

THESIS FOR THE DEGREE OF LICENTIATE OF PHILOSOPHY

Multivariate Aspects of Phylogenetic Comparative Methods

KRZYSZTOF BARTOSZEK

CHALMERS |  **UNIVERSITY OF GOTHENBURG**

Division of Mathematical Statistics
Department of Mathematical Sciences
CHALMERS UNIVERSITY OF TECHNOLOGY
AND UNIVERSITY OF GOTHENBURG
Göteborg, Sweden 2011

Multivariate Aspects of Phylogenetic Comparative Methods
Krzysztof Bartoszek

Copyright © Krzysztof Bartoszek, 2011.

ISSN 1652-9715

Report 2011:21

Department of Mathematical Sciences
Division of Mathematical Statistics
Chalmers University of Technology
and University of Gothenburg
SE-412 96 GÖTEBORG, Sweden
Phone: +46 (0)31-772 10 00

Author e-mail: krzbar@chalmers.se

Typeset with L^AT_EX.
Department of Mathematical Sciences
Printed in Göteborg, Sweden 2011

Multivariate Aspects of Phylogenetic Comparative Methods

Krzysztof Bartoszek

Abstract

This thesis concerns multivariate phylogenetic comparative methods. We investigate two aspects of them. The first is the bias caused by measurement error in regression studies of comparative data. We calculate the formula for the bias and show how to correct for it. We also study whether it is always advantageous to correct for the bias as correction can increase the mean square error of the estimate. We propose a criterion, which depends on the observed data, that indicates whether it is beneficial to correct or not. Accompanying the results is an R program that offers the bias correction tool.

The second topic is a multivariate model for trait evolution which is based on an Ornstein–Uhlenbeck type stochastic process, often used for studying trait adaptation, co–evolution, allometry or trade–offs. Alongside the description of the model and presentation of its most important features we present an R program estimating the model’s parameters. To the best of our knowledge our program is the first program that allows for nearly all combinations of key model parameters providing the biologist with a flexible tool for studying multiple interacting traits in the Ornstein–Uhlenbeck framework. There are numerous packages available that include the Ornstein–Uhlenbeck process but their multivariate capabilities seem limited.

Keywords: General Linear Model, Ornstein–Uhlenbeck process, Multivariate phylogenetic comparative method, Evolutionary model, Adaptation, Optimality, Measurement error, Regression, Adaptation, Major-axis regression, Reduced major-axis regression, Structural equation, Allometry, Phylogenetic inertia.

Acknowledgment

I would like to thank Petter Mostad, Olle Nerman and Serik Sagitov for their support in writing this thesis and Patrik Albin, and Thomas Ericsson for many helpful comments.

I would also like to thank Staffan Andersson, Mohammad Asadzadeh, Jakob Augustin, Wojciech Bartoszek, Joachim Domsta, Peter Gennemark, Thomas F. Hansen, Michał Krzemiński, Pietro Lió, Jason Pienaar, Maria Prager, Małgorzata Pułka, Jarosław Skokowski, Anil Sorathiya, Franciska Steinhoff, Anna Stokowska and Kjetil Voje for fruitful collaborations.

I finally thank my colleagues at the Department of Mathematical Sciences for a wonderful working environment and my parents for their constant support.

Krzysztof Bartoszek
Göteborg, April 23, 2012

List of Papers

The licenciate thesis is based on the following papers.

- I. Hansen T. and **Bartoszek K.** (2011). Interpreting the evolutionary regression: the interplay between observational and biological errors in phylogenetic comparative studies. *Systematic Biology in press*
- II. **Bartoszek K.**, Pienaar J., Mostad P., Andersson S., Hansen T. (2011). A comparative method for studying multivariate adaptation. *Working paper*

Contents

1	Introduction	3
2	The general regression framework	6
3	Presentation of Paper I	9
3.1	The measurement error bias	9
3.2	Mean square error analysis	11
4	Presentation of Paper II	13
4.1	Brownian Motion	13
4.2	Ornstein–Uhlenbeck process	14
4.3	The mvSLOUCH package	15
4.4	Example data analysis	16
5	Future work	18
6	Supplementary material A: Bias caused by measurement error; to correct or not to correct (Paper I)	20
6.1	Independent observations of predictors	20
6.2	Independent predictors	21
6.3	Single predictor	22
6.4	Mean square error (MSE) analysis	23
6.4.1	Mean square error for different GLS estimators . .	24
6.4.2	Mean square error for a single predictor	25
6.5	Effect on intercept, fixed effects	28
6.5.1	Fixed effects, random effects measured with and without error	30
6.5.2	Single predictor	31
6.5.3	Mean square error analysis	33
6.6	Unconditional bias of GLS estimator with measurement error in predictors	34
7	Supplementary material B: Stochastic differential equations for the Ornstein–Uhlenbeck model (Paper II)	38
7.1	Convergence of stochastic process	38
7.2	Matrix exponential and related integral	39
7.3	Ornstein–Uhlenbeck model	40
7.3.1	Asymptotic properties of unitrait Ornstein–Uhlenbeck model	41
7.3.2	Asymptotic properties of multitrait Ornstein–Uhlenbeck model	42
7.4	Ornstein–Uhlenbeck Brownian Motion model	42
7.5	Multivariate Ornstein–Uhlenbeck Brownian Motion model	43
7.5.1	Stationary regime of the mvOUBM model	46

7.6	Relations between the models	47
8	Supplementary material C: Matrix parametrizations used in mvSLOUCH package (Paper II)	49
8.1	Real eigenvalues case	49
8.2	Complex eigenvalues case	50
8.3	Matrix parametrization implemented in package	51

1 Introduction

In evolutionary biology one of the fundamental concerns is how different inherited traits depend on each other and react to the changing environment. The natural approach to answer questions related to these problems is to record trait measurements and environmental conditions, and look how they relate to each other, *e.g.* via a regression analysis.

There is however a problem that is not taken into account with just this approach. Namely, the sample is not independent as the different species are related by a common evolutionary history. We assume throughout this thesis that the underlying phylogeny is fully known. This means that we know the topology and branch lengths of the phylogenetic tree describing the common evolutionary history of the species.

The immediate consequence of this common evolutionary history is that closely related species will have similar trait values. As it is very probable that closely related species will be living in similar environments we could observe also a false dependency of traits and environment. One needs to correct for the this common phylogenetic history to be able to distinguish between effects stemming from similarity and evolutionary effects. This is of particular importance when studying the adaptation of traits towards environmental niches. Are species living in a similar environment similar because they have adapted to it or is it because they are closely related?

In statistical terms forgetting about the inter-species dependency means that we would be analyzing data under a wrong model. Such an analysis from the perspective of the correct model could lead to biased estimates and misleading parameter interpretations. It would certainly result in wrong confidence intervals and p-values.

The thesis presented here does not focus on the biological motivations for comparative analysis but on the mathematics and computations involved in estimating parameters of a multivariate model for continuous comparative data. Biological background for the subject can be found in *e.g.* [27] which also discusses methodologies for discrete data and simple continuous settings. Evolution of discrete characters is also discussed in [37] or [39]. Another approach to studying comparative data is using general estimating equations, see *e.g.* [21], [18], [36] or [41].

One of the underlying assumptions is that the different time scales match. This means that the speed of losing the ancestral effects is not orders of magnitude faster than species diversification. If it would be so then we would essentially have an independent sample with no need to take the phylogeny into account. This problem is considered in *e.g.* [14] or [38]. The book [40] is about various R [45] packages for analyzing data connected to phylogenies including comparative data.

This thesis is based on two papers. With each of the papers is an R program related to it.

Paper I which follows this thesis concerns the problem of measurement error (or observational variance, as most comparative studies take average values from a number of individuals) in phylogenetic comparative analysis studies. The problem of measurement error is widely addressed in the literature (see *e.g.* [9] or [15]) and in the multivariate case well studied (*e.g.* [19]). However in the general case of dependent errors in the predictor variables or dependently evolving predictor variables we didn't come across any results in the literature. In Paper I we consider general covariance matrices $\mathbf{V}_d$ and $\mathbf{V}_u$ relating the predictors and their measurement errors respectively between observations. Also we mention that if one assumes some structure on these matrices then substantial simplifications can be made in the bias formula. In chapter 6 of this thesis we consider different covariance structures resulting in the bias formulas (3.3), (3.5), (6.1), (6.2).

We also study in Paper I whether it is always advantageous to correct for the bias, as correcting can increase the mean square error of the estimate. We propose a criterion, which depends on the observed data, that indicates whether it is beneficial to correct or not. In chapter 6.4 of the thesis we look in more detail at conditions in the single predictor case for when it is better to correct, see equation (6.5). We also derive similar conditions for the multipredictor case. We study how the situation changes when one introduces fixed effects into the linear model, see equations (6.7) and (6.8). Lastly we see what can be said about the unconditional (of the observed with error design matrix) bias in regression estimates.

Paper II presents and develops a multivariate model of trait evolution which is based on an Ornstein–Uhlenbeck type stochastic process, often used for studying trait adaptation, co–evolution, allometry or trade–offs. Alongside the description of the model and discussion of its most important features is a software package to estimate its parameters. We are not limited to looking at interactions between traits and environment. Within the presented multivariate framework it is possible to pose hypotheses about adaptation or trade–offs which can be rigorously tested with the correction for the phylogeny.

The thesis itself is made up of eight chapters. The first is this one, the introduction. Chapter 2 reminds the reader of the general linear regression model which is fundamental to the findings of the two papers. Chapters 3 and 4 are introductions to Paper I and Paper II, respectively. Possible directions of future development are outlined in chapter 5. Chapter 6 contains detailed derivations concerning bias due to measurement error for specific covariance structures in predictor variables and their measurement errors. Chapters 7 and 8 concern properties of the stochastic processes considered in Paper II and matrix parametrizations needed for the estimation procedures in the software package accompanying Pa-

per II.

In what follows we use the convention that matrices are written in bold-face, vectors as columns in normal font with an arrow above (*i.e.* $\vec{Z}$) and scalar values in normal font. By $\mathbf{I}$ we denote the identity matrix of appropriate size. For a matrix $\mathbf{M}$, $\mathbf{M}^T$ denotes the transpose of $\mathbf{M}$, $\mathbf{M}^{-1}$ the inverse and $\mathbf{M}^{-T}$ the transpose of the inverse of $\mathbf{M}$. The Hadamard product between two matrices is denoted by $*$ and the Kronecker product by $\otimes$. For a matrix $\mathbf{M}$ the operation $\text{vec}(\mathbf{M})$ means the vectorization, *i.e.* stacking of columns onto each other, of the matrix $\mathbf{M}$. The notation $\mathcal{N}(\vec{\mu}, \mathbf{V})$ means a normal distribution with mean vector $\vec{\mu}$ and covariance matrix $\mathbf{V}$. For two random variables X and Y the notation $X \perp\!\!\!\perp Y$ means that they are independent.

We assume that the phylogenetic tree is rooted and has a time arrow attached to it. We use the word time to describe the branch lengths allowing for different biological interpretations of it. On the phylogenetic tree, different times connected to species require distinct notation,

- t_i is the time of the i -th species on the phylogeny (we do not need to assume anywhere an ultrametric tree),
- $t_{a_{ij}}$ is time of divergence of tip species i and j ,
- $t_{1_i}, t_{2_i}, \dots, t_{m_i}$ are the consecutive times of changes of niches on the lineage ending at tip species i , with t_{0_i} denoting the time of the root and $t_{m_i} = t_i$.

See Figure 1 in Paper II for the illustration of t_i and $t_{a_{ij}}$ on the phylogenetic tree.

2 The general regression framework

The linear regression model can be considered to be one of the most common tools to study relationships between variables. It has been studied for over two hundred years now with Gauss describing the least squares method at the end of the eighteenth century. Due to the variety of fields and situations where it is applied a multitude of variations and particular cases of it have been considered. The phylogenetic comparative methods field is no exception.

We begin by reminding the reader of the multivariate linear regression model,

$$\mathbf{Y} = \mathbf{X}\mathbf{B} + \mathcal{E},$$

where $\mathbf{X} := [\vec{X}_1; \dots; \vec{X}_n]^T$, $\mathbf{Y} := [\vec{Y}_1; \dots; \vec{Y}_n]^T$, n is the number of observations and $\text{vec}(\mathcal{E}) \sim \mathcal{N}(\mathbf{0}, \mathbf{V})$ where $\mathbf{V} = \mathbf{I} \otimes \mathbf{\Sigma}_e$. We assume that each observed $\vec{Y}_i$ is of length k_Y and each observed $\vec{X}_i$ is of length k_X so that the matrix $\mathbf{B}$ has k_X rows and k_Y columns. The k_X predictor variables can be either continuous (*e.g.* certain trait values) or categorical (*e.g.* different environments). To estimate the unknown $\mathbf{B}$ and $\mathbf{\Sigma}_e$ parameters the least squares (LS) estimation is most commonly used, which is,

$$\begin{aligned} \hat{\mathbf{B}} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \\ \hat{\mathbf{\Sigma}}_e &= \frac{1}{n-1} \sum_{i=1}^n (\vec{Y}_i - \hat{\mathbf{B}} \vec{X}_i)(\vec{Y}_i - \hat{\mathbf{B}} \vec{X}_i)^T. \end{aligned} \quad (2.1)$$

This estimator (after [11]) is unbiased

$$\mathbb{E} [\hat{\mathbf{B}}] = \mathbb{E} [(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}] = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} \mathbf{B} = \mathbf{B}.$$

We will assume (to avoid cumbersome notation) for the moment that $\mathbf{\Sigma}_e = \sigma^2 \mathbf{I}$ and then the variance of the estimate of the j -th column of $\mathbf{B}$ (since $\mathbf{\Sigma}_e$ is diagonal this will be the same for each column) is,

$$\text{Var} [\hat{\mathbf{B}}_{\cdot j}] = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \sigma^2 \mathbf{I} \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

The confidence bound for the i -th element of each column j of $\mathbf{B}$ from equation (2.1) is $\hat{\mathbf{B}}_{ij} \pm t_{\alpha/2} S \sqrt{((\mathbf{X}^T \mathbf{X})^{-1})_{ii}}$, where S^2 is the usual unbiased estimate for σ^2 and $t_{\alpha/2}$ is the appropriate quantile of the t distribution with $k_Y \cdot n - k_X \cdot k_Y - 1$ degrees of freedom.

