 0.5, -0.05, -0.1, 0.2, 0.05, 0.005)^T$ and $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$ used to generate the three lots of 100 data sets. Variances are on the diagonals, covariances in the lower triangle, and correlations in the upper triangle.

(a) $\psi = 0.01$						
0.0200	0.69	-0.32	-0.50	0.63	0.32	0.04
0.0500	0.2600	-0.44	-0.69	0.88	0.44	0.05
-0.0050	-0.0250	0.0125	0.32	-0.40	-0.20	-0.02
-0.0100	-0.0500	0.0050	0.0200	-0.63	-0.32	-0.04
0.0200	0.1000	-0.0100	-0.0200	0.0500	0.40	0.04
0.0050	0.0250	-0.0025	-0.0050	0.0100	0.0125	0.02
0.0005	0.0025	-0.0002	-0.0005	0.0010	0.0002	0.0100

(b) $\psi = 0.05$						
0.0600	0.37	-0.09	-0.17	0.27	0.09	0.01
0.0500	0.3000	-0.20	-0.37	0.61	0.20	0.02
-0.0050	-0.0250	0.0525	0.09	-0.15	-0.05	-0.00
-0.0100	-0.0500	0.0050	0.0600	-0.27	-0.09	-0.01
0.0200	0.1000	-0.0100	-0.0200	0.0900	0.15	0.01
0.0050	0.0250	-0.0025	-0.0050	0.0100	0.0525	0.00
0.0005	0.0025	-0.0002	-0.0005	0.0010	0.0002	0.0500

(c) $\psi = 0.25$						
0.2600	0.14	-0.02	-0.04	0.07	0.02	0.00
0.0500	0.5000	-0.07	-0.14	0.26	0.07	0.01
-0.0050	-0.0250	0.2525	0.02	-0.04	-0.01	-0.00
-0.0100	-0.0500	0.0050	0.2600	-0.07	-0.02	-0.00
0.0200	0.1000	-0.0100	-0.0200	0.2900	0.04	0.00
0.0050	0.0250	-0.0025	-0.0050	0.0100	0.2525	0.00
0.0005	0.0025	-0.0002	-0.0005	0.0010	0.0002	0.2500

Table 8.3: Minimum, median, and maximum number of iterations until convergence for Model $\mathcal{O}$, $\mathcal{M}$, $\mathcal{X}$, and $\mathcal{Z}$ using two different missing data specifications and different levels of ϵ . A 100 data sets were generated for each level of the common specific variance, i.e., $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$.

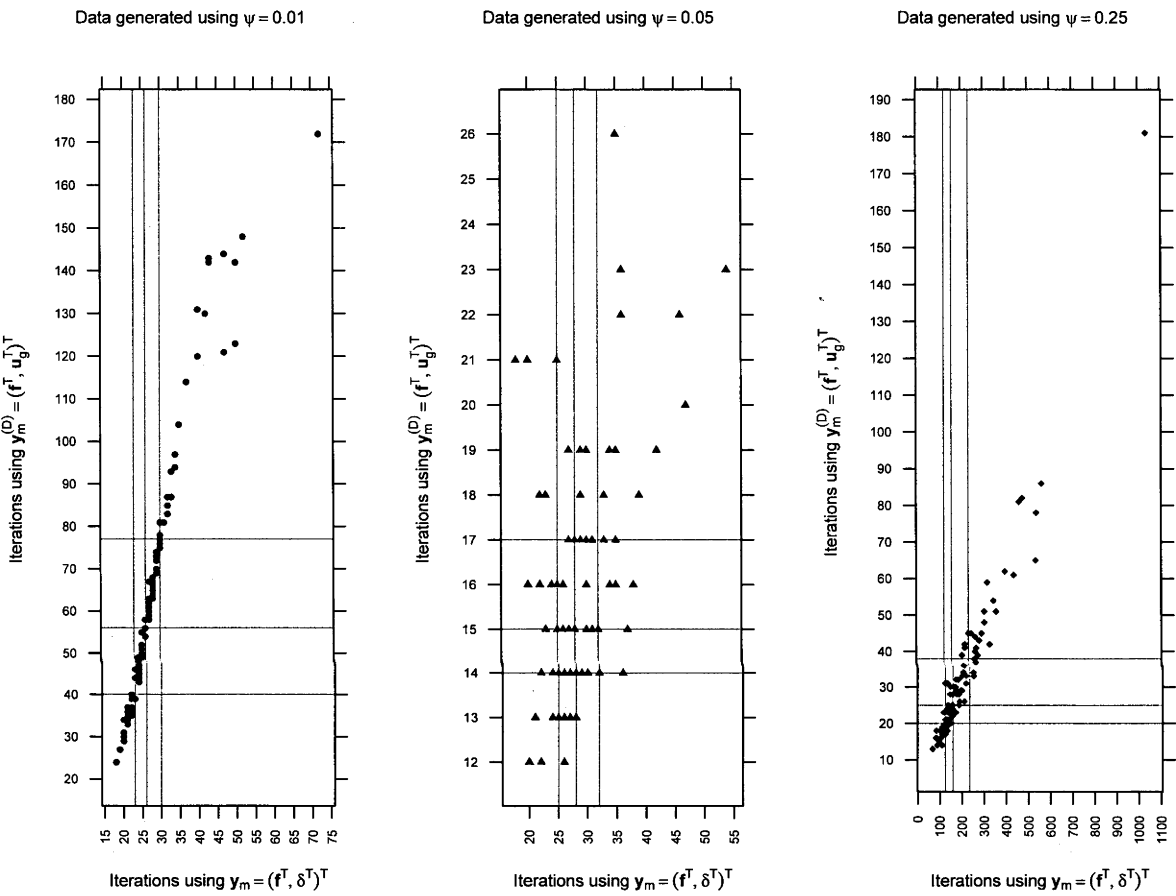
	Missing data specification $\mathbf{y}_m = (f^T, \delta^T)^T$												Missing data specification $\mathbf{y}_m^{(D)} = (f^T, \mathbf{u}_g^T)^T$											
	Model $\mathcal{O}$			Model $\mathcal{M}$			Model $\mathcal{X}$			Model $\mathcal{Z}$			Model $\mathcal{O}$			Model $\mathcal{M}$			Model $\mathcal{X}$			Model $\mathcal{Z}$		
	min	med	max	min	med	max	min	med	max	min	med	max	min	med	max	min	med	max	min	med	max	min	med	max
$\epsilon = 10^{-8}$																								
$\psi = 0.01$	119	299	2500	94	289	2500	94	295	2500	36	57	816	137	339	2500	93	290	2500	93	294	2500	51	157	2500
$\psi = 0.05$	86	144	235	48	58	114	51	64	130	34	54	114	50	76	121	33	44	77	32	44	77	23	28	50
$\psi = 0.25$	265	447	2500	145	387	2500	177	451	2500	146	387	2500	35	50	364	26	47	339	27	50	364	26	46	340
$\epsilon = 10^{-7}$																								
$\psi = 0.01$	100	232	2500	77	222	2500	76	229	2500	29	46	479	115	276	2500	76	224	2500	75	228	2500	41	123	2500
$\psi = 0.05$	73	122	198	40	48	93	42	54	109	28	45	93	42	64	102	27	36	63	27	36	63	19	23	41
$\psi = 0.25$	224	370	2492	116	315	2067	145	376	2496	117	315	2058	30	41	311	22	38	285	23	42	311	21	38	286
$\epsilon = 10^{-6}$																								
$\psi = 0.01$	82	167	1126	61	157	873	60	163	1125	24	36	208	94	214	1550	59	159	970	59	162	1124	33	89	758
$\psi = 0.05$	61	101	162	33	39	74	34	45	90	23	36	74	36	54	84	22	30	50	22	30	50	16	19	33
$\psi = 0.25$	184	300	1982	91	231	1557	118	306	1986	91	231	1548	25	35	258	18	32	232	19	35	258	17	32	233
$\epsilon = 10^{-5}$																								
$\psi = 0.01$	64	104	249	44	92	203	44	98	247	18	26	72	72	154	651	43	93	223	42	97	246	24	56	172
$\psi = 0.05$	49	81	126	25	31	54	28	37	70	18	28	54	29	43	66	17	23	37	17	23	37	12	15	26
$\psi = 0.25$	144	226	1472	67	157	1046	92	225	1476	67	158	1037	21	28	205	15	25	180	15	28	205	13	25	181
$\epsilon = 10^{-4}$																								
$\psi = 0.01$	47	59	74	28	41	60	28	44	63	13	16	24	51	93	229	27	42	66	26	43	62	16	28	48
$\psi = 0.05$	38	60	90	19	22	35	21	28	51	13	20	35	23	32	47	12	17	25	12	16	25	9	11	18
$\psi = 0.25$	104	156	961	44	90	536	68	155	965	44	91	528	16	22	153	10	19	127	11	22	153	10	19	128
$\epsilon = 10^{-3}$																								
$\psi = 0.01$	29	35	46	13	15	19	13	15	19	8	9	10	30	43	59	12	15	20	12	14	18	8	11	16
$\psi = 0.05$	26	39	54	13	15	20	14	20	32	9	13	19	16	21	29	8	10	15	8	11	16	6	8	11
$\psi = 0.25$	59	92	441	28	48	260	44	88	445	29	49	261	11	16	100	7	12	74	8	16	100	7	12	75