The above description assumed that the different observations are independent of each other. This will not always be the case and there are many applications where the noise terms are correlated, for example time series data. We denote the noise covariance structure by the matrix $\mathbf{V}$ and remind the reader, after [8], of the theory in this situation.

It is known that the naïve least squares estimator, $\hat{\mathbf{B}}$ as in equation (2.1), is still an unbiased estimator of $\mathbf{B}$ even with arbitrary dependencies

in the noise. However the second moment properties will not hold in this situation. If $\mathbf{V}$ is known up to a scalar, σ^2 , then one can use the generalized least squares (GLS) estimator,

$$\begin{aligned}\hat{\mathbf{B}} &= (\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}^{-1} \mathbf{V}^{-1} \mathbf{Y} \\ \hat{\sigma}^2 &= \frac{1}{n - k_X \cdot k_Y} (\vec{Y} - \text{vec}(\mathbf{X} \hat{\mathbf{B}}))^T \mathbf{V}^{-1} (\vec{Y} - \text{vec}(\mathbf{X} \hat{\mathbf{B}})),\end{aligned}$$

where $\vec{Y} := \text{vec}(\mathbf{Y})$ and this will have the exact same nice properties as the LS estimator.

Following [8] (see also [11]) we assume that our data model is $\vec{Y} = \mathbf{D} \vec{\beta} + \vec{\epsilon}$, where the response vector $\vec{Y} \in \mathbb{R}^{n \cdot k_Y}$ (this corresponds to $\text{vec}(\mathbf{Y})$ from the independent observations case), the design matrix $\mathbf{D} \in \mathbb{R}^{(n \cdot k_Y) \times k_X}$ ($\mathbf{D} = \mathbf{X} \otimes \mathbf{I}$ in the independent observations case) is of full rank k_X , the unknown coefficient vector $\vec{\beta} \in \mathbb{R}^{k_Y \cdot k_X}$ ($\vec{\beta} = \text{vec}(\mathbf{B}^T)$ in the independent observations case) and the $\vec{\epsilon}$ vector is distributed as $\mathcal{N}(\vec{0}, \mathbf{V})$, with the covariance matrix $\mathbf{V} \in \mathbb{R}^{(n \cdot k_Y) \times (n \cdot k_Y)}$ known up to a positive constant σ^2 ($\mathbf{V} = \mathbf{I} \otimes \Sigma_e$ in the independent observations case). The aim is to estimate the unknown $\vec{\beta}$ vector and via a linear transformation we can restate the problem as an ordinary least squares (OLS) one.

As $\mathbf{V}^{-1}$ is symmetric positive definite it can be decomposed by the Cholesky decomposition as $\mathbf{V}^{-1} = \mathbf{P} \mathbf{P}^T / \sigma^2$, where $\mathbf{P}$ is lower-triangular. Then we can write,

$$\mathbf{P}^T \vec{Y} = \mathbf{P}^T \mathbf{D} \vec{\beta} + \mathbf{P}^T \vec{\epsilon},$$

and denoting $\vec{Y}^* := \mathbf{P}^T \vec{Y}$, $\mathbf{D}^* := \mathbf{P}^T \mathbf{D}$ and $\vec{\epsilon}^* := \mathbf{P}^T \vec{\epsilon}$ we get

$$\vec{Y}^* = \mathbf{D}^* \vec{\beta} + \vec{\epsilon}^*.$$

As now, $\text{Var}[\vec{\epsilon}^*] = \text{Var}[\mathbf{P}^T \vec{\epsilon}] = \mathbf{P}^T \mathbf{V} \mathbf{P} = \mathbf{P}^T \mathbf{P}^{-T} \mathbf{P}^{-1} \mathbf{P} = \mathbf{I}$, so $\vec{\epsilon}^* \sim \mathcal{N}(\vec{0}, \mathbf{I})$ and this means that the noise terms and thereby the entries of the transformed response vector $\vec{Y}^*$ are independent. We get that the estimate of $\vec{\beta}$ is,

$$\begin{aligned}\hat{\vec{\beta}} &= (\mathbf{D}^{*T} \mathbf{D}^*)^{-1} \mathbf{D}^{*T} \vec{Y}^* \\ &= (\mathbf{D} \mathbf{P} \mathbf{P}^T \mathbf{X})^{-1} \mathbf{D} \mathbf{P} \mathbf{P}^T \vec{Y}\end{aligned}$$

which gives the formula for the generalized least squares estimator,

$$\hat{\vec{\beta}} = (\mathbf{D} \mathbf{V}^{-1} \mathbf{D})^{-1} \mathbf{D} \mathbf{V}^{-1} \vec{Y}. \quad (2.2)$$

To estimate σ^2 ,

$$\begin{aligned}\hat{\sigma}^2 &= \frac{1}{n - k_X \cdot k_Y} (\vec{Y}^* - \mathbf{D}^* \hat{\vec{\beta}})^T (\vec{Y}^* - \mathbf{D}^* \hat{\vec{\beta}}) \\ &= \frac{1}{n - k_X \cdot k_Y} (\mathbf{P}^T \vec{Y} - \mathbf{P}^T \mathbf{D} \hat{\vec{\beta}})^T (\mathbf{P}^T \mathbf{Y} - \mathbf{P}^T \mathbf{D} \hat{\vec{\beta}}) \\ &= \frac{1}{n - k_X \cdot k_Y} (\vec{Y} - \mathbf{D} \hat{\vec{\beta}})^T \mathbf{P} \mathbf{P}^T (\vec{Y} - \mathbf{D} \hat{\vec{\beta}}) \\ &= \frac{1}{n - k_X \cdot k_Y} (\vec{Y} - \mathbf{D} \hat{\vec{\beta}})^T \mathbf{V}^{-1} (\vec{Y} - \mathbf{D} \hat{\vec{\beta}})\end{aligned}$$

and all the properties can be derived in the same way,

$$\begin{aligned} \mathbb{E} \left[\hat{\vec{\beta}} \right] &= \vec{\beta} \\ \text{Var} \left[\hat{\vec{\beta}} \right] &= \sigma^2 (\mathbf{D}^T \mathbf{V}^{-1} \mathbf{D})^{-1} \end{aligned}$$

with the confidence bounds for the i -th element of $\hat{\vec{\beta}}$ being, $\hat{\beta}_i \pm t_{\alpha/2} \sqrt{\hat{\sigma}^2 \sqrt{((\mathbf{D}^T \mathbf{V}^{-1} \mathbf{D})^{-1})_{ii}}}$, where $t_{\alpha/2}$ is the appropriate point from the t distribution with $n \cdot k_Y - k_Y \cdot k_X - 1$ degrees of freedom.

In phylogenetic comparative analysis we run into exactly this issue that our data points are dependent. If we are doing a phylogenetic regression then the noise covariance structure $\mathbf{V}$ is non-diagonal. The covariance matrix $\mathbf{V}$ will depend on the phylogeny but it is not obvious how. To deal with this one needs to assume some sort of stochastic model of evolution. If for example we assume a Brownian motion model of evolution then $\mathbf{V} = \mathbf{T} \otimes (\boldsymbol{\Sigma} \boldsymbol{\Sigma}^T)$ where $\mathbf{T}$ is the matrix of divergence times of species on the phylogeny, and $\boldsymbol{\Sigma}$ the diffusion matrix, see also chapter 4.1 and [13]. Using the earlier introduced notation the i -th, j -th entry of $\mathbf{T}$ will equal $t_{a_{ij}}$.

3 Presentation of Paper I

Paper I concerns the issue of correcting for measurement error in regression where the predictor variables and measurement errors can be correlated in an arbitrary way. The bias caused by measurement error in regression is a well studied topic in the case of independent observations. However as we show in the Paper I everything becomes more complicated when the predictor variables are dependent. If we allow for the measurement errors to be dependent between observations we get an even more complicated situation.

In comparative studies both cases are important. The predictor variables are usually also species' traits evolving on the phylogeny. The "measurement errors" are often due to intra-species variability. The most common trait data in a comparative analysis are species means from a number of observations. Attached to these means is the variability which can be treated as "measurement error". As the species are related by their phylogeny the intra-species variability should be too. Whether in practice this will be possible to quantify is another matter but we provide a theoretical framework and a method to deal with this. Included with Paper I is a short R script, GLSME.R (General Least Squares with Measurement Error) that implements the described bias correction. The code is available for download from <http://www.math.chalmers.se/~krzbar/GLSME>

In the next section we present the description of the general model considered in Paper I.

3.1 The measurement error bias

In Paper I we use the following notation: A variable with subscript o will denote an observed (with error) variable, while the subscript t will mean the true, unobserved value of the variable. We consider the general linear model,

$$\vec{Y}_t = \mathbf{D}_t \vec{\beta} + \vec{r}_t, \quad \vec{r}_t \sim \mathcal{N}(\vec{0}, \mathbf{V}_t),$$

where $\vec{Y}_t$ is a vector of n observations of the dependent variable, $\vec{\beta}$ is a vector of m parameters to be estimated, $\mathbf{D}_t$ is an $n \times m$ design matrix and $\vec{r}_t$ is a vector of n noise terms with $n \times n$ covariance matrix $\mathbf{V}_t$. Due to the correlation in the noise for estimating β we use the generalized least squares estimator of equation (2.2). To the above general linear model we want to introduce an error measurement model. We can write the model with errors in the design matrix and the response variables as,

$$\begin{aligned} \mathbf{D}_o &= \mathbf{D}_t + \mathbf{U} \\ \vec{Y}_o &= \vec{Y}_t + \vec{e}_y, \end{aligned}$$

where $\mathbf{U}$ is a $n \times m$ matrix of random observation errors in the elements of $\mathbf{D}_t$, and $\vec{e}_y$ is a vector of length n of observation errors in $\vec{Y}_t$. Each

column of $\mathbf{U}$ is a vector of observation errors for a predictor variable. Furthermore we assume that $\text{vec}(\mathbf{D}_t)$ and $\text{vec}(\mathbf{U})$ are normal random vectors, *i.e.* $\text{vec}(\mathbf{D}_t) \sim \mathcal{N}(\vec{0}, \mathbf{V}_d)$ and $\text{vec}(\mathbf{U}) \sim \mathcal{N}(\vec{0}, \mathbf{V}_u)$. To simplify the derivations we assume that $\mathbf{D}_t$ has zero-mean. Summing up, the regression with observation error model is the following,

$$\begin{aligned} \vec{Y}_t &= \mathbf{D}_t \vec{\beta} + \vec{r}_t & \text{vec}(\mathbf{D}_t) &\sim \mathcal{N}(\vec{0}, \mathbf{V}_d), \quad \vec{r}_t \sim \mathcal{N}(\vec{0}, \mathbf{V}_t) \\ \vec{Y}_o &= \vec{Y}_t + \vec{e}_y & \vec{e}_y &\sim \mathcal{N}(\vec{0}, \mathbf{V}_e) \\ \mathbf{D}_o &= \mathbf{D}_t + \mathbf{U} & \text{vec}(\mathbf{U}) &\sim \mathcal{N}(\vec{0}, \mathbf{V}_u) \\ \vec{Y}_o &= (\mathbf{D}_o - \mathbf{U}) \vec{\beta} + \vec{r}_t + \vec{e}_y. \end{aligned}$$

We assume that all the errors are independent of each other and the other variables, *i.e.* $\vec{r}_t \perp \vec{e}_y$, $\vec{r}_t \perp \mathbf{U}$, $\mathbf{U} \perp \vec{e}_y$, $\mathbf{U} \perp \mathbf{D}_t$ and $\vec{Y}_t \perp \vec{e}_y$. We can write the last equation as, $\vec{Y}_o = (\mathbf{D}_o - \mathbf{U}) \vec{\beta} + \vec{r}$, where $\vec{r} = \vec{r}_t + \vec{e}_y$, so $\vec{r} \sim \mathcal{N}(\vec{0}, \mathbf{V})$, where $\mathbf{V} = \mathbf{V}_t + \mathbf{V}_e$. In Paper I we write the noise as $\vec{r} = \vec{r}_t + \vec{e}_y - \mathbf{U} \vec{\beta}$ and consider the regression model $\vec{Y}_o = \mathbf{D}_o \vec{\beta} + \vec{r}$ but here we do not do this so that we do not have to consider iterative procedures. We do not assume any structure (in particular diagonality) of the covariance matrices $\mathbf{V}$, $\mathbf{V}_t$, $\mathbf{V}_e$, $\mathbf{V}_d$ or $\mathbf{V}_u$. The generalized least squares estimator of $\vec{\beta}$ from equation (2.2) can be written, as

$$\begin{aligned} \hat{\vec{\beta}} &= (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} ((\mathbf{D}_o - \mathbf{U}) \vec{\beta} + \vec{r}) \\ &= (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} (\mathbf{D}_t \vec{\beta} + \vec{r}). \end{aligned}$$