Table 8.4: The absolute difference between $\|\hat{\mathbf{G}}_e\|_F$ using $\varepsilon = 10^{-8}$ and $\|\hat{\mathbf{G}}_e\|_F$ using another level of ε was calculated for the 100 generated data sets for each level of the common specific variance, i.e., $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$. The minimum, median, and maximum of these absolute differences are provided below. We have used the notation $|(8) - (a)|$ where (8) refers to $\varepsilon = 10^{-8}$ and (a) is one of the other levels of ε . These summary statistics have been calculated for Model $\mathcal{O}$, $\mathcal{M}$, $\mathcal{X}$, and $\mathcal{Z}$ using the two different missing data specifications.

	Model $\mathcal{O}$			Missing data specification $\mathbf{y}_m = (\mathbf{f}^T, \boldsymbol{\delta}^T)^T$			Model $\mathcal{X}$			Model $\mathcal{Z}$		
	min	med	max	min	med	max	min	med	max	min	med	max
$\psi = 0.01$												
(8) - (7)	0.000000	0.000000	0.000008	0.000000	0.000001	0.000016	0.000000	0.000000	0.000009	0.000000	0.000000	0.000003
(8) - (6)	0.000000	0.000004	0.000143	0.000001	0.000006	0.000211	0.000001	0.000005	0.000145	0.000000	0.000001	0.000028
(8) - (5)	0.000000	0.000041	0.000524	0.000012	0.000062	0.000677	0.000012	0.000051	0.000559	0.000005	0.000011	0.000140
(8) - (4)	0.000006	0.000336	0.001168	0.000113	0.000488	0.001591	0.000113	0.000434	0.001418	0.000043	0.000107	0.000513
(8) - (3)	0.000649	0.004479	0.012572	0.000885	0.001959	0.003221	0.000760	0.002009	0.003221	0.000014	0.000581	0.001471
$\psi = 0.05$												
(8) - (7)	0.000000	0.000001	0.000001	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
(8) - (6)	0.000001	0.000006	0.000015	0.000000	0.000000	0.000000	0.000000	0.000000	0.000002	0.000000	0.000000	0.000002
(8) - (5)	0.000009	0.000061	0.000155	0.000000	0.000003	0.000023	0.000000	0.000001	0.000015	0.000000	0.000002	0.000018
(8) - (4)	0.000084	0.000602	0.001539	0.000000	0.000035	0.000191	0.000000	0.000005	0.000152	0.000000	0.000027	0.000195
(8) - (3)	0.000965	0.005895	0.015257	0.000004	0.000277	0.001178	0.000000	0.000056	0.001086	0.000004	0.000452	0.001667
$\psi = 0.25$												
(8) - (7)	0.000000	0.000000	0.000003	0.000000	0.000000	0.000001	0.000000	0.000000	0.000000	0.000000	0.000000	0.000001
(8) - (6)	0.000000	0.000001	0.000033	0.000000	0.000001	0.000007	0.000000	0.000000	0.000001	0.000000	0.000001	0.000008
(8) - (5)	0.000000	0.000020	0.000331	0.000000	0.000011	0.000072	0.000000	0.000001	0.000006	0.000000	0.000012	0.000077
(8) - (4)	0.000002	0.000616	0.003154	0.000005	0.000155	0.000533	0.000000	0.000013	0.000061	0.000000	0.000158	0.000540
(8) - (3)	0.000177	0.008385	0.029612	0.000458	0.002320	0.009164	0.000008	0.000199	0.000758	0.000483	0.002363	0.009201
$(\mathbf{f}^T, \mathbf{u}_g^T)^T$												
$\psi = 0.01$												
(8) - (7)	0.000000	0.000001	0.000004	0.000000	0.000001	0.000014	0.000000	0.000000	0.000009	0.000000	0.000000	0.000012
(8) - (6)	0.000000	0.000011	0.000305	0.000001	0.000006	0.000194	0.000001	0.000005	0.000145	0.000000	0.000003	0.000169
(8) - (5)	0.000004	0.000122	0.002836	0.000011	0.000056	0.000644	0.000013	0.000050	0.000559	0.000006	0.000028	0.000529
(8) - (4)	0.000408	0.001283	0.015645	0.000107	0.000440	0.001515	0.000124	0.000430	0.001415	0.000053	0.000245	0.001234
(8) - (3)	0.004177	0.011778	0.062254	0.000808	0.001862	0.003068	0.000931	0.001962	0.003187	0.000315	0.001167	0.002429
$\psi = 0.05$												
(8) - (7)	0.000000	0.000000	0.000001	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
(8) - (6)	0.000000	0.000004	0.000008	0.000000	0.000001	0.000001	0.000000	0.000001	0.000001	0.000000	0.000000	0.000001
(8) - (5)	0.000005	0.000036	0.000093	0.000000	0.000006	0.000015	0.000000	0.000007	0.000015	0.000000	0.000002	0.000014
(8) - (4)	0.000048	0.000363	0.000920	0.000002	0.000061	0.000149	0.000000	0.000060	0.000149	0.000000	0.000029	0.000127
(8) - (3)	0.000567	0.003547	0.009039	0.000012	0.000392	0.001092	0.000012	0.000360	0.001312	0.000000	0.000592	0.004037
$\psi = 0.25$												
(8) - (7)	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
(8) - (6)	0.000000	0.000000	0.000005	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000001
(8) - (5)	0.000000	0.000001	0.000040	0.000000	0.000001	0.000009	0.000000	0.000000	0.000008	0.000000	0.000001	0.000019
(8) - (4)	0.000000	0.000018	0.000481	0.000000	0.000013	0.000312	0.000000	0.000003	0.000141	0.000000	0.000012	0.000145
(8) - (3)	0.000000	0.000390	0.004412	0.000002	0.000244	0.004239	0.000000	0.000055	0.001695	0.000009	0.000189	0.004155