We want to compute the expectation of this estimator conditional on the observed predictor variables $\mathbf{D}_o$. We assumed $\vec{r}$ had zero mean and is independent of $\mathbf{D}_o$ so we have,

$$\begin{aligned} \mathbb{E} [\hat{\vec{\beta}} | \mathbf{D}_o] &= (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbb{E} [\mathbf{D}_t | \mathbf{D}_o] \vec{\beta} \\ &= (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbb{E} [\mathbf{D}_o - \mathbf{U} | \mathbf{D}_o] \vec{\beta} \\ &= (\mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbb{E} [\mathbf{U} | \mathbf{D}_o]) \vec{\beta}. \end{aligned}$$

It remains to calculate $\mathbb{E} [\mathbf{U} | \mathbf{D}_o]$, to do this we need to work with its vectorized form, $\mathbb{E} [\text{vec}(\mathbf{U}) | \mathbf{D}_o]$. Using common facts about the multivariate normal distribution and that $\mathbf{D}_o = \mathbf{D}_t + \mathbf{U}$ where $\mathbf{D}_t$, $\mathbf{U}$ have zero mean and are independent of each other we derive,

$$\begin{aligned} &\mathbb{E} [\text{vec}(\mathbf{U}) | \mathbf{D}_o] \\ &= \mathbb{E} [\text{vec}(\mathbf{U})] + \text{Cov} [\text{vec}(\mathbf{U}), \text{vec}(\mathbf{D}_o)] \text{Var} [\text{vec}(\mathbf{D}_o)]^{-1} (\text{vec}(\mathbf{D}_o) - \mathbb{E} [\text{vec}(\mathbf{D}_o)]) \\ &= \text{Var} [\text{vec}(\mathbf{U})] \text{Var} [\text{vec}(\mathbf{D}_o)]^{-1} \text{vec}(\mathbf{D}_o) \\ &= \mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o), \end{aligned}$$

where $\mathbf{V}_o = \mathbf{V}_d + \mathbf{V}_u$. (In setup of the model of [26] $\mathbf{V}_d$ would be the covariance matrix resulting from the assumed Brownian-motion of the predictors). Therefore we have that,

$$\mathbb{E} [\hat{\vec{\beta}} | \mathbf{D}_o] = \mathbf{K} \vec{\beta}, \quad (3.1)$$

where the so-called reliability matrix $\mathbf{K}$ is given by

$$\mathbf{K} = \mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \text{vec}^{-1}(\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o)), \quad (3.2)$$

where $\text{vec}^{-1}(\mathbf{v})$ for a vector $\mathbf{v}$ means the inverse of the vec operation with appropriate matrix size.

The above is a general formula for any covariance matrices $\mathbf{V}_o$ and $\mathbf{V}_u$. If we impose some structure on them then it will simplify. No assumptions are needed on the covariance matrices $\mathbf{V}_t$ and $\mathbf{V}_e$. If we write out the matrix multiplications elementwise we will see that all the predictors influence the bias of all the others. However the directions and magnitudes of these biases depend on the entries of $\mathbf{D}_o$ and the covariance matrices and so they can be arbitrary.

We write below some special cases of $\mathbf{V}_o$ and $\mathbf{V}_u$ where $\mathbf{K}$ simplifies with the details behind them in section 6.

Independent observations of predictors where Σ_d is the covariance matrix of the true (unobserved) predictors in a given observation and Σ_u is the covariance matrix of the measurement errors in predictors in a given observation

$$\mathbf{K} = (\Sigma_d + \Sigma_u)^{-1} \Sigma_d. \quad (3.3)$$

If the predictor is one dimensional instead of $\mathbf{K}$, Σ_d , Σ_u we use κ , σ_d^2 and σ_u^2 respectively and the above becomes,

$$\kappa = \frac{\sigma_d^2}{\sigma_d^2 + \sigma_u^2}. \quad (3.4)$$

Independent predictors (i.e. the predictors and their errors are independent between each other but not between observations), now Σ_{d_i} is the covariance matrix of the true (unobserved) i th predictor between observations and Σ_{u_i} the covariance matrix of the measurement error and $\mathbf{d}_{o_i}$ denotes the i th column of $\mathbf{D}_o$,

$$\mathbf{K} = \mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} [\Sigma_{u_1} \Sigma_{o_1}^{-1} \mathbf{d}_{o_1}; \dots; \Sigma_{u_m} \Sigma_{o_m}^{-1} \mathbf{d}_{o_m}]. \quad (3.5)$$

Further simplifications of this case are possible if one assumes independence of errors or proportionality between all of the Σ_{u_i} matrices and these are shown in section 6.

Single predictor

$$\mathbf{K} = \mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o \quad (3.6)$$

3.2 Mean square error analysis

In the case of a single predictor equation with independent observations it is well known that $\kappa \in (-1, 1)$, see equation (3.4). Therefore the

bias corrected estimator has a larger variance than the uncorrected one. This immediately implies a trade-off if we use the mean square error, $E[(\hat{\beta} - \beta)^2]$, to compare estimators.

In Paper I we study how the situation looks with dependent observations of a single predictor. We transform the formulae for the mean square errors of both estimators into a criterion in terms of estimable quantities, $\sigma_\beta^2/\kappa^2 < \sigma_\beta^2 + (\kappa - 1)^2\beta^2$, where σ_β^2 is the variance of the estimator of β .

Unlike in the independent observations case κ can take values outside the interval $(-1, 1)$. This results in a much more complicated situation as is described in section 6.4.2 of this thesis.

In the supplementary material to Paper I we present simulation results that show that this criterion works well even when all of its terms are estimated from the data.

4 Presentation of Paper II

Paper II is devoted to a multivariate stochastic differential equation model of trait evolution. With Paper II is an R software package, `mvSLOUCH` (multivariate Stochastic Linear Ornstein–Uhlenbeck models for phylogenetic Comparative Hypothesis) which covers nearly all cases possible in the framework of [24][25][34] (with the exception of some cases where the drift matrix is singular or does not have an eigendecomposition). The package will be available for download from <http://www.math.chalmers.se/~krzbar/mvSLOUCH> with a detailed description of the input and output routines.

We start our presentation of Paper II with a brief introduction to the most relevant stochastic processes in the field of phylogenetic comparative methods.

4.1 Brownian Motion

The first modelling approach via stochastic differential equations for comparative methods in a continuous trait setting is due to Felsenstein [13]. It proposed that the traits evolve as a Brownian motion along the phylogeny. A trait X is evolving as a Brownian motion if it can be described by the following stochastic differential equation,

$$dX(t) = \sigma dW(t),$$

where $dW(t)$ is white noise. More formally $W(t)$ is the Wiener process, it has independent increments with $W(t) - W(s) \sim \mathcal{N}(0, t - s), 0 \leq s \leq t$. This means that the displacement of the trait after time t from its initial value is $X(t) \sim \mathcal{N}(X(0), \sigma^2 t)$. This model assumes that there is absolutely no selective pressure on trait X whatsoever, so that the trait value just randomly fluctuates from its initial starting point.

A multi-dimensional Brownian motion model is immediate,

$$d\vec{X}(t) = \Sigma d\vec{W}(t),$$

where Σ will be the diffusion matrix, $\vec{W}$ will be the multidimensional Wiener process. All components of the random vector are normally distributed as $\vec{W}(t) - \vec{W}(s) \sim \mathcal{N}(\vec{0}, (t - s)\mathbf{I})$ and so $\vec{X}(t) \sim \mathcal{N}(\vec{X}(0), t\Sigma\Sigma^T)$. Such a trait evolution has the property that it is time reversible. It does not matter whether we look at it taking $\vec{X}(0)$ (moving forward) or $\vec{X}(t)$ (moving backward) as the starting point. Both $\vec{X}(t) - \vec{X}(0)$ and $\vec{X}(0) - \vec{X}(t)$ will be identically distributed. From this one can find an independent sample on the phylogenetic tree, *i.e.* the contrasts between nodes such that no pair of nodes shares a branch in the path connecting them. If the distance between two nodes is $t_i - t_j$ then the contrast vector will be distributed as $\vec{X}(t_i) - \vec{X}(t_j) \sim \mathcal{N}(\vec{0}, (t_i - t_j)\Sigma\Sigma^T)$. The

covariance matrix of all of the data will be $\mathbf{V} = \mathbf{T} \otimes (\mathbf{\Sigma} \mathbf{\Sigma}^T)$ where $\mathbf{T}$ is the matrix of divergence times of species on the phylogeny. In order to estimate the ancestral state and $\mathbf{\Sigma}$ in [13] the independent contrasts algorithm was proposed which is similar to the Cholesky decomposition described in section 2.

4.2 Ornstein–Uhlenbeck process

The lack of the possibility of adaptation in the Brownian Motion model of [13] was noticed in [24] and therefore a more complex Ornstein–Uhlenbeck type of framework was proposed,

$$d\vec{Z}(t) = -\mathbf{F}(\vec{Z}(t) - \vec{\Psi}(t))dt + \mathbf{\Sigma}d\vec{W}(t), \quad (4.1)$$

termed here the Ornstein–Uhlenbeck (OU) model. The $\mathbf{F}$ matrix is called the drift matrix, $\vec{\Psi}(t)$ the drift vector and $\mathbf{\Sigma}$ the diffusion matrix. This stochastic differential equation is a linear equation in the narrow sense. The general model (4.1) has been only partially implemented in some R packages and in this work we present a nearly complete implementation of it. We use the word nearly, as the package does not allow for certain very specific parameter combinations which should not be of biological significance.

We also consider in detail model with a very important decomposition of $\mathbf{F}$ from equation (4.1) which results in the multivariate Ornstein–Uhlenbeck Brownian motion (mvOUBM) model

$$d \begin{bmatrix} \vec{Y} \\ \vec{X} \end{bmatrix} (t) = - \left[\begin{array}{c|c} \mathbf{A} & \mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \left(\begin{bmatrix} \vec{Y} \\ \vec{X} \end{bmatrix} (t) - \begin{bmatrix} \vec{\psi} \\ \vec{0} \end{bmatrix} (t) \right) dt + \left[\begin{array}{c|c} \mathbf{\Sigma}^{yy} & \mathbf{\Sigma}^{yx} \\ \hline \mathbf{\Sigma}^{xy} & \mathbf{\Sigma}^{xx} \end{array} \right] d\mathbf{W}(t). \quad (4.2)$$

This is a multivariate generalization of the OUBM model from [26].

The different parameters in the above model have biological interpretations but one must be careful how to interpret them. The $\vec{\Psi}(t)$ function represents a deterministic optimum value for $\vec{Z}(t)$, equation (4.1) and $\vec{\psi}(t)$ the deterministic part of the optimum for $\vec{Y}(t)$ in equation (4.2) given that $\mathbf{F}$ or $\mathbf{A}$ have positive real part eigenvalues respectively. The individual entries in $\mathbf{F}$ and $\mathbf{A}$ can be tricky to interpret. However their eigendecomposition is much easier. The eigenvalues, if they have positive real part, can be understood to control the speed of the traits’ approach to their optima while the eigenvectors indicate the path towards the optimum. The diffusion matrix in general represents stochastic perturbations to the traits’ approaching their optimum. These perturbations can come from many sources internal or external, *e.g.* they can represent unknown/unmeasured components of the system under study or some random genetic changes linked to the traits.

We have described here the processes as if they were evolving on a single lineage. In the phylogenetic comparative setting these processes are evolving on a phylogenetic tree. At each internal node they branch off into independently evolving components. We assume that we know the tree and include these branching times in the estimation procedure. The mathematics of this is straightforward but it is non-trivial to effectively implement the calculation of the log-likelihood function.

4.3 The mvSLOUCH package

With Paper II we present an R package, mvSLOUCH, that implements a (heuristic) maximum likelihood estimation to estimate the parameters of the stochastic differential equation (4.1). The package also allows for the estimation of parameters of a special, but very appealing, submodel of equation (4.1) where some of the variables behave as Brownian motions, equation (4.2). This is a multivariate generalization of the model presented in [26].

The implementation of the estimation procedure requires combining advanced linear algebra and computing techniques with probabilistic and statistical knowledge. This is due to the high model dimensionality and nearly arbitrary dependency structure between observations. The most interesting techniques employed are described in Paper II, its appendix, as well as in sections 7 and 8 of this thesis.

The package can estimate not only the parameters but simulate the trajectories of described stochastic differential equation models, both on an arbitrary phylogeny and on a single lineage.

In Figure 1 we can see such a trajectory simulated from the model of equation (4.2). The simulation reveals that small changes in the predictor process, X (*e.g.* environment), cause very large changes in the optimal processes for the two traits Y_1 and Y_2 (see Paper II for the used nomenclature). This kind of behaviour could be biologically interpreted as specialization, when small changes of the surroundings require the organism to make large changes in order to be best adapted to it. We can also see that the traits are tracking their respective optimal process, but with a lag and have a very gentle response to extremes in them. This lag means that the organism is unable to change quickly enough to adapt and misses out on narrow (in time) optimal niches which results in being close to optimality when the optimal process becomes toned down again. The trait and optimal processes are obviously negatively correlated. Other examples of simulated trajectories are given in section 7, Figure 4 and in Paper II. Through these examples one can see that these processes can exhibit all sorts of dynamics.

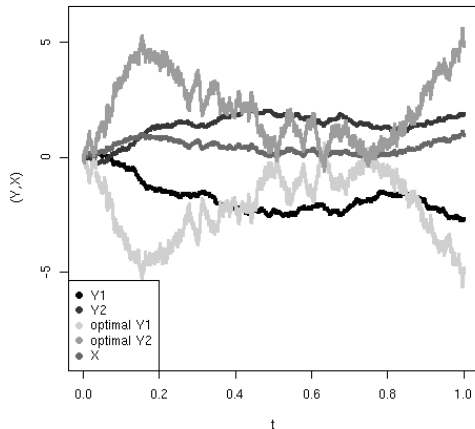


Figure 1: Trajectory and optimal values of the multivariate Ornstein-Uhlenbeck Brownian Motion (mvOUBM) stochastic process.

4.4 Example data analysis

In Paper II we reanalyze a data set of deer antler length, and male and female body masses from [44] under our presented models. Our results confirm those in [44], that (i) there is positive linear relationship between the logarithms of antler length and male body mass and (ii) the mating tactic does not influence directly the antler length and also not the male body mass. The allometry between antler length and male body mass is greater than $1/3$ indicating that there is more than just a proportional increase of antler length (one dimensional) when body mass increases (three dimensional). We can also observe that the estimates of the effect of breeding group size on the antler length and male body mass are larger with the increase in breeding group size.

However our analysis in addition shows that antler length and male body mass are adapting very rapidly to changes in female body mass. We can also see that they adapt independently of each other. There is no direct influence of one variable on the primary optimum of the other. All dependencies between antler length and male body mass are due to the common female body mass predictor variable and correlations in noise pushing them away from their respective optimum values. The latter could be *e.g.* due to mutations in some common genetic regulatory mechanism. This however is pure speculation as our models show only that there is some general “noise” mechanism and not what it is

due to. This adaptation result was not observed in the study of [44] simply because they did not consider a model that allowed for it. To account for phylogenetic effects they assumed a Brownian motion model of phenotypic trait evolution.

5 Future work

The results of Paper I can be developed in the direction of studying the effects of measurement error on the estimation of covariance structures under a Brownian motion model of evolution and possibly more complex models afterwards.

Furthermore we have three directions of development from Paper II. The first one concerns modelling trait evolution and interactions among the traits and the environment. All the models based on continuous stochastic differential equations considered so far have assumed that the drift function is a deterministic step function. The regimes on the phylogeny where it is constant are presumed to be known or preestimated. The next step in developing these models would be to assume a random process for the drift function coupled with the stochastic differential model.

The randomness can be added on a number of levels. The simplest case is to assume two possible regime types with exponential times between regime changes (if the times of change are known then this would be the same as estimating what regime was at the root of the tree). Mathematically speaking we have a function taking on two possible values (which are unknown) driven by a Poisson process and this is joined with a stochastic differential equation. Another possibility is that the drift function takes on a finite number of (unknown) values assuming that a Markov structure governs their changes. To connect this with the traits, we can assume that the corresponding probability transition matrix functionally depends on the current trait values.

The second direction is to address the various statistical questions that arise from considering such a dependency structure between different observations and aspects of the stochastic processes used for modelling that might not have been traditionally looked at. To the best of our knowledge issues of parameter estimability, effective sample size nor the correct way of constructing confidence intervals have not been addressed in a satisfactory manner. The need to capture interactions between the modelled traits requires us to study the behaviour of the mean and covariance functions. Their asymptotics are well known but it remains to describe *e.g.* how their sign changes with time or how do the magnitudes of the regression slopes evolve.

The third idea of development would be to join the contents of the two papers presented in the thesis together, to correct for bias due to measurement error when estimating the parameters of the processes evolving on the phylogeny. This might seem straightforward at first but it is not. The measurement error model considered in Paper I assumes that we essentially know all covariance structures (covariance between predictors, covariance between measurement errors and noise covariance structure) up to at most a scalar. The estimation procedures implemented in this

package can estimate the covariance structure between predictors and noise. However in Paper I it is noticed that correcting for bias is not always a good idea as it might increase the mean square error. This was in a situation where the noise and predictor covariance structures were assumed to be known. It remains to study how the mean square error is affected by the bias correction in a situation where the other covariance structures are estimated.

We are aware that conditioning on the phylogeny is a gross simplification and it would be ideal to combine phenotypic and molecular data in order to jointly study the times of species diversification with the development of their traits. However this would be very complex and computationally extremely demanding. At the moment the estimation of the phylogeny with estimation of parameters of the process modelling trait evolution is only done for the Brownian motion process [29]. We hope to be able to develop this into more complex stochastic models.

6 Supplementary material A: Bias caused by measurement error; to correct or not to correct (Paper I)

In section 3.1 we derived a general formula for the reliability matrix $\mathbf{K}$, equation (3.2) and afterwards presented formulae in situations where some kind of independence was assumed. In sections 6.1 – 6.3. we treat the subject in more detail. Section 6.4 answers the question on whether it is worth to correct for the bias caused by measurement error. Section 6.5 considers the case where we have fixed effects and more generally predictors observed without error. Finally section 6.6 finds discusses conditions when it is possible to derive analytical formulae for the unconditional (of the observed design matrix) bias caused by measurement error in predictors.

6.1 Independent observations of predictors

The case of independent observations is widely described in the literature, see *e.g.* [1] (where they show that in the multivariate case the bias can be in their words “*of arbitrary sign and magnitude*” and “*There do not seem to be any rules of thumb that would permit one to make even qualitative statements about the nature of this bias.*”), [15], [16] (where the effect of error in data analysis is shown on examples) or [19] (where some statistical results are derived *e.g.* mean square error results, connection to maximum likelihood estimates, estimability and distributional properties).

The situation with a single predictor is described in Paper I (see equation (12) there). Here we will briefly describe the multipredictor case. Let each row of $\mathbf{D}_t$ (the vector of predictors) be distributed as $\mathcal{N}(0, \mathbf{\Sigma}_d)$ and each row of $\mathbf{U}$ (vector of errors for each observation) be distributed as $\mathcal{N}(0, \mathbf{\Sigma}_u)$. Then each row of $\mathbf{D}_o$ is distributed as $\mathcal{N}(0, \mathbf{\Sigma}_d + \mathbf{\Sigma}_u)$ and we have the distributions $\text{vec}(\mathbf{D}_o) \sim \mathcal{N}(0, (\mathbf{\Sigma}_d + \mathbf{\Sigma}_u) \otimes \mathbf{I})$ and $\text{vec}(\mathbf{U}) \sim \mathcal{N}(0, \mathbf{\Sigma}_u \otimes \mathbf{I})$. Putting $\mathbf{\Sigma}_o := \mathbf{\Sigma}_d + \mathbf{\Sigma}_u$ we arrive at equation (3.3),

$$\mathbb{E} \left[\hat{\vec{\beta}} | \mathbf{D}_o \right] = (\mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o \mathbf{\Sigma}_o^{-1} \mathbf{\Sigma}_u) \vec{\beta} = (\mathbf{I} - \mathbf{\Sigma}_o^{-1} \mathbf{\Sigma}_u) \vec{\beta}$$

which simplifies to $\mathbf{\Sigma}_o^{-1} \mathbf{\Sigma}_d \vec{\beta}$. This also gives us the unconditional bias as in this case the formula does not depend on $\mathbf{D}_o$, $\mathbb{E} \left[\hat{\vec{\beta}} \right] = \mathbf{\Sigma}_o^{-1} \mathbf{\Sigma}_d \vec{\beta}$.

We can see that the above formula is a straightforward generalization of the single predictor case. Some discussion about this situation can be found in [9], with detailed derivations given in [19]. There are examples (see [1]) showing that the one-dimensional shrinkage to zero does not generalize and no qualitative statements about the bias in the general case

can be made. In particular in [19] it is shown that if the reliability matrix ($\Sigma_o^{-1} \Sigma_t$) is estimated by maximum likelihood, then the bias corrected estimates will be the maximum likelihood ones.

In the multivariate setting we can get an upward and downward bias of arbitrary size in a component depending on its correlation with another component (compare with example 2 in [1] even for predictors that have a slope of 0, *i.e.* they are not related to the response). To illustrate this we will write out the formula for the two dimensional case. Denote,

$$\Sigma_u = \begin{bmatrix} \sigma_{11} & \sigma_{12} \\ \sigma_{12} & \sigma_{22} \end{bmatrix} \quad \Sigma_o^{-1} = \begin{bmatrix} p_{11} & p_{12} \\ p_{12} & p_{22} \end{bmatrix}$$

and let $\vec{\beta} = [\beta_1, \beta_2]^T$. We then have

$$\mathbb{E} \left[\begin{bmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{bmatrix} \right] = \begin{bmatrix} (1 - (p_{11}\sigma_{11} + p_{12}\sigma_{12}))\beta_1 - (p_{11}\sigma_{12} + p_{12}\sigma_{22})\beta_2 \\ (1 - (p_{22}\sigma_{22} + p_{12}\sigma_{12}))\beta_2 - (p_{22}\sigma_{12} + p_{12}\sigma_{11})\beta_1 \end{bmatrix}.$$

So we can see that both predictors cause bias in each other but the direction of the bias depends on the covariance between the predictors and between errors (as shown in [1] and compare to a similar example on pages 109–110 of [9]).

6.2 Independent predictors

In the general case of dependent observations the formula for $\mathbb{E} [\hat{\vec{\beta}} | \mathbf{D}_o]$ does not simplify so well. However if we assume that the predictors are independent (the columns of $\mathbf{D}_t$ are independent) and that the errors in different predictors are also independent (the columns of $\mathbf{U}$ are independent) then some simplifications can be made.

Let us assume that the i -th column of $\mathbf{D}_t$ is distributed as $\mathcal{N}(\vec{0}, \Sigma_{d_i})$ (Σ_{d_i} is the covariance matrix between the i -th predictor in different observations) and that the error in the i -th predictor is distributed as $\mathcal{N}(\vec{0}, \Sigma_{u_i})$. The distribution of the i -th column of $\mathbf{D}_o$ will be $\mathcal{N}(\vec{0}, \Sigma_{d_i} + \Sigma_{u_i})$ and we will denote $\Sigma_{o_i} := \Sigma_{d_i} + \Sigma_{u_i}$. We then have that,

$$\begin{aligned} \text{vec}(\mathbf{D}_o) &\sim \mathcal{N} \left(\vec{0}, \begin{bmatrix} \Sigma_{o_1} & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \Sigma_{o_2} & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \Sigma_{o_m} \end{bmatrix} \right) \\ \text{vec}(\mathbf{U}) &\sim \mathcal{N} \left(\vec{0}, \begin{bmatrix} \Sigma_{u_1} & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \Sigma_{u_2} & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \Sigma_{u_m} \end{bmatrix} \right). \end{aligned}$$

This then gives us after performing the matrix multiplication in the formula for the conditional expectation that the i -th column of $\mathbb{E} [\mathbf{U} | \mathbf{D}_o]$ is

$\Sigma_{u_i} \Sigma_{o_i}^{-1} \mathbf{d}_{o_i}$, where $\mathbf{d}_{o_i}$ is the i -th column of $\mathbf{D}_o$. This gives us equation (3.5),

$$\mathbb{E} [\hat{\beta} | \mathbf{D}_o] = (\mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} [\Sigma_{u_1} \Sigma_{o_1}^{-1} \mathbf{d}_{o_1}; \dots; \Sigma_{u_m} \Sigma_{o_m}^{-1} \mathbf{d}_{o_m}]) \vec{\beta}.$$

If we assume further that errors in different observations are independent and that the covariance matrices are of the form $\Sigma_{u_i} = \sigma_i^2 \sigma_u^2 \mathbf{I}$ and $\Sigma_{d_i} = \sigma_i^2 \mathbf{T}$, where $\mathbf{T}$ is some symmetric positive-definite matrix (predictors evolve as independent Brownian Motions along the phylogenetic tree), then $\Sigma_{d_i} + \Sigma_{u_i} = \sigma_i^2 (\mathbf{T} + \sigma_u^2 \mathbf{I})$,

$\mathbf{V}_u = \text{diag}(\sigma_1^2, \dots, \sigma_m^2) \otimes \sigma_u^2 \mathbf{I}$ and $\mathbf{V}_o = \text{diag}(\sigma_1^2, \dots, \sigma_m^2) \otimes (\mathbf{T} + \sigma_u^2 \mathbf{I})$, where $\text{diag}(\mathbf{v})$ is a diagonal matrix with the vector $\mathbf{v}$ on the diagonal.