CHAPTER 8. A NEW MISSING DATA SPECIFICATION FOR FACTOR ANALYTIC MIXED MODELS

Figure 8.1: Model $\mathcal{Z}$ number of iterations until convergence using the new missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \mathbf{u}_g^T)^T$ versus iterations until convergence using the missing data specification $\mathbf{y}_m^{(D)} = (\mathbf{f}^T, \delta^T)^T$. Iterations until convergence for both is based on $\varepsilon = 10^{-5}$. Each data point represents 1 generated data set out of 100 for each of $\psi = 0.01$, $\psi = 0.05$, and $\psi = 0.25$. Grid lines represent the 25th, 50th, and 75th percentiles.



Chapter 9

Conclusion

In this thesis we have advanced the theory underlying the application of REML EM and REML PX EM algorithms to linear mixed models and factor analytic mixed models. A new incomplete data specification has been introduced based on the transformed observed data vector associated with the marginal distribution in the conditional derivation of REML (Verbyla, 1990). This transformed observed data vector has been referred to as $\mathbf{y}_2$. There are two benefits in defining the incomplete data in a REML EM or REML PX EM algorithm as $\mathbf{y}_2$ rather than the observed data vector $\mathbf{y}_o$. Firstly, the new incomplete data specification results in a REML EM or REML PX EM algorithm that is computationally more efficient. Secondly, in our experience the new incomplete data specification will result in a REML EM or REML PX EM algorithm having a faster rate of convergence. We have used Loewner partial ordering to describe the conditions necessary for the new incomplete data specification to result in a faster rate of convergence.

There are two related aspects to the improvement in computational efficiency when specifying the incomplete data as $\mathbf{y}_2$. Firstly, the E-step of a REML EM or REML PX EM algorithm using the new incomplete data specification only requires knowledge of one conditional distribution, namely $\mathbf{u}|\mathbf{y}_2$. In contrast, the E-step of a REML EM and REML PX EM algorithm using $\mathbf{y}_o$ as the incomplete data requires two conditional distributions, namely $\mathbf{u}|\mathbf{y}_2$ and typically $\mathbf{e}|\mathbf{y}_2$. This aspect of the new incomplete data specification may have an application in situations where the E-step is intractable. This is a potential area of future research.

The second aspect to the improvement in computational efficiency relates to the first. REML EM and REML PX EM algorithm parameter updates involve computing the

trace of a matrix. For the new incomplete data specification this matrix is either the $b \times b$ matrix $\mathbf{C}^{ZZ}$ or the $b \times b$ matrix $\mathbf{Z}^T \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ}$ where $\mathbf{U}$ is a projection matrix. When the incomplete data is specified as $\mathbf{y}_o$ this matrix is either $\mathbf{C}^{ZZ}$ or the $n \times n$ matrix $\mathbf{W} \mathbf{C}^{-1} \mathbf{W}^T$ where $\mathbf{W} = (\mathbf{X} \ \mathbf{Z})$. The latter involves the inverse of the entire coefficient matrix associated with Henderson's mixed model equations whereas the former only involves a smaller partition of this matrix. It is computing the trace of $\mathbf{W} \mathbf{C}^{-1} \mathbf{W}^T$ which makes a REML EM or REML PX EM algorithm where the incomplete data is specified as $\mathbf{y}_o$ less computationally efficient than a REML EM or REML PX EM algorithm based on the new incomplete data specification.

Although the REML EM and REML PX EM algorithm for linear mixed models based on the new incomplete data specification is a step forward in terms of computational efficiency further work is required before either of these two algorithms can be easily implemented within existing software packages using Newton-Raphson type algorithms. In the case of the average information algorithm (Gilmour et al., 1995) the only term computed in terms of the inverse of the coefficient matrix associated with Henderson's mixed model equations is $\mathbf{C}^{ZZ}$. The REML EM and REML PX EM algorithm require computing $\mathbf{Z}^T \mathbf{U} \mathbf{Z} \mathbf{C}^{ZZ}$. A potential area of future research would be to study approximations of the trace of this matrix only using terms that form part of the average information algorithm computing strategy.

Factor analytic mixed models are routinely applied to data arising from plant improvement programmes. In statistical software packages using Newton-Raphson type algorithms these models can be extremely difficult to fit. Successive iterations of these algorithms are not guaranteed to increase the residual log-likelihood and parameter updates may be outside their parameter space. REML EM or REML PX EM algorithms for factor analytic mixed models solve these problems. However, they do suffer from slow convergence. Besides specifying the incomplete data for these types of models as $\mathbf{y}_2$ this thesis has addressed the problem of slow convergence by considering model specification, a new parameter expansion, a new missing data specification, and the efficacy of a less stringent stopping rule.