From

$$\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o) = (\mathbf{I} \otimes (\frac{1}{\sigma_u^2} \mathbf{T} + \mathbf{I})^{-1}) \text{vec}(\mathbf{D}_o),$$

we get

$$\mathbb{E} [\mathbf{U} | \mathbf{D}_o] = (\frac{1}{\sigma_u^2} \mathbf{T} + \mathbf{I})^{-1} \mathbf{D}_o,$$

so that

$$\mathbb{E} [\hat{\beta} | \mathbf{D}_o] = (\mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} (\frac{1}{\sigma_u^2} \mathbf{T} + \mathbf{I})^{-1} \mathbf{D}_o) \vec{\beta}. \quad (6.1)$$

If we do not assume that the errors between different observations are independent, *i.e.* $\Sigma_{u_i} = \sigma_i^2 \Sigma_u$, this will change as,

$\mathbf{V}_u = \text{diag}(\sigma_1^2, \dots, \sigma_m^2) \otimes \Sigma_u$ and $\mathbf{V}_o = \text{diag}(\sigma_1^2, \dots, \sigma_m^2) \otimes (\mathbf{T} + \Sigma_u)$. It follows

$$\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o) = (\mathbf{I} \otimes (\mathbf{T} \Sigma_u^{-1} + \mathbf{I})^{-1}) \text{vec}(\mathbf{D}_o),$$

implying

$$\mathbb{E} [\mathbf{U} | \mathbf{D}_o] = (\mathbf{T} \Sigma_u^{-1} + \mathbf{I})^{-1} \mathbf{D}_o,$$

which gives

$$\mathbb{E} [\hat{\beta} | \mathbf{D}_o] = (\mathbf{I} - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} (\mathbf{T} \Sigma_u^{-1} + \mathbf{I})^{-1} \mathbf{D}_o) \vec{\beta}. \quad (6.2)$$

6.3 Single predictor

The single predictor case of the general formula (3.1) presents an illustrative example which will be treated here in detail. The matrix $\mathbf{D}_o$ is a single column matrix and the generalized least squares estimator is,

$$\hat{\beta} = (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} (\mathbf{D}_t \beta + \vec{r}_t + \vec{e}_y),$$

and its conditional expectation (as $\vec{r}_t$ and $\vec{e}_y$ are independent of the rest) is,

$$\mathbb{E} [\hat{\beta} | \mathbf{D}_o] = (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbb{E} [\mathbf{D}_t | \mathbf{D}_o] \beta,$$

and using $E[\mathbf{D}_t|\mathbf{D}_o] = \mathbf{D}_o - E[\mathbf{U}|\mathbf{D}_o]$ (here as the matrix contains one column we can work with it directly there is no need to vectorize it) we get

$$E[\hat{\beta}|\mathbf{D}_o] = (1 - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o])\beta.$$

The trick again is to compute $E[\mathbf{U}|\mathbf{D}_o]$, as we assumed that all the variables are multivariate normal, then together the $\mathbf{U}$ and $\mathbf{D}_o$ will also be multivariate normal and so we get,

$$E[\mathbf{U}|\mathbf{D}_o] = E[\mathbf{U}] + \text{Cov}[\mathbf{U}, \mathbf{D}_o] \text{Var}[\mathbf{D}_o]^{-1} (\mathbf{D}_o - E[\mathbf{D}_o])$$

Since we assumed that $E[\mathbf{U}] = E[\mathbf{D}_o] = \mathbf{0}$ we have,

$$\text{Cov}[\mathbf{U}, \mathbf{D}_o] = \text{Cov}[\mathbf{U}, \mathbf{D}_t + \mathbf{U}] = \text{Var}[\mathbf{U}] = \mathbf{V}_u,$$

which gives $E[\mathbf{U}|\mathbf{D}_o] = \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o$ and from this we can write the final formula of equation (3.6),

$$E[\hat{\beta}|\mathbf{D}_o] = (1 - (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o)\beta.$$

If $\mathbf{V}_u = \sigma_u^2 \mathbf{I}$ and $\mathbf{V}_o = (\sigma_d^2 + \sigma_u^2) \mathbf{I}$ then we get equation (3.4),

$$E[\hat{\beta}|\mathbf{D}_o] = (1 - \frac{\sigma_u^2}{\sigma_d^2 + \sigma_u^2})\beta = \frac{\sigma_d^2}{\sigma_d^2 + \sigma_u^2}\beta = E[\hat{\beta}],$$

which is consistent with the standard regression case if we estimate σ_u^2 and σ_o^2 from the sum of squares.

6.4 Mean square error (MSE) analysis

In the derivation of the conditional expectation of the generalized least squares estimator we used the matrix $\mathbf{V}$. One could however be oblivious to the fact of error in the response and use $\mathbf{V}_t$ instead. As we shall see here this does not have an effect on the form of bias derivations (but will numerically change the results). Let us first consider a more general version of the generalized least squares estimator, $\hat{\vec{\beta}} = (\mathbf{D}_o^T \mathbf{Z} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{Z} \vec{Y}_o$, where $\mathbf{Z}$ is some symmetric, positive definite matrix. Then doing the same calculations as before we get that,

$$E[\hat{\vec{\beta}}|\mathbf{D}_o] = (\mathbf{I} - (\mathbf{D}_o^T \mathbf{Z} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{Z} \text{vec}^{-1}(\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o))) \vec{\beta},$$

so the bias correction is of the same form for any symmetric, positive definite matrix $\mathbf{Z}$, however the covariance matrix of $\hat{\vec{\beta}}$ can substantially change and then so will the mean square error. The mean square error of an estimator is defined as,

$$\text{MSE}[\hat{\vec{\beta}}|\mathbf{D}_o] = E[(\hat{\vec{\beta}} - \vec{\beta})^T (\hat{\vec{\beta}} - \vec{\beta})|\mathbf{D}_o]$$

and equals to $\text{Tr}(\text{Var}[\hat{\vec{\beta}}|\mathbf{D}_o]) + \mathbf{E}[\hat{\vec{\beta}} - \vec{\beta}|\mathbf{D}_o]^T \mathbf{E}[\hat{\vec{\beta}} - \vec{\beta}|\mathbf{D}_o]$, where $\text{Tr}(\mathbf{M})$ is the trace of a matrix $\mathbf{M}$.

The covariance of the estimator is,

$$(\mathbf{D}_o^T \mathbf{Z} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{Z} \left(\text{Var}[\mathbf{U}\vec{\beta}|\mathbf{D}_o] + \mathbf{V} \right) \mathbf{Z} \mathbf{D}_o (\mathbf{D}_o^T \mathbf{Z} \mathbf{D}_o)^{-1}.$$

It remains to calculate $\text{Var}[\mathbf{U}\vec{\beta}|\mathbf{D}_o]$, as we assumed both $\mathbf{U}$ and $\mathbf{D}_o$ are multivariate normal we get this from the formula for the conditional covariance,

$$\text{Var}[\mathbf{U}\vec{\beta}|\mathbf{D}_o] = \text{Var}[\mathbf{U}\vec{\beta}] - \text{Cov}[\mathbf{U}\vec{\beta}, \mathbf{D}_o] \text{Var}[\mathbf{D}_o]^{-1} \text{Cov}[\mathbf{D}_o, \mathbf{U}\vec{\beta}],$$

where element i, j of $\text{Var}[\mathbf{U}\vec{\beta}]$ is $\sum_{r=1}^m \sum_{k=1}^m \vec{\beta}_k \vec{\beta}_r \mathbf{V}_u[(r-1)n+i, (k-1)n+j]$ and element $i, (k-1)n+j$ of $\text{Cov}[\mathbf{U}\vec{\beta}, \mathbf{D}_o]$ is $\sum_{r=1}^m \vec{\beta}_r \mathbf{V}_u[(r-1)n+i, (k-1)n+j]$.

6.4.1 Mean square error for different GLS estimators

We will consider the mean square error for the general case of some $\mathbf{Z}$ matrix as the residual covariance structure. To see that there is no qualitative difference when we take $\mathbf{Z} = \mathbf{V}^{-1}$ or $\mathbf{Z} = \mathbf{V}_t^{-1}$ or any other matrix we introduce the general notation,

$$\begin{aligned} \mathbf{A} &:= (\mathbf{D}_o^T \mathbf{Z}^{-1} \mathbf{D}_o)^{-1}, \\ \mathbf{B} &:= (\text{vec}^{-1}(\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o)))^T \mathbf{Z}^{-1} \mathbf{D}_o. \end{aligned}$$

We introduce a third matrix $\mathbf{C}$ whose definition will depend on the choice of $\mathbf{Z}$. If $\mathbf{Z} = \mathbf{V}^{-1}$ then

$$\mathbf{C} := \mathbf{D}_o^T \mathbf{V}^{-1} \text{Var}[\mathbf{U}\vec{\beta}|\mathbf{D}_o] \mathbf{V}^{-1} \mathbf{D}_o$$

and if $\mathbf{Z} = \mathbf{V}_t^{-1}$ then

$$\mathbf{C} := \mathbf{D}_o^T \mathbf{V}_t^{-1} \left(\text{Var}[\mathbf{U}\vec{\beta}|\mathbf{D}_o] + \mathbf{V}_e \right) \mathbf{V}_t^{-1} \mathbf{D}_o.$$

We can notice immediately that $\mathbf{A}$ and $\mathbf{C}$ are symmetric positive definite matrices and recognize $\mathbf{A}$ as the estimated covariance of the estimate of $\vec{\beta}$.

In the case of the uncorrected estimator,

$$\text{MSE}[\hat{\vec{\beta}}|\mathbf{D}_o] = \text{Tr}(\mathbf{A} + \mathbf{A}\mathbf{C}\mathbf{A}) + \vec{\beta}^T (\mathbf{B}\mathbf{A}\mathbf{A}\mathbf{B}^T) \vec{\beta} \quad (6.3)$$

and for the corrected estimator the mean square error becomes,

$$\text{MSE}[\hat{\vec{\beta}}_{\text{corr}}|\mathbf{D}_o] = \text{Tr} \left((\mathbf{I} - \mathbf{A}\mathbf{B})^{-1} (\mathbf{A} + \mathbf{A}\mathbf{C}\mathbf{A}) (\mathbf{I} - \mathbf{A}\mathbf{B})^{-T} \right). \quad (6.4)$$

We notice that in both cases the formula for the mean square error depends on $\vec{\beta}$, in the second formula only through $\mathbf{C}$.

6.4.2 Mean square error for a single predictor

Turning to the single predictor case one can find more visible conditions for the corrected estimator having a smaller mean square error. Our notation for the matrix components $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{C}$ of the mean square error formulae simplify to,

$$\begin{aligned} a &:= (\mathbf{D}_o^T \mathbf{Z}^{-1} \mathbf{D}_o)^{-1} > 0, \\ b &:= \mathbf{D}_o^T \mathbf{Z}_o^{-1} \mathbf{V}_u \mathbf{V}^{-1} \mathbf{D}_o = \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o. \end{aligned}$$

If $\mathbf{Z} = \mathbf{V}^{-1}$ then

$$c := \mathbf{D}_o^T \mathbf{V}^{-1} \text{Var} [\mathbf{U} | \mathbf{D}_o] \mathbf{V}^{-1} \mathbf{D}_o > 0$$

and if $\mathbf{Z} = \mathbf{V}_t^{-1}$ then

$$c := \mathbf{D}_o^T \mathbf{V}_t^{-1} (\text{Var} [\mathbf{U} | \mathbf{D}_o] + \mathbf{V}_e) \mathbf{V}_t^{-1} \mathbf{D}_o > 0.$$

The corrected estimator has a smaller mean square error if the following condition, derived from combining the univariate versions of equations (6.3) and (6.4), holds,

$$a^2 b^2 \beta^2 + (a + a^2 c \beta^2) \left(1 - \frac{1}{(1-ab)^2}\right) > 0. \quad (6.5)$$

Notice that $(1-ab)$ is the bias coefficient in equation (3.6). We can immediately see that if $b = 0$ (this will be when we have *e.g.* no measurement error) then both estimators have equal mean square errors. The other conclusion is that if $ab \in (-\infty, 0) \cup [2, \infty)$ then the corrected estimator is better which implies that for all $b < 0$ it is better to correct. This is because for these values the variance of the corrected estimator is smaller than that of the uncorrected one. When $ab \in (0, 2)$ we have a trade-off between the bias and variance. The condition can be reformulated as,

$$a^2 (b^2 + c \frac{(1-ab)^2 - 1}{(1-ab)^2}) \beta^2 + a (1 - \frac{1}{(1-ab)^2}) > 0,$$

and depending on the sign of $a^2 (b^2 + c \frac{(1-ab)^2 - 1}{(1-ab)^2})$ (when $0 < ab < 2$), we will have one of the two situations illustrated in Figure 2. Notice that the formula is symmetric around 0 in β . We can see that if $0 < ab < 2$ the magnitude of c will determine whether it is worth correcting or not for a given region of β s. If c is large then the coefficient will be negative and the second degree polynomial will always be negative so it is not worth correcting. If c is small then it will be worth to correct for large β s while not for small ones. We can see that if $0 < ab < 2$ then for $\beta = 0$ the biased estimator is better.

We can rewrite the mean square error difference condition (6.5) in terms of the bias coefficient $(1 - ab)$ which we will label as $\kappa_x := 1 - ab$. The condition then becomes,

$$(\kappa_x - 1)(\kappa_x^3 - \kappa_x^2 + (\frac{a}{\beta^2} + a^2 c)\kappa_x + \frac{a}{\beta^2} + a^2 c) > 0 \quad (6.6)$$

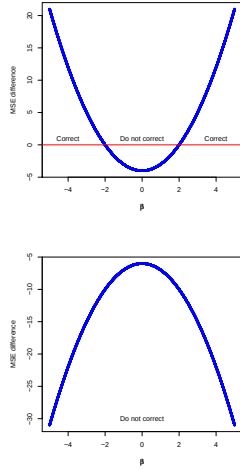


Figure 2: Depending whether the second degree polynomial is positive or negative when $0 < ab < 2$ for a given β it might be beneficial to correct for the bias or not, top: sign of coefficient in front of β^2 positive, bottom: negative.

which can be also written as,

$$(\kappa_x - 1)((a^2c(\kappa_x + 1) + \kappa_x^2(\kappa_x - 1))\beta^2 + a(\kappa_x + 1)) > 0.$$

We know from previous calculations when $0 < ab < 2$ then there is a potential trade – off between the variance and the bias. In terms of κ_x this is $-1 < \kappa_x < 1$. We can also see that when $\kappa_x = 0$ then it is always better not to correct and when $\kappa_x = -1$ it is always better to correct. This means that equation (6.6) has to have a root between -1 and 0 (another root is $\kappa_x = 1$). It remains to study whether it has any other real roots. If we have that $a^2c > \kappa_x^2(1 - \kappa_x)/(\kappa_x + 1)$ then for $\kappa_x \in (0, 1)$ the function will always be negative. This will happen if $a^2c > 0.09017$ as $\kappa_x^2(1 - \kappa_x)/(\kappa_x + 1)$ when $\kappa_x \in (0, 1)$ attains a maximum at $\frac{\sqrt{5}-1}{2}$ equal to just below 0.09017 . As we know that we have to have a root between $(-1, 0)$ let us assume that it is $-\kappa_0$ (so $\kappa_0 > 0$). Then we can write the condition polynomial as, $(\kappa_x - 1)(\kappa_x + \kappa_0)(\kappa_x^2 - (1 + \kappa_0)\kappa_x + (a/\beta^2 + a^2c)/\kappa_0)$. If it has two more real roots they both have to be positive (and therefore between 0 and 1 as 1 will not be a root of the equation), as they will be,

$$\kappa_{x1,2} = \frac{1}{2} \left((1 + \kappa_0) \pm \sqrt{(1 + \kappa_0)^2 - 4 \frac{a/\beta^2 + a^2c}{\kappa_0}} \right),$$

The two possibilities for how the mean square error difference can look like are presented symbolically in Figure 3.

The above result can be extended further for other values of the $\mathbf{Z}$ matrix. The mean square error difference can be written in a general

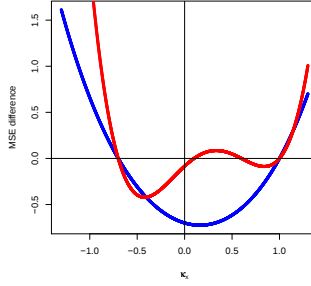


Figure 3: The two possibilities how the mean square error difference will be depending on the parameters.

form as,

$$(\kappa_x - 1) \left(\kappa_x^3 - \kappa_x^2 + \kappa_x \frac{\text{Var}[\hat{\beta}|\mathbf{D}_o]}{\beta_2^2} + \frac{\text{Var}[\hat{\beta}|\mathbf{D}_o]}{\beta_2^2} \right)$$

and as $\text{Var}[\hat{\beta}|\mathbf{D}_o] / \beta_2^2 > 0$ all of the previous discussion holds with just $a/\beta^2 + a^2c$ changed to $\text{Var}[\hat{\beta}|\mathbf{D}_o] / \beta_2^2$. The condition for which it will be better never to correct when $\kappa_x \in (0, 1)$ changes to

$\text{Var}[\hat{\beta}|\mathbf{D}_o] / \beta_2^2 > \kappa_x^2(1 - \kappa_x)/(\kappa_x + 1)$ and by the previous discussion this will be if $\text{Var}[\hat{\beta}|\mathbf{D}_o] / \beta_2^2 > 0.09017$.

If we make the simplifying assumption that $\mathbf{V}_u \mathbf{V}_o^{-1} = \kappa \mathbf{I}$, then $b = \kappa/a$ and the mean square error difference becomes,

$$(\kappa^2 - \frac{a^2 c \kappa (2 - \kappa)}{(1 - \kappa)^2}) \beta^2 - \frac{a \kappa (2 - \kappa)}{(1 - \kappa)^2} > 0,$$

we can immediately see that if the true slope is 0 then the biased estimator is superior in terms of mean square error (compare to the discussed later result of [19]). Because $0 < \kappa < 1$ the free term is always negative, so there will always be a region of small β s where the correction increases the mean square error. Therefore depending on the sign of the coefficient in front of β^2 it will be better sometimes to correct and sometimes not to correct. If it is negative then the biased estimator will be always better, if it is positive then for small β s it will be better not to correct, while for large ones it will be better.

The coefficient in front of β^2 can be negative if a (the estimate of the estimator's variance in the case of no measurement error) or c is large. This gives us a rule of thumb, if our estimator has large variance

it is better not to correct. However if we compare this to the situation, discussed previously, where our observations are dependent then this rule of thumb does not hold, the dependencies confound it.

The multipredictor case with independent observations is studied in detail in [19], here we just state the mean square error condition in our formulation. In the multipredictor case if we assume that $\mathbf{V}_u = \boldsymbol{\Sigma}_u \otimes \mathbf{I}$ and $\mathbf{V}_o = (\boldsymbol{\Sigma}_d + \boldsymbol{\Sigma}_u) \otimes \mathbf{I}$ (our observations of predictors are independent) and denoting,

$$\begin{aligned} \mathbf{A} &= (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \\ \mathbf{B} &= \boldsymbol{\Sigma}_u - \boldsymbol{\Sigma}_u (\boldsymbol{\Sigma}_d + \boldsymbol{\Sigma}_u)^{-1} \boldsymbol{\Sigma}_u \\ \mathbf{C} &= \mathbf{D}_o^T \mathbf{V}^{-2} \mathbf{D}_o \\ \mathbf{G} &= (\mathbf{I} - (\boldsymbol{\Sigma}_d + \boldsymbol{\Sigma}_u)^{-1} \boldsymbol{\Sigma}_u)^{-1}, \end{aligned}$$

we get that the difference between the mean square error of the uncorrected and corrected estimator is,

$$\begin{aligned} & \text{MSE}[\hat{\beta}|\mathbf{D}_o] - \text{MSE}[\hat{\beta}_{corr}|\mathbf{D}_o] \\ &= \text{Tr} \left(\mathbf{A} + \vec{\beta}^T \mathbf{B} \vec{\beta} \mathbf{A} \mathbf{C} \mathbf{A} - \mathbf{G} (\mathbf{A} + \vec{\beta}^T \mathbf{B} \vec{\beta} \mathbf{A} \mathbf{C} \mathbf{A}) \mathbf{G}^T \right) \\ &+ \vec{\beta}^T \boldsymbol{\Sigma}_u (\boldsymbol{\Sigma}_u + \boldsymbol{\Sigma}_d)^{-2} \boldsymbol{\Sigma}_u \vec{\beta}. \end{aligned}$$

The mean square error of the situation with independent observations is considered in [19]. The author states the result that the trace of the covariance matrix of the corrected estimator is greater than the trace of the uncorrected one and so there is a trade-off between the bias and variance in the mean square error. Notice that from the single predictor case discussion we can see that this result does not generalize to the case of dependent predictors. It is shown that if the sample size is large, the bias overwhelms the variance so then it is better to correct. Also the mean square error of estimating $\mathbf{T}^{-1} \vec{\beta}$ is considered, where $\mathbf{T} \mathbf{J} \mathbf{T}^{-1} = (\mathbf{I} - (\boldsymbol{\Sigma}_d + \boldsymbol{\Sigma}_u)^{-1} \boldsymbol{\Sigma}_u)$, where $\mathbf{J}$ is a diagonal eigenvalue matrix. If the true slope is zero then it is shown that the biased estimator of $\mathbf{T}^{-1} \vec{\beta}$ has smaller risk otherwise a condition is given which estimator to choose. In [16] the regression with a single predictor, intercept and independent observations is considered. The mean square error is shortly studied in relation to sample size, as the main topic of applications were survey studies.

6.5 Effect on intercept, fixed effects

The previous section considered the situation where $\mathbb{E}[\mathbf{D}_t] = \mathbb{E}[\mathbf{D}_o] = \mathbf{0}$, this in particular implies that there is no intercept in the model which is of course not realistic. Here we consider a more general case, where the observed design matrix $\mathbf{D}_o$ consists of fixed effects (*e.g.* an intercept), random effects observed without error and random effects observed with error.

We first consider the case where all random effects are observed with error. Let $\mathbf{D}_t = [\mathbf{F}; \mathbf{X}_t]$, where $\mathbf{F}$ is meant to denote the fixed effect part of the matrix and $\mathbf{X}_t$ the random part with measurement errors. Then $\mathbf{D}_o = [\mathbf{F}; \mathbf{X}_t] + [\mathbf{0}; \mathbf{U}] = [\mathbf{F}; \mathbf{X}_o]$ and the model is,

$$\vec{Y}_o = \mathbf{F}\vec{\beta}_F + (\mathbf{X}_o - \mathbf{U})\vec{\beta}_X + \vec{r},$$

we denote by $\vec{\beta} = [\vec{\beta}_F^T; \vec{\beta}_X^T]^T$. We need to calculate $E[\text{vec}([\mathbf{0}; \mathbf{U}]|\mathbf{D}_o)]$. Below we divide the matrices into blocks of appropriate sizes (relating to fixed and random effects),

$$E[\text{vec}([\mathbf{0}; \mathbf{U}]|\mathbf{D}_o)] = \begin{bmatrix} \mathbf{0} \\ \text{Cov}[\text{vec}(\mathbf{U}), \text{vec}(\mathbf{X}_o)] \text{Var}[\text{vec}(\mathbf{X}_o)]^{-1} \text{vec}(\mathbf{X}_o) \end{bmatrix}.$$

Using equation (3.1), we write $E[\hat{\vec{\beta}}|\mathbf{D}_o]$ as

$$\begin{aligned} & \left(\mathbf{I} - ([\mathbf{F}; \mathbf{X}_o]^T \mathbf{V}^{-1} [\mathbf{F}; \mathbf{X}_o])^{-1} [\mathbf{F}; \mathbf{X}_o]^T \mathbf{V}^{-1} \begin{bmatrix} \mathbf{0} \\ E[\mathbf{U}|\mathbf{D}_o] \end{bmatrix} \right) \vec{\beta} \\ &= \left(\mathbf{I} - \begin{bmatrix} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{F} & \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o \\ \vec{X}_o^T \mathbf{V}^{-1} \mathbf{F} & \vec{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{0} & \mathbf{F}^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] \\ \mathbf{0} & \vec{X}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] \end{bmatrix} \right) \vec{\beta} \\ &= \begin{bmatrix} \mathbf{I} & -\mathbf{A} \mathbf{F}^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] - \mathbf{B} \mathbf{X}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] \\ \mathbf{0} & -\mathbf{C} \mathbf{F}^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] - \mathbf{G} \mathbf{X}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] \end{bmatrix} \vec{\beta}, \end{aligned}$$

where

$$\begin{aligned} \mathbf{A} &= (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F} - \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o (\mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o)^{-1} \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F})^{-1}, \\ \mathbf{B} &= -(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F} - \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o (\mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o)^{-1} \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F})^{-1} \\ &\quad \times \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o (\mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o)^{-1}, \\ \mathbf{C} &= -(\mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o - \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F} (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o)^{-1} \\ &\quad \times \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F} (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1}, \\ \mathbf{G} &= (\mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o - \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F} (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o)^{-1}, \end{aligned}$$

coming from the formula for the inverse of a block matrix. So the final formula for $E[\hat{\vec{\beta}}|\mathbf{D}_o]$ is,

$$\begin{bmatrix} \vec{\beta}_F - (\mathbf{A} \mathbf{F}^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] + \mathbf{B} \mathbf{X}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o]) \vec{\beta}_X \\ (\mathbf{I} - (\mathbf{C} \mathbf{F}^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o] + \mathbf{G} \mathbf{X}_o^T \mathbf{V}^{-1} E[\mathbf{U}|\mathbf{D}_o])) \vec{\beta}_X \end{bmatrix}. \quad (6.7)$$

We can see that there is an influence on the estimate of the fixed effects from the measurement error and the fixed effects coefficient does not have any influence on the bias of the random effects estimates (compare with example 7 in [1]). Equation (6.7) additionally shows how to correct for the bias in the fixed effects. We can ask what assumptions are needed for formula (6.7) to simplify.

First we will consider whether it is possible to transform this equation so that the bias of the coefficients of the random effects is of the same functional form as the one in equation (3.1). It can be shown that in general these biases will not be equal but if we consider the following transformation of the random part of the design matrix $\mathbf{X}_o^* := \mathbf{X}_o - \mathbf{F}(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o$ (notice that $\mathbb{E}[\mathbf{X}_o^*] = \mathbf{0}$ which is crucial for the derivations) then, $\mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o^* = \mathbf{0}$ which gives that $\mathbb{E} \left[\begin{bmatrix} \hat{\beta}_F^{*T} & \hat{\beta}_X^{*T} \end{bmatrix}^T | \mathbf{D}_o \right]$ is,

$$\begin{bmatrix} \vec{\beta}_F^* - (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbb{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \\ (\mathbf{I} - (\mathbf{X}_o^{*T} \mathbf{V}^{-1} \mathbf{X}_o^*)^{-1} \mathbf{X}_o^{*T} \mathbf{V}^{-1} \mathbb{E}[\mathbf{U} | \mathbf{D}_o]) \vec{\beta}_X \end{bmatrix},$$

where $\vec{\beta}_F^* := \vec{\beta}_F + (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o \vec{\beta}_X$.