Factor analytic mixed models with $k = 1$ factors and a common specific variance were considered in an example and a simulation study. The recommendation of this thesis is to use a REML PX EM algorithm based on a model specification where the variance of the factor scores is assumed to be an unstructured positive definite matrix rather than the identity matrix, a new missing data specification based on the "dependent" formulation of factor analytic variance models presented in Thomp-

son et al. (2003), a new parameter expansion which is associated with the specific variances, and a less stringent convergence criterion based on $\varepsilon = 10^{-5}$ rather than $\varepsilon = 10^{-8}$. In many instances this results in iterations to convergence being less than 10% of the number of iterations to convergence using an alternative model specification, missing data specification, parameter expansion, and convergence criterion. Future research is required to determine if these results apply to factor analytic mixed models which don't have a common specific variance and with $k > 1$ factors. The research presented in this thesis suggests that a REML PX EM algorithm has the potential to be a practical option when fitting factor analytic mixed models.

Appendix A

Vector and matrix results

A.1 Conditional distribution between two vectors

If $\mathbf{x} = (\mathbf{x}_1^T, \mathbf{x}_2^T)^T$ is distributed as $N_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ where $\boldsymbol{\mu} = (\boldsymbol{\mu}_1^T, \boldsymbol{\mu}_2^T)^T$ and

$$\boldsymbol{\Sigma} = \begin{pmatrix} \boldsymbol{\Sigma}_{11} & \boldsymbol{\Sigma}_{12} \\ \boldsymbol{\Sigma}_{21} & \boldsymbol{\Sigma}_{22} \end{pmatrix}$$

where $\det(\boldsymbol{\Sigma}_{22}) > 0$ then for a given value of $\mathbf{x}_2$

$$\mathbf{x}_1|\mathbf{x}_2 \sim N(\boldsymbol{\mu}_1 + \boldsymbol{\Sigma}_{12}\boldsymbol{\Sigma}_{22}^{-1}(\mathbf{x}_2 - \boldsymbol{\mu}_2), \boldsymbol{\Sigma}_{11} - \boldsymbol{\Sigma}_{12}\boldsymbol{\Sigma}_{22}^{-1}\boldsymbol{\Sigma}_{21})$$

A.2 Covariance between two vectors

If $\mathbf{x} = (\mathbf{x}_1^T, \mathbf{x}_2^T)^T$ is distributed as $N_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ where $\boldsymbol{\mu} = (\boldsymbol{\mu}_1^T, \boldsymbol{\mu}_2^T)^T$ and

$$\boldsymbol{\Sigma} = \begin{pmatrix} \boldsymbol{\Sigma}_{11} & \boldsymbol{\Sigma}_{12} \\ \boldsymbol{\Sigma}_{21} & \boldsymbol{\Sigma}_{22} \end{pmatrix}$$

then we can write

$$\boldsymbol{\Sigma}_{12} = E(\mathbf{x}_1\mathbf{x}_2^T) - E(\mathbf{x}_1)E(\mathbf{x}_2)^T.$$

Rearranging we have

$$\begin{aligned} E(\mathbf{x}_1 \mathbf{x}_2^T) &= E(\mathbf{x}_1) E(\mathbf{x}_2)^T + \boldsymbol{\Sigma}_{12} \\ &= \boldsymbol{\mu}_1 \boldsymbol{\mu}_2^T + \boldsymbol{\Sigma}_{12}. \end{aligned}$$

A.3 Matrix identities

A.3.1 Trace of a matrix

The following is taken from Mardia et al. (1979). For a $n \times n$ matrix $\mathbf{A}$ with elements a_{ij} the trace of $\mathbf{A}$ is

$$\text{tr}(\mathbf{A}) = \sum_{i=1}^n a_{ii} = \sum_{i=1}^n \lambda_i$$

where λ_i is the i -th eigenvalue of $\mathbf{A}$.

A.3.2 Quadratic forms

The following result is taken from Searle et al. (1992). For $\mathbf{y} \sim N(\boldsymbol{\mu}, \mathbf{V})$ and symmetric matrix $\mathbf{A}$,

$$E(\mathbf{y}^T \mathbf{A} \mathbf{y}) = \boldsymbol{\mu}^T \mathbf{A} \boldsymbol{\mu} + \text{tr}(\mathbf{A} \mathbf{V})$$

A.3.3 Determinants

The following results are taken from Harville (1997).

Matrix multiplied by a scalar

For any $n \times n$ matrix $\mathbf{A}$ and scalar k

$$\det(k\mathbf{A}) = k^n \det(\mathbf{A})$$

Invertible matrix

If $\mathbf{A}$ is non-singular then

$$\det(\mathbf{A}^{-1}) = \frac{1}{\det(\mathbf{A})}$$

Matrix multiplication

For $n \times n$ matrices $\mathbf{A}$ and $\mathbf{B}$

$$\det(\mathbf{AB}) = \det(\mathbf{A}) \det(\mathbf{B})$$

Partitioned matrix

For a non-singular matrix $\mathbf{A}$ partitioned as

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_{11} & \mathbf{A}_{12} \\ \mathbf{A}_{21} & \mathbf{A}_{22} \end{pmatrix}$$

where $\mathbf{A}_{11}$ is a $m \times m$ matrix, $\mathbf{A}_{12}$ is a $m \times n$ matrix, $\mathbf{A}_{21}$ is a $n \times m$ matrix and $\mathbf{A}_{22}$ is a $n \times n$ matrix. The determinant of $\mathbf{A}$ is

$$\det(\mathbf{A}) = \det(\mathbf{A}_{11}) \det(\mathbf{A}_{22} - \mathbf{A}_{21} \mathbf{A}_{11}^{-1} \mathbf{A}_{12})$$

A.3.4 Inverses of matrices of the form $\mathbf{A} + \mathbf{B}$

The following result is taken from Harville (1997). Let $\mathbf{A}$ and $\mathbf{B}$ be $n \times n$ matrices and $\mathbf{A} + \mathbf{B}$ be nonsingular then

$$\begin{aligned} (\mathbf{A} + \mathbf{B})^{-1} \mathbf{A} &= \mathbf{I}_n - (\mathbf{A} + \mathbf{B})^{-1} \mathbf{B} \\ \mathbf{A}(\mathbf{A} + \mathbf{B})^{-1} &= \mathbf{I}_n - \mathbf{B}(\mathbf{A} + \mathbf{B})^{-1} \end{aligned}$$