In section 6.6 we will look at what assumptions are needed for the bias to have a simplified formula. If we follow the same chain of reasoning for the fixed effects to be unbiased we require that $\mathbb{E}[\mathbf{U} | \mathbf{D}_o] = \mathbf{X}_o \mathbf{K}$ for some matrix $\mathbf{K}$ (the situation when this will happen was discussed in section 6.1). If this is the case then we will have the fixed effects unbiased and the reliability matrix for the random effects,

$$\mathbb{E} \left[\hat{\vec{\beta}} | \mathbf{D}_o \right] = \begin{bmatrix} \vec{\beta}_F \\ (\mathbf{I} - \mathbf{K}) \vec{\beta}_X \end{bmatrix}.$$

6.5.1 Fixed effects, random effects measured with and without error

The above derivations did not consider the situation that there could be random effects measured without error. We now denote $\mathbf{D}_t = [\mathbf{F}; \mathbf{X}_c; \mathbf{X}_t]$ and $\mathbf{D}_o = \mathbf{D}_t + [\mathbf{0}; \mathbf{0}; \mathbf{U}]$, where $\mathbf{X}_c$ are the random effects measured without error. The random effects covariance matrix is $\mathbf{V}_o = \begin{bmatrix} \mathbf{V}_c & \mathbf{V}_{cx} \\ \mathbf{V}_{cx}^T & \mathbf{V}_x \end{bmatrix}$ (where $\mathbf{V}_x = \text{Var}[\text{vec}(\mathbf{X}_t)] + \mathbf{V}_u$) and the regression model is,

$$\vec{Y}_o = \mathbf{F} \vec{\beta}_F + \mathbf{X}_c \vec{\beta}_c + (\mathbf{X}_o - \mathbf{U}) \vec{\beta}_X + \vec{r}.$$

As before we need $\mathbb{E}[\mathbf{U} | \mathbf{D}_o]$, which can be calculated as,

$$\text{vec}^{-1}(\mathbf{V}_u \text{Var}[\text{vec}(\mathbf{X}_o) | \mathbf{X}_c]^{-1} (\text{vec}(\mathbf{X}_o) - \mathbf{V}_{cx} \mathbf{V}_c^{-1} \text{vec}(\mathbf{X}_c))),$$

and this results in the following formula for the bias,

$$\mathbb{E} \left[\hat{\vec{\beta}} | \mathbf{D}_o \right] = \begin{bmatrix} \vec{\beta}_F - (\mathbf{A} \mathbf{F}^T + \mathbf{B} \mathbf{X}_c^T + \mathbf{C} \mathbf{X}_o^T) \mathbf{V}^{-1} \mathbb{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \\ \vec{\beta}_c - (\mathbf{G} \mathbf{F}^T + \mathbf{H} \mathbf{X}_c^T + \mathbf{L} \mathbf{X}_o^T) \mathbf{V}^{-1} \mathbb{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \\ (\mathbf{I} - (\mathbf{P} \mathbf{F}^T + \mathbf{R} \mathbf{X}_c^T + \mathbf{Q} \mathbf{X}_o^T) \mathbf{V}^{-1} \mathbb{E}[\mathbf{U} | \mathbf{D}_o]) \vec{\beta}_X \end{bmatrix}, \quad (6.8)$$

where

$$\begin{bmatrix} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{F} & \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_c & \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o \\ \mathbf{X}_c^T \mathbf{V}^{-1} \mathbf{F} & \mathbf{X}_c^T \mathbf{V}^{-1} \mathbf{X}_c & \mathbf{X}_c^T \mathbf{V}^{-1} \mathbf{X}_o \\ \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{F} & \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_c & \mathbf{X}_o^T \mathbf{V}^{-1} \mathbf{X}_o \end{bmatrix}^{-1} = \begin{bmatrix} \mathbf{A} & \mathbf{B} & \mathbf{C} \\ \mathbf{G} & \mathbf{H} & \mathbf{L} \\ \mathbf{P} & \mathbf{R} & \mathbf{Q} \end{bmatrix}.$$

We can see the same property that the coefficients of the correctly measured effects do not influence the bias. It is only the coefficient corresponding to the predictors measured with error that causes a bias in all the other coefficients and this bias can have any size and sign. One can compare this to the examples on pages 109 – 113 in [9] and we have the same conclusion that “*measurement error can bias the coefficients that are measured without error*” [9].

We will do the same kind of formula transformation as in the previous section. We perform the following transformation of the random part of the design matrix $\mathbf{X}_c^* := \mathbf{X}_c - \mathbf{F}(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_c$, $\mathbf{X}_o^* := \mathbf{X}_o - \mathbf{F}(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o$ (notice that $\mathbf{E}[\mathbf{X}_c^*] = \mathbf{0}$, $\mathbf{E}[\mathbf{X}_o^*] = \mathbf{0}$ still which again is crucial for the derivations) then, $\mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_c^* = \mathbf{0}$ and $\mathbf{F}^T \mathbf{V}^{-1} \mathbf{X}_o^* = \mathbf{0}$ which give that $\mathbf{E} \left[[\hat{\beta}_F^{*T} \ \hat{\beta}_c^{*T} \ \hat{\beta}_X^{*T}]^T | \mathbf{D}_o \right]$ equals,

$$\begin{bmatrix} \vec{\beta}_F^* - (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \\ \vec{\beta}_c - (\mathbf{A} \mathbf{X}_c^{*T} + \mathbf{B} \mathbf{X}_o^{*T}) \mathbf{V}^{-1} \mathbf{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \\ (\mathbf{I} - \mathbf{C} \mathbf{X}_c^{*T} + \mathbf{G} \mathbf{X}_o^{*T}) \mathbf{V}^{-1} \mathbf{E}[\mathbf{U} | \mathbf{D}_o] \vec{\beta}_X \end{bmatrix},$$

where $\vec{\beta}_F^* := \vec{\beta}_F + (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} (\mathbf{X}_c \vec{\beta}_c + \mathbf{X}_o \vec{\beta}_X)$ and

$$\begin{bmatrix} \mathbf{X}_c^{*T} \mathbf{V}^{-1} \mathbf{X}_c^* & \mathbf{X}_c^{*T} \mathbf{V}^{-1} \mathbf{X}_o^* \\ \mathbf{X}_o^{*T} \mathbf{V}^{-1} \mathbf{X}_c^* & \mathbf{X}_o^{*T} \mathbf{V}^{-1} \mathbf{X}_o^* \end{bmatrix}^{-1} = \begin{bmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{C} & \mathbf{G} \end{bmatrix}.$$

6.5.2 Single predictor

The single predictor with intercept model can be written as,

$$\vec{Y}_o = \vec{1} \beta_1 + (\vec{X}_o - \mathbf{U}) \beta_2 + \vec{r},$$

where $\vec{1}$ denotes of vector of 1s of length n . Since $\mathbf{E}[\mathbf{U} | \mathbf{D}_o] = \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o$ (where $\text{Var}[\vec{X}_o] = \mathbf{V}_o$), we have,

$$\mathbf{E}[\hat{\beta}_1 \ \hat{\beta}_2^T | \mathbf{D}_o] = \begin{bmatrix} \beta_1 - \kappa_1 \beta_2 \\ (1 - \kappa_2) \beta_2 \end{bmatrix},$$

where

$$\begin{aligned} \kappa_1 &= \frac{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\vec{1}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o) - (\vec{X}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)(\vec{1}^T \mathbf{V}^{-1} \vec{X}_o)}{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\vec{1}^T \mathbf{V}^{-1} \vec{1}) - (\vec{1}^T \mathbf{V}^{-1} \vec{X}_o)^2}, \\ \kappa_2 &= \frac{(\vec{X}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)(\vec{1}^T \mathbf{V}^{-1} \vec{1}) - (\vec{X}_o^T \mathbf{V}^{-1} \vec{1})(\vec{1}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)}{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\vec{1}^T \mathbf{V}^{-1} \vec{1}) - (\vec{1}^T \mathbf{V}^{-1} \vec{X}_o)^2}. \end{aligned}$$

We can ask as in the general case how does the coefficient in front of the slope, β_2 , $1 - \kappa_2$ compare to the corresponding one in equation (3.6). It can be shown that in the general dependent case (unless $\mathbf{V}_u = \kappa \mathbf{V}_o$) they are not equal (assuming that $\vec{X}_o = \mathbf{D}_o$). This is another difference caused by dependence between predictors. However one can as before transform the bias coefficient in front of β_2 to be of the same functional form as the one in equation (3.6). If we transform $\vec{X}_o$ to $\vec{X}_o^* := \vec{X}_o - (\vec{1}^T \mathbf{V}^{-1} \vec{X}_o) / (\vec{1}^T \mathbf{V}^{-1} \vec{1}) \vec{1}$ then we still have a zero mean design matrix as $E[\vec{X}_o^*] = \mathbf{0}$ and $\vec{1}^T \mathbf{V}^{-1} \vec{X}_o^* = 0$. The model becomes,

$$\vec{Y}_o = \vec{1}(\beta_1 + \frac{\vec{1}^T \mathbf{V}^{-1} \vec{X}_o}{\vec{1}^T \mathbf{V}^{-1} \vec{1}}) + (\vec{X}_o^* - \mathbf{U})\beta_2 + \vec{r}.$$

Defining $\beta_1^* := \beta_1 + (\vec{1}^T \mathbf{V}^{-1} \vec{X}_o) / (\vec{1}^T \mathbf{V}^{-1} \vec{1})$, we get,

$$E \left[\begin{bmatrix} \hat{\beta}_1^* \\ \hat{\beta}_2 \end{bmatrix} | \mathbf{D}_o \right] = \begin{bmatrix} \beta_1^* - (\vec{1}^T \mathbf{V}^{-1} \vec{1})^{-1} \vec{1}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o \beta_2 \\ (1 - (\vec{X}_o^{*T} \mathbf{V}^{-1} \vec{X}_o^*)^{-1} \vec{X}_o^{*T} \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o) \beta_2 \end{bmatrix}.$$

Notice that in the above formula we have both $\vec{X}_o$ and $\vec{X}_o^*$.

If we do not have an intercept but some general fixed effects vector then we write the model as,

$$\vec{Y}_o = \mathbf{F} \vec{\beta}_F + (\vec{X}_o - \mathbf{U})\beta_2 + \vec{r},$$

$E[\mathbf{U} | \mathbf{D}_o]$ does not change, as this only depends on the random effects and the formula ,

$$E \left[[\hat{\vec{\beta}}_F^T \hat{\beta}_2]^T | \mathbf{D}_o \right] = \begin{bmatrix} \vec{\beta}_F - \kappa_1 \beta_2 \\ (1 - \kappa_2) \beta_2 \end{bmatrix},$$

with

$$\begin{aligned} \kappa_1 &= \frac{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o) - (\vec{X}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)(\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o)}{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F}) - (\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o)^2}, \\ \kappa_2 &= \frac{(\vec{X}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F}) - (\vec{X}_o^T \mathbf{V}^{-1} \mathbf{F})(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)}{(\vec{X}_o^T \mathbf{V}^{-1} \vec{X}_o)(\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F}) - (\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o)^2}. \end{aligned}$$

We can do the same trick as with the intercept and instead of $\vec{X}_o$ consider, $\vec{X}_o^* := \vec{X}_o - (\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o) / (\mathbf{F}^T \mathbf{V}^{-1} \vec{1}) \vec{1}$. This has the effect that $\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o^* = 0$ but still $E[\vec{X}_o^*] = \mathbf{0}$. The formula then becomes,

$$E \left[\begin{bmatrix} \hat{\vec{\beta}}_F^* \\ \hat{\beta}_2 \end{bmatrix} | \mathbf{D}_o \right] = \begin{bmatrix} \vec{\beta}_F^* - (\mathbf{F}^T \mathbf{V}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o \beta_2 \\ (1 - (\vec{X}_o^{*T} \mathbf{V}^{-1} \vec{X}_o^*)^{-1} \vec{X}_o^{*T} \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o) \beta_2 \end{bmatrix},$$

where $\vec{\beta}_F^* := \vec{\beta}_F + (\mathbf{F}^T \mathbf{V}^{-1} \vec{X}_o) / (\mathbf{F}^T \mathbf{V}^{-1} \vec{1}) \vec{1}$ and so the slope's coefficient is of the same form as that in equation (3.6). Notice again that in the last formula we have both $\vec{X}_o$ and $\vec{X}_o^*$.

If $\mathbf{F}$ is not a vector but a fixed effects matrix (*i.e.* $\vec{\beta}_F$ is a vector and not a scalar) then again $E[\mathbf{U}|\mathbf{D}_o]$ does not change and the formula becomes,

$$\begin{aligned} E[\hat{\vec{\beta}}|\mathbf{D}_o] &= \begin{bmatrix} \vec{\beta}_F - \left(\mathbf{A}\mathbf{F}^T\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o + \mathbf{B}\vec{X}_o^T\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o \right) \beta_2 \\ \left(1 - (\mathbf{C}\mathbf{F}^T\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o + \mathbf{G}\vec{X}_o^T\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o) \right) \beta_2 \end{bmatrix} \\ &= \begin{bmatrix} \beta_1 - \mathbf{K}_1\beta_2 \\ (1 - \kappa_2)\beta_2 \end{bmatrix}, \end{aligned}$$

where, as previously

$$\begin{aligned} \mathbf{A} &= \left(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F} - \mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o(\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o)^{-1}\vec{X}_o^T\mathbf{V}^{-1}\mathbf{F} \right)^{-1}, \\ \mathbf{B} &= - \left(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F} - \mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o(\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o)^{-1}\vec{X}_o^T\mathbf{V}^{-1}\mathbf{F} \right)^{-1} \\ &\quad \times \mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o(\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o)^{-1}, \\ \mathbf{C} &= - \left(\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o - \vec{X}_o^T\mathbf{V}^{-1}\mathbf{F}(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o \right)^{-1} \\ &\quad \times \vec{X}_o^T\mathbf{V}^{-1}\mathbf{F}(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}, \\ \mathbf{G} &= \left(\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o - \vec{X}_o^T\mathbf{V}^{-1}\mathbf{F}(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o \right)^{-1}, \end{aligned}$$

We can find a simplification of the formula as before. We notice that after the study of the single predictor case it might be tempting to also consider $\vec{X}_o^* := \vec{X}_o - \mathbf{J}(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{J})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o$, where $\mathbf{J}$ is a matrix of ones, with number of rows equaling the number of observations and columns equaling the number of columns of $\mathbf{F}$. However one cannot use this, as $\mathbf{F}^T\mathbf{V}^{-1}\mathbf{J}$ is singular. We define as in the general case $\vec{X}_o^* := \vec{X}_o - \mathbf{F}(\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o$ and as before $\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o^* = \mathbf{0}$. This results in,

$$\begin{aligned} \mathbf{A} &= (\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}, \\ \mathbf{B} &= \mathbf{0}, \mathbf{C} = \mathbf{0}, \\ \mathbf{G} &= (\vec{X}_o^T\mathbf{V}^{-1}\vec{X}_o)^{-1} \end{aligned}$$

and the formula becomes,

$$E\left[\begin{bmatrix} \hat{\vec{\beta}}_F^* \\ \beta_2 \end{bmatrix} | \mathbf{D}_o\right] = \begin{bmatrix} \vec{\beta}_F^* - (\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o\beta_2 \\ \left(1 - (\vec{X}_o^{*T}\mathbf{V}^{-1}\vec{X}_o^*)^{-1}\vec{X}_o^{*T}\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1}\vec{X}_o \right) \beta_2 \end{bmatrix},$$

where $\vec{\beta}_F^* := \beta_F + (\mathbf{F}^T\mathbf{V}^{-1}\mathbf{F})^{-1}\mathbf{F}^T\mathbf{V}^{-1}\vec{X}_o$. This simplification depends on $E[\vec{X}_o] = \mathbf{0}$.