A.3.5 Inverses of matrices of the form $A + BCD$

The following result is taken from Harville (1997). Let A be a nonsingular $n \times n$ matrix, B a $n \times m$ matrix, C a nonsingular $m \times m$ matrix, and D a $m \times n$ matrix then

$$(A + BCD)^{-1} = A^{-1} - A^{-1}B(C^{-1} + DA^{-1}B)^{-1}DA^{-1}$$

A.3.6 Inverses of partitioned matrices

The following results are taken from Harville (1997). Define a partitioned non-singular matrix A as

$$A = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix}$$

where A_{11} is a $m \times m$ matrix, A_{12} is a $m \times n$ matrix, A_{21} is a $n \times m$ matrix and A_{22} is a $n \times n$ matrix. Define the the non-singular matrix A^{-1} as

$$A^{-1} = \begin{pmatrix} A^{11} & A^{12} \\ A^{21} & A^{22} \end{pmatrix}$$

where $A^{11}, A^{12}, A^{21}, A^{22}$ have the same dimensions as $A_{11}, A_{12}, A_{21}, A_{22}$ respectively. If A_{11} is non-singular then A^{11} is non-singular and

$$\begin{aligned} A^{11} &= (A_{11} - A_{12}A_{22}^{-1}A_{21})^{-1} \\ A^{12} &= -(A_{11} - A_{12}A_{22}^{-1}A_{21})^{-1}A_{12}A_{22}^{-1} \\ A^{21} &= -A_{22}^{-1}A_{21}(A_{11} - A_{12}A_{22}^{-1}A_{21})^{-1} \\ A^{22} &= A_{22}^{-1} + A_{22}^{-1}A_{21}(A_{11} - A_{12}A_{22}^{-1}A_{21})^{-1}A_{12}A_{22}^{-1} \end{aligned}$$

$$\text{A.3.7} \quad \mathbf{G}\mathbf{Z}^T(\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} = (\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1}$$

Using Result A.3.5 and Result A.3.4 we have

$$\begin{aligned} (\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} &= \mathbf{R}^{-1} - \mathbf{R}^{-1}\mathbf{Z}(\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1} \\ \mathbf{R}(\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} &= \mathbf{I}_n - \mathbf{Z}(\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1} \\ \mathbf{I}_n - \mathbf{Z}\mathbf{G}\mathbf{Z}^T(\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} &= \mathbf{I}_n - \mathbf{Z}(\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1} \\ \mathbf{Z}\mathbf{G}\mathbf{Z}^T(\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} &= \mathbf{Z}(\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1} \\ \mathbf{G}\mathbf{Z}^T(\mathbf{R} + \mathbf{Z}\mathbf{G}\mathbf{Z}^T)^{-1} &= (\mathbf{G}^{-1} + \mathbf{Z}^T\mathbf{R}^{-1}\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{R}^{-1} \end{aligned}$$

as required.

$$\text{A.3.8} \quad \mathbf{P} = \mathbf{L}_2(\mathbf{L}_2^T\mathbf{H}\mathbf{L}_2)^{-1}\mathbf{L}_2^T = \mathbf{H}^{-1} - \mathbf{H}^{-1}\mathbf{X}(\mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-1}$$

If $\mathbf{L}_2^T\mathbf{X} = \mathbf{0}$ and $\mathbf{H}$ is symmetric positive definite then it is possible to define a projection matrix $\mathbf{P}$ where

$$\begin{aligned} \mathbf{P} &= \mathbf{L}_2(\mathbf{L}_2^T\mathbf{H}\mathbf{L}_2)^{-1}\mathbf{L}_2^T \\ &= \mathbf{H}^{-1} - \mathbf{H}^{-1}\mathbf{X}(\mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-1} \end{aligned}$$

Part of the proof is presented below and is taken from Searle et al. (1992). We start by noting that $\mathbf{L}_2(\mathbf{L}_2^T\mathbf{L}_2)^{-1}\mathbf{L}_2^T$ and $\mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T$ are symmetric and idempotent and that it is possible to show that

$$\mathbf{I}_n - \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T = \mathbf{L}_2(\mathbf{L}_2^T\mathbf{L}_2)^{-1}\mathbf{L}_2^T$$

Since $\mathbf{H}$ is symmetric positive definite a matrix $\mathbf{H}^{\frac{1}{2}}$ exists such that $\mathbf{H} = (\mathbf{H}^{\frac{1}{2}})^2$ and we can write

$$\mathbf{I}_n - \mathbf{H}^{-\frac{1}{2}}\mathbf{X}(\mathbf{X}^T\mathbf{H}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\mathbf{H}^{-\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}\mathbf{L}_2(\mathbf{L}_2^T\mathbf{H}\mathbf{L}_2)^{-1}\mathbf{L}_2^T\mathbf{H}^{\frac{1}{2}}$$

Pre and post-multiplying both sides by $\mathbf{H}^{-\frac{1}{2}}$ gives

$$P = H^{-1} - H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1} = L_2(L_2^T H L_2)^{-1}L_2^T$$

Since R and Σ are symmetric positive definite it follows that we can also define

$$S = R^{-1} - R^{-1}X(X^T R^{-1}X)^{-1}X^T R^{-1} = L_2(L_2^T R L_2)^{-1}L_2^T$$

$$U = \Sigma^{-1} - \Sigma^{-1}X(X^T \Sigma^{-1}X)^{-1}X^T \Sigma^{-1} = L_2(L_2^T \Sigma L_2)^{-1}L_2^T$$

A.3.9 $P = R^{-1} - R^{-1}WC^{-1}W^T R^{-1}$

From Result A.3.8 we have

$$P = H^{-1} - H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1}$$

Now consider

$$\begin{aligned} H^{-1} &= R^{-1} - R^{-1}Z(G^{-1} + Z^T R^{-1}Z)^{-1}Z^T R^{-1} \\ &= R^{-1} - R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1} \end{aligned}$$

and

$$\begin{aligned} &H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1} \\ &= (R^{-1} - R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1})XC^{XX}X^T(R^{-1} - R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1}) \\ &= R^{-1}XC^{XX}X^T R^{-1} - R^{-1}XC^{XX}X^T R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1} \\ &\quad - R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1}XC^{XX}X^T R^{-1} \\ &\quad + R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1}XC^{XX}X^T R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1} \\ &= R^{-1}(XC^{XX}X^T - XC^{XX}C_{XZ}C_{ZZ}^{-1}Z^T - ZC_{ZZ}^{-1}C_{ZX}C^{XX}X^T \\ &\quad + ZC_{ZZ}^{-1}C_{ZX}C^{XX}C_{XZ}C_{ZZ}^{-1}Z^T)R^{-1} \end{aligned}$$

Therefore,

$$\begin{aligned}
 P &= H^{-1} - H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1} \\
 &= R^{-1} - R^{-1}ZC_{ZZ}^{-1}Z^T R^{-1} - R^{-1}(XC^{XX}X^T - XC^{XX}C_{XZ}C_{ZZ}^{-1}Z^T \\
 &\quad - ZC_{ZZ}^{-1}C_{ZX}C^{XX}X^T + ZC_{ZZ}^{-1}C_{ZX}C^{XX}C_{XZ}C_{ZZ}^{-1}Z^T)R^{-1} \\
 &= R^{-1} - R^{-1}(ZC_{ZZ}^{-1}Z^T + XC^{XX}X^T - XC^{XX}C_{XZ}C_{ZZ}^{-1}Z^T \\
 &\quad - ZC_{ZZ}^{-1}C_{ZX}C^{XX}X^T + ZC_{ZZ}^{-1}C_{ZX}C^{XX}C_{XZ}C_{ZZ}^{-1}Z^T)R^{-1} \\
 &= R^{-1} - R^{-1}WC^{-1}W^T R^{-1}
 \end{aligned}$$

where $W = (X \ Z)$.

A.3.10 $P = S - SZC^{ZZ}Z^T S$

Begin by noting that P can be written as

$$\begin{aligned}
 &L_2(L_2^T H L_2)^{-1}L_2^T \\
 &= L_2(L_2^T R L_2 + L_2 Z G Z^T L_2)^{-1}L_2^T \\
 &= L_2\{(L_2^T R L_2)^{-1} \\
 &\quad - (L_2^T R L_2)^{-1}L_2^T Z\{G^{-1} + Z^T L_2(L_2^T R L_2)^{-1}L_2^T Z\}^{-1}Z^T L_2(L_2^T R L_2)^{-1}\}L_2^T \\
 &\quad \text{using Result A.3.5} \\
 &= S - SZ(G^{-1} + Z^T S Z)^{-1}Z^T S \\
 &= S - SZC^{ZZ}Z^T S
 \end{aligned}$$

If we define $L_2^T y_o = y_2$ then we can use the result presented above to re-express $G - GZ^T P Z G$ as

$$\begin{aligned}
G - GZ^T P Z G &= G - G \{ Z^T S Z - Z^T S Z (G^{-1} + Z^T S Z)^{-1} Z^T S Z \} G \\
&= G - G \{ I_b - Z^T S Z (G^{-1} + Z^T S Z)^{-1} \} Z^T S Z G \\
&= G - G [I_b - \{ I_b - G^{-1} (G^{-1} + Z^T S Z)^{-1} \}] Z^T S Z G \\
&\quad \text{using Result A.3.4} \\
&= G - (G^{-1} + Z^T S Z)^{-1} Z^T S Z G \\
&= \{ I_b - (G^{-1} + Z^T S Z)^{-1} Z^T S Z \} G \\
&= [I_b - \{ I_b - (G^{-1} + Z^T S Z)^{-1} G^{-1} \}] G \\
&\quad \text{using Result A.3.4} \\
&= (G^{-1} + Z^T S Z)^{-1} \\
&= C^{ZZ}
\end{aligned}$$

A.3.11 $PHP = P$

Note that $P = L_2(L_2^T H L_2)^{-1} L_2^T$. Therefore,

$$\begin{aligned}
PHP &= L_2(L_2^T H L_2)^{-1} L_2^T H L_2 (L_2^T H L_2)^{-1} L_2^T \\
&= L_2(L_2^T H L_2)^{-1} L_2^T \\
&= P
\end{aligned}$$

as required. Similarly $SRS = S$ and $U\Sigma U = U$.

A.3.12 $PHS = S$

Noting that

$$\begin{aligned}
P &= L_2(L_2^T H L_2)^{-1} L_2^T \\
S &= L_2(L_2^T R L_2)^{-1} L_2^T
\end{aligned}$$

then

$$\begin{aligned}
 PHS &= L_2(L_2^T H L_2)^{-1} L_2^T H L_2 (L_2^T R L_2)^{-1} L_2^T \\
 &= L_2(L_2^T R L_2)^{-1} L_2^T \\
 &= S
 \end{aligned}$$

as required.

A.3.13 $PRS = P$

Noting that

$$\begin{aligned}
 P &= L_2(L_2^T H L_2)^{-1} L_2^T \\
 S &= L_2(L_2^T R L_2)^{-1} L_2^T
 \end{aligned}$$

then

$$\begin{aligned}
 PRS &= L_2(L_2^T H L_2)^{-1} L_2^T R L_2 (L_2^T R L_2)^{-1} L_2^T \\
 &= L_2(L_2^T H L_2)^{-1} L_2^T \\
 &= P
 \end{aligned}$$

as required.

A.3.14 $\text{tr}(HP) = n - t$

If we define $P = H^{-1} - H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1}$ then

$$\begin{aligned}
 \text{tr}(HP) &= \text{tr}[H\{H^{-1} - H^{-1}X(X^T H^{-1}X)^{-1}X^T H^{-1}\}] \\
 &= \text{tr}\{I_n - X(X^T H^{-1}X)^{-1}X^T H^{-1}\} \\
 &= \text{tr}(I_n) - \text{tr}(X(X^T H^{-1}X)^{-1}X^T H^{-1}) \\
 &= n - \text{tr}\{(X^T H^{-1}X)^{-1}X^T H^{-1}X\} \\
 &= n - \text{tr}(I_t) \\
 &= n - t
 \end{aligned}$$

as required. Similarly $\text{tr}(\mathbf{R}\mathbf{S}) = \text{tr}(\mathbf{\Sigma}\mathbf{U}) = n - t$.

A.4 Matrix differentiation results

A.4.1 Derivatives of determinants

The following result is taken from Harville (1997). If $\mathbf{A}$ is a $p \times p$ matrix where $\mathbf{A} = \mathbf{A}(x) = \{b_{ij}\}$ is a non-singular matrix whose elements are functions of x and is continuously differentiable at every x then

$$\frac{\partial \log\{\det(\mathbf{A})\}}{\partial x_j} = \text{tr} \left(\mathbf{A}^{-1} \frac{\partial \mathbf{A}}{\partial x_j} \right)$$

A.4.2 Derivatives of inverses

The following result is taken from Harville (1997). If $\mathbf{A}$ is a $p \times p$ matrix where $\mathbf{A} = \mathbf{A}(x) = \{b_{ij}\}$ is a non-singular matrix whose elements are functions of x and is continuously differentiable at every x then

$$\frac{\partial \mathbf{A}^{-1}}{\partial x_j} = -\mathbf{A}^{-1} \frac{\partial \mathbf{A}}{\partial x_j} \mathbf{A}^{-1}$$

A.4.3 Derivatives of $\mathbf{P}$

Define a vector of variance parameters $\boldsymbol{\kappa} = (\sigma^2, \boldsymbol{\phi}^T, \boldsymbol{\gamma}^T)^T$ and consider the partial derivative of $\mathbf{P} = \mathbf{L}_2(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T$ (Result A.3.8) with respect to κ_i , i.e.,

$$\begin{aligned} \frac{\partial \mathbf{P}}{\partial \kappa_i} &= \frac{\partial}{\partial \kappa_i} \left\{ \mathbf{L}_2(\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \right\} \\ &= \mathbf{L}_2 \frac{\partial (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1}}{\partial \kappa_i} \mathbf{L}_2^T \\ &= -\mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \frac{\partial \mathbf{H}}{\partial \kappa_i} \mathbf{L}_2 (\mathbf{L}_2^T \mathbf{H} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \quad \text{using Result A.4.2} \\ &= -\mathbf{P} \frac{\partial \mathbf{H}}{\partial \kappa_i} \mathbf{P} \end{aligned}$$

Likewise,

$$\begin{aligned}
 \frac{\partial \mathbf{U}}{\partial \phi_{ij}} &= \frac{\partial}{\partial \phi_{ij}} \left\{ \mathbf{L}_2 (\mathbf{L}_2^T \boldsymbol{\Sigma} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \right\} \\
 &= \mathbf{L}_2 \frac{\partial (\mathbf{L}_2^T \boldsymbol{\Sigma} \mathbf{L}_2)^{-1}}{\partial \phi_{ij}} \mathbf{L}_2^T \\
 &= -\mathbf{L}_2 (\mathbf{L}_2^T \boldsymbol{\Sigma} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \frac{\partial \boldsymbol{\Sigma}}{\partial \phi_{ij}} \mathbf{L}_2 (\mathbf{L}_2^T \boldsymbol{\Sigma} \mathbf{L}_2)^{-1} \mathbf{L}_2^T \quad \text{using Result A.4.2} \\
 &= -\mathbf{U} \frac{\partial \boldsymbol{\Sigma}}{\partial \phi_{ij}} \mathbf{U}
 \end{aligned}$$

Appendix B

Derivation of identity in Oakes (1999)

A key result in Oakes (1999) is the identity

$$\begin{aligned} & E \left[\frac{\partial^2 \log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\ &= -E \left(\frac{\partial}{\partial \boldsymbol{\theta}^{(w)}} \left[\log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right) \end{aligned}$$

Noting that $\boldsymbol{\theta} = (\theta_1, \dots, \theta_d)^T$, defining $i = 1, \dots, d$ and $j = 1, \dots, d$ we can write

$$\begin{aligned} \frac{\partial^2 \log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \theta_i^{(w)} \partial \theta_j^{(w)}} &= \frac{\partial}{\partial \theta_i^{(w)}} \left[\frac{\partial \log \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \theta_j^{(w)}} \right] \\ &= \frac{\partial}{\partial \theta_i^{(w)}} \left[\frac{\frac{\partial}{\partial \theta_j^{(w)}} \{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})} \right] \end{aligned}$$

Using the quotient rule we have

$$\begin{aligned}
 & \frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \theta_i^{(w)} \partial \theta_j^{(w)}} \\
 &= \frac{\frac{\partial^2}{\partial \theta_i^{(w)} \partial \theta_j^{(w)}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\} k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) - \frac{\partial}{\partial \theta_j^{(w)}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\} \frac{\partial}{\partial \theta_i^{(w)}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})^2} \\
 &= \frac{\frac{\partial^2}{\partial \theta_i^{(w)} \partial \theta_j^{(w)}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})} - \frac{\partial}{\partial \theta_j^{(w)}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \frac{\partial}{\partial \theta_i^{(w)}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right]
 \end{aligned}$$

Which can be written using matrix notation as

$$\begin{aligned}
 & \frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \\
 &= \frac{\frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})} - \frac{\partial}{\partial \boldsymbol{\theta}^{(w)}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right]
 \end{aligned}$$

Taking expectations on both sides over $k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})$ we have

$$\begin{aligned}
 & E \left[\frac{\partial^2 \log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\}}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
 &= E \left[\frac{\frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \left\{ k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
 &\quad - E \left(\frac{\partial}{\partial \boldsymbol{\theta}^{(w)}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \frac{\partial}{\partial \boldsymbol{\theta}^{(w)T}} \left[\log\{k(\mathbf{y}_m|\mathbf{y}_o; \boldsymbol{\theta}^{(w)})\} \right] \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right)
 \end{aligned}$$

The required identity is obtained if the first term on the right hand side is zero. We have

$$\begin{aligned}
& E \left[\frac{\frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \left\{ k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})} \middle| \mathbf{y}_o; \boldsymbol{\theta}^{(w)} \right] \\
&= \int_{S(\mathbf{y}_m)} \frac{\frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \left\{ k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) \right\}}{k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)})} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \\
&= \frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \left\{ \int_{S(\mathbf{y}_m)} k(\mathbf{y}_m | \mathbf{y}_o; \boldsymbol{\theta}^{(w)}) d\mathbf{y}_m \right\} \\
&= \frac{\partial^2}{\partial \boldsymbol{\theta}^{(w)} \partial \boldsymbol{\theta}^{(w)T}} \begin{pmatrix} 1 \end{pmatrix} \\
&= \mathbf{0}
\end{aligned}$$

as required.

Appendix C

Loewner partial ordering of symmetric positive semidefinite matrices

Loewner (sometimes referred to as Löwner) partial ordering of matrices is analagous to ordering numbers on the real number line using the inequalities $>$, $\geq$, $<$, and $\leq$. Loewner partial ordering has its' mathematical foundations in convex geometry. Much of what follows has been taken from Horn and Johnson (1990)

We begin by defining $\mathbb{S}_+^n$ as the convex cone containing the set of all positive semidefinite symmetric $n \times n$ matrices and the interior of $\mathbb{S}_+^n$, denoted $\mathbb{S}_{++}^n$, as the set of all positive definite symmetric $n \times n$ matrices. For $\mathbf{A} \in \mathbb{S}_+^n$ we can write

$$\mathbf{A} \succeq 0 \iff \mathbf{x}^T \mathbf{A} \mathbf{x} \geq 0 \quad \forall \mathbf{x} \in \mathbb{R}^n$$

where $\mathbb{R}^n$ refers to Euclidean n -dimensional real vector space and the symbols $\in$, $\forall$ and $\iff$ can be read as “in”, “for all” and “material equivalence” respectively. Later we introduce the symbol $\implies$ which can be read as “material implication”. For $\mathbf{A} \in \mathbb{S}_{++}^n$ we can write

$$\mathbf{A} \succ 0 \iff \mathbf{x}^T \mathbf{A} \mathbf{x} > 0 \quad \forall \mathbf{x} \in \mathbb{R}^n$$

If $\mathbf{A} - \mathbf{B} \succeq 0$ then $\mathbf{A} \succeq \mathbf{B}$ where $\mathbf{A} \succeq \mathbf{B}$ defines a so-called Loewner partial ordering (see for example Jarre (2000) or Dattorro (2005) Appendix A). The corollaries we

are particularly interested in are (Horn and Johnson (1990) Chapter 7)

Corollary C.1. *If $\mathbf{A}, \mathbf{B} \in \mathbb{S}_{++}^n$, then*

- (a) $\mathbf{A} \succeq \mathbf{B} \iff \mathbf{A}^{-1} \preceq \mathbf{B}^{-1}$;
- (b) $\mathbf{A} \succeq \mathbf{B} \implies \det(\mathbf{A}) \geq \det(\mathbf{B})$; and
- (c) $\mathbf{A} \succeq \mathbf{B} \implies \lambda_k(\mathbf{A}) \geq \lambda_k(\mathbf{B})$ for all $k = 1, 2, \dots, n$ if the respective eigenvalues of $\mathbf{A}$ and $\mathbf{B}$ are arranged in the same (increasing or decreasing) order.

Note that Corollary C.1 (b) and (c) imply that for $\mathbf{A}, \mathbf{B} \in \mathbb{S}_{++}^n$ then

$$\det(\mathbf{A}) \geq \det(\mathbf{B}) \iff \lambda_k(\mathbf{A}) \geq \lambda_k(\mathbf{B}) \quad \text{for all } k = 1, 2, \dots, n$$

According to Dattorro (2005) the requirement that $\mathbf{A}$ and $\mathbf{B}$ are positive definite is only required for Corollary C.1 (a) whereas $\mathbf{A}, \mathbf{B} \in \mathbb{S}_+^N$ is sufficient for Corollary C.1 (b) and (c).

Bibliography

- Butler, D., Cullis, B., Gilmour, A., and Gogel, B. (2007). *ASReml-R reference manual*. ISSN 0812-0005.
- Callanan, T. P. and Harville, D. A. (1991). Some new algorithms for computing restricted maximum likelihood estimates of variance components. *Journal of Statistical Computation and Simulation*, 38:239–259.
- Cullis, B., Smith, A., and Thompson, R. (2004). *Perspectives of ANOVA, REML, and a general linear mixed model*. Methods and Models in Statistics - In honour of Professor John Nelder, FRS, Editors N. Adams, M. Crowder, D. J. Hand and D. Stephens. Imperial College Press, London.
- Dattorro, J. (2005). *Convex Optimization and Euclidean Distance Geometry*. Meboo Publishing. v2009.09.22.
- Dempster, A., Laird, N., and Rubin, D. (1977). Maximum Likelihood from Incomplete Data via the EM Algorithm. *Journal of the Royal Statistical Society, Series B*, 39:1–38.
- Foulley, J.-L., Jaffrezic, F., and Robert-Cranie, C. (2000). EM-REML estimation of covariance parameters in Gaussian mixed models for longitudinal data analysis. *Genetics Selection Evolution*, 32:129–141.
- Foulley, J.-L. and Van Dyk, D. (2000). The PX-EM algorithm for fast stable fitting of Henderson’s mixed model. *Genetics Selection Evolution*, 32:129–141.
- Gilmour, A., Thompson, R., and Cullis, B. (1995). AI, an efficient algorithm for REML estimation in linear mixed models. *Biometrics*, 51:1440–1450.
- Harville, D. (1997). *Matrix algebra from a statistician’s perspective*. Springer-Verlag, New York, Inc.

BIBLIOGRAPHY

- Henderson, C. (1953). Estimation of variance and covariance components. *Biometrics*, 9(2):226–252.
- Henderson, C., Kempthorne, O., Searle, S., and von Krosigk, C. (1959). The estimation of environmental and genetic trends from records subject to culling. *Biometrics*, 15(2):192–218.
- Horn, R. and Johnson, C. (1990). *Matrix analysis*. Cambridge University Press.
- Jarre, F. (2000). *Convex analysis on symmetric matrices*, chapter 2. The geometry of semidefinite programming, Editors H. Wolkowicz, R. Saigal and L. Vandenberghe. Kluwer.
- Kalbfleisch, J. and Sprott, A. (1970). Application of likelihood methods to models involving large numbers of parameters. *Journal of the Royal Statistical Society, Series B*, 32(2):175–208.
- Knight, E. (2008). *Improved iterative schemes for REML estimation of variance parameters in linear mixed models*. PhD thesis, The University of Adelaide, School of Agriculture, Food and Wine.
- Laird, N. and Ware, J. (1982). Random-Effects Models for Longitudinal Data. *Biometrics*, 38:963–974.
- Liu, C., Rubin, D., and Wu, Y. (1998). Parameter expansion to accelerated EM: The PX-EM algorithm. *Biometrika*, 85:755–770.
- Louis, T. (1982). Finding the observed information matrix when using the EM algorithm. *Journal of the Royal Statistical Society, B*, 44:226–233.
- Mardia, K., Kent, J., and Bibby, J. (1979). *Multivariate analysis*. Academic Press (London) Limited.
- McLachlan, G. and Krishnan, K. (1997). *The EM Algorithm and Extensions*. John-Wiley and Sons Inc.
- Oakes, D. (1999). Direct calculation of the information matrix via the EM algorithm. *Journal of the Royal Statistical Society, B*, 61:479–482.
- Orchard, T. and Woodbury, M. (1972). A missing information principle: theory and applications. In *Proceedings of the 6th Berkely Symposium on Mathematical Statistics and Probability*, volume 1, pages 697–715, Berkeley, California. University of California Press.

- Patterson, H. and Thompson, R. (1971). Recovery of inter-block information when block sizes are unequal. *Biometrika*, 58:545–554.
- R Development Core Team (2011). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. ISBN 3-900051-07-0.
- Searle, S., Casella, G., and McCulloch, C. (1992). *Variance Components*. John-Wiley and Sons Inc.
- Smith, A., Cullis, B., and Thompson, R. (2001). Analyzing variety by environment data using multiplicative mixed models and adjustment for spatial field trend. *Biometrics*, 57(4):1138–1147.
- Thompson, R., Cullis, B., Smith, A., and Gilmour, A. (2003). A sparse implementation of the average information algorithm for factor analytic and reduced rank variance models. *Australian and New Zealand Journal of Statistics*, 45:445–459.
- Verbyla, A. (1990). A conditional derivation of residual maximum likelihood. *Australian Journal of Statistics*, 32(2):227–230.
- Wald, A. (1949). Note on the consistency of the maximum likelihood estimate. *Annals of Mathematical Statistics*, 20:595.