6.5.3 Mean square error analysis

If we consider the mean square error of both the intercept and slope then we get the following conditions using the results and notation from

section 6.3. The corrected estimator of the slope is better if,

$$(\kappa_x - 1)(\kappa_x^3 - \kappa_x^2 + (\frac{a}{\beta^2} + a^2c)\kappa_x + \frac{a}{\beta^2} + a^2c) > 0$$

which is the same condition, equation (6.6), as in the no intercept case. Thus all the results from the no intercept case can be carried over to the case when we have an intercept. The condition for the corrected estimator of the intercept to have a smaller mean square error is that

$$\begin{aligned} & (\tilde{\mathbf{I}}^T \mathbf{V}^{-1} \tilde{\mathbf{I}})^{-2} (\tilde{\mathbf{I}}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o)^2 [(1 - a^2c)\beta_2^2 - a] \\ & - 2 \frac{(\tilde{\mathbf{I}}^T \mathbf{V}^{-1} \tilde{\mathbf{I}})^{-2}}{a} \tilde{\mathbf{I}}^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \vec{X}_o \tilde{\mathbf{I}}^T \mathbf{V}^{-1} \text{Var} \left[\mathbf{U} | \vec{X}_o \right] \mathbf{V}^{-1} \vec{X}_o^* \beta_2 \end{aligned}$$

is positive. As we can analyze the estimators of the intercept and slope separately it might in some situations pay-off to use for one parameter the corrected estimator and for the other the uncorrected one.

In all of the discussion we assumed that we know all the covariance matrices. In practice this might of course not be the case. The only way to know the measurement error covariance matrix is to make multiple measurements. In comparative analysis studies this is natural as for each species we have a mean from a number of individuals. This will cause $\mathbf{V}_e$ and $\mathbf{V}_u$ to be diagonal or of the form $\Sigma \otimes \mathbf{I}$ (in the multivariate case) and we can use this in the estimation procedure.

6.6 Unconditional bias of GLS estimator with measurement error in predictors

The main problem with calculating the bias of the generalized least squares estimator is that one needs to be able to calculate the expectation of

$$(\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o$$

for arbitrary covariance matrices $\mathbf{V}$, $\mathbf{V}_u$ and $\mathbf{V}_o$. This is difficult even in the case of one predictor as the matrix product $\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1}$ seems to have few useful for us properties. We will now consider the very simple case when β is one dimensional, *i.e.* $\mathbf{D}_o$ is of dimension $n \times 1$ which means that we can write the expectation of the generalized least squares estimator as,

$$\mathbb{E} \left[\hat{\beta} \right] = \beta - \mathbb{E} \left[\frac{\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o}{\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o} \right] \beta$$

and we need to consider,

$$\mathbb{E} \left[\frac{\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} \mathbf{D}_o}{\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o} \right],$$

where $\mathbf{D}_o \sim \mathcal{N}(\vec{0}, \mathbf{V}_o)$. It can be easily shown that the matrix

$$\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} = \mathbf{V}^{-1} \mathbf{V}_u (\mathbf{V}_d + \mathbf{V}_u)^{-1} = \mathbf{V}^{-1} (\mathbf{V}_u^{-1} \mathbf{V}_d + \mathbf{I})^{-1}$$

has positive real eigenvalues.

Lemma 6.1. *If $\mathbf{B}$ and $\mathbf{C}$ are symmetric positive definite matrices then $\mathbf{A} = \mathbf{BC}$ has positive real eigenvalues.*

Proof $\mathbf{C} = \mathbf{B}^{-1}\mathbf{A}$ is by assumption symmetric positive definite (and so Hermitian) so for any complex $\mathbf{z}$, $\mathbf{z}^T \mathbf{B}^{-1} \mathbf{A} \mathbf{z} \in \mathbb{R}_+$. Assume that $\mathbf{A}$ has a complex negative real part eigenvalue $\lambda = (-a) + ib$ ($a > 0$) and let $\mathbf{z}$ be its (complex) eigenvector. Then

$$\begin{aligned} \mathbb{R}_+ \in \mathbf{z}^T \mathbf{C} \mathbf{z} &= \mathbf{z}^T \mathbf{B}^{-1} \mathbf{A} \mathbf{z} = \mathbf{z}^T \mathbf{B}^{-1} \lambda \mathbf{z} = \mathbf{z}^T \mathbf{B}^{-1} (-a + ib) \mathbf{z} \\ &= -a \mathbf{z}^T \mathbf{B}^{-1} \mathbf{z} + ib \mathbf{z}^T \mathbf{B}^{-1} \mathbf{z}, \end{aligned}$$

which is complex with negative real part as $\mathbf{z}^T \mathbf{B}^{-1} \mathbf{z} \in \mathbb{R}_+$ a contradiction.

Q.E.D.

Results on the products of symmetric positive definite matrices can be found in *e.g.* [3], [4], [5], [6] or [56]. However positive real eigenvalues are not sufficient for the matrix to be positive definite (see *e.g.* [32]) and we cannot have symmetry guaranteed as there is a product of symmetric matrices. The properties of positive definiteness and symmetry of $\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1}$ are needed as results for distributions of ratios of quadratic forms assume this. Let us therefore assume that $\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1}$ is symmetric positive definite and define

$$\begin{aligned} \mathbf{Q} &= (\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1} + (\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1})^T) / 2 \\ \mathbf{P} &= \mathbf{V}^{-1}, \end{aligned}$$

and then the expectation we are interested in equals

$$\mathbb{E} [(\mathbf{D}_o^T \mathbf{Q} \mathbf{D}_o) / (\mathbf{D}_o^T \mathbf{P} \mathbf{D}_o)],$$

where both $\mathbf{Q}$ and $\mathbf{P}$ are symmetric positive definite. Results on expectations of such ratios can be found in *e.g.* [23], [46], [47], [48] but they are rather complex and in most cases (even if $\mathbf{V}_u = \sigma^2 \mathbf{I}$) we cannot expect $\mathbf{V}^{-1} \mathbf{V}_u \mathbf{V}_o^{-1}$ to be symmetric positive definite.

We can also ask what conditions on $\mathbf{V}_u$ and $\mathbf{V}_o$ are necessary and sufficient for the bias to be of the form $\mathbb{E} [\hat{\beta} | \mathbf{D}_o] = (\mathbf{I} - \mathbf{K}) \vec{\beta}$ (a linear combination of the true parameters) for some matrix $\mathbf{K}$ independent of $\mathbf{D}_o$ (in section 6.3 we showed the condition for the single predictor case). From equation 3.1 we can see that we need $\text{vec}^{-1}(\mathbf{V}_u \mathbf{V}_o^{-1} \text{vec}(\mathbf{D}_o)) = \mathbf{D}_o \mathbf{K}$ for some appropriate $m \times m$ matrix $\mathbf{K}$ (as the inverse of a matrix is unique). This will be achieved for those covariance matrix $\mathbf{V}_u$, $\mathbf{V}_o$ pairs for which there exists an appropriate $m \times m$ matrix $\mathbf{K}$ such that $\mathbf{V}_u \mathbf{V}_o^{-1} = \mathbf{K}^T \otimes \mathbf{I}$. As $\mathbf{V}_o = \mathbf{V}_u + \mathbf{V}_d$ this gives $\mathbf{V}_d = \mathbf{V}_u (\mathbf{I} - \mathbf{K}^T \otimes \mathbf{I})$ which simplifies to

$$\mathbf{V}_u = \mathbf{V}_d ((\mathbf{K}^T + \mathbf{I})^{-1} \otimes \mathbf{I}) = ((\mathbf{K}^{-T} - \mathbf{I})^{-1} \otimes \mathbf{I}) \mathbf{V}_d.$$

The condition $\mathbf{V}_u \mathbf{V}_o^{-1} = \mathbf{K}^T \otimes \mathbf{I}$ can be understood in the following manner. The matrices $\mathbf{V}_u, \mathbf{V}_o^{-1}$ have full rank so their rows of $\mathbf{V}_u$ and columns of $\mathbf{V}_o^{-1}$ generate the vector space $\mathbb{R}^{nm}$. Each subset of the rows/columns generates some subspace, in particular consider the following subspaces,

$$\Psi_1 = \text{span}\{\mathbf{V}_u[(r-1)n+1,], r=1, \dots, m\}$$

$$\vdots$$

$$\Psi_h = \text{span}\{\mathbf{V}_u[(r-1)n+k,], r=1, \dots, m\}$$

$$\vdots$$

$$\Psi_n = \text{span}\{\mathbf{V}_u[(r-1)n+n,], r=1, \dots, m\}$$

and

$$\Phi_1 = \text{span}\{\mathbf{V}_o^{-1}[:, (r-1)n+1], r=1, \dots, m\}$$

$$\vdots$$

$$\Phi_h = \text{span}\{\mathbf{V}_o^{-1}[:, (r-1)n+k], r=1, \dots, m\}$$

$$\vdots$$

$$\Phi_n = \text{span}\{\mathbf{V}_o^{-1}[:, (r-1)n+n], r=1, \dots, m\},$$

where for a matrix $\mathbf{M}$, $\mathbf{M}[i,]$ means the i th row of $\mathbf{M}$ and $\mathbf{M}[:, j]$ means the j th column of $\mathbf{M}$. The first condition linking the matrices $\mathbf{V}_u$ and $\mathbf{V}_o$ is that for each $h=1, \dots, n$ $\Psi_h = \Phi_h$ and for each $i \neq j$ Ψ_i and Φ_j are orthogonal.

The second condition is that there exists a matrix $\mathbf{K} = [k_{ij}]$ of size $m \times m$ such that for each pair i and j we have for each $h=1, \dots, n$ that the dot product of the i th vector of the h th Ψ base and the j th vector of the h th Φ base is k_{ji} (so we are also assuming there is an ordering of base vectors). As the relationship between $\mathbf{V}_d$ and $\mathbf{V}_u$ is of the same form with the Kronecker product, the same interpretation is valid for linking $\mathbf{V}_d$ and $\mathbf{V}_u$. We used the wording that $\mathbf{K}$ must be appropriate. This is so that the matrices $\mathbf{V}_d$ and $\mathbf{V}_u$ will be covariance matrices. A necessary condition is in lemma 6.1, $\mathbf{K}$ must have positive real eigenvalues. The family of permissible $\mathbf{K}$ matrices in the case of independent observations is described in [1].

We can further ask what are the necessary and sufficient conditions on $\mathbf{V}_u$ and $\mathbf{V}_o$ for the bias to be of the form,

$\mathbb{E}[\hat{\beta}|\mathbf{D}_o] = (\mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{D}_o)^{-1} \mathbf{D}_o^T \mathbf{V}^{-1} \mathbf{H} \mathbf{D}_o \mathbf{K}$ for some appropriate matrices $\mathbf{H} \in \mathbb{R}^{n \times n}$ and $\mathbf{K} \in \mathbb{R}^{m \times m}$. The needed relation is $\mathbf{V}_u \mathbf{V}_o^{-1} = \mathbf{K}^T \otimes \mathbf{H}$ which using that $\mathbf{V}_o = \mathbf{V}_u + \mathbf{V}_d$ this becomes,

$$\mathbf{V}_u = \mathbf{V}_d (\mathbf{I} - \mathbf{K}^T \otimes \mathbf{H})^{-1} = (\mathbf{K}^{-T} \otimes \mathbf{H}^{-1} - \mathbf{I})^{-1} \mathbf{V}_d.$$

If we look at the bias in the single predictor case, from equation (3.6) we can see that for the bias to be of the “classical” form $\mathbb{E}[\hat{\beta}|\mathbf{D}_o] =$

$(1 - \kappa)\beta$ (where κ does not depend on $\mathbf{D}_o$), we need that (as $\mathbf{D}_o$ can take on any (random) values) $\mathbf{V}^{-1}\mathbf{V}_u\mathbf{V}_o^{-1} = \kappa\mathbf{V}^{-1}$ or $\mathbf{V}_u = \kappa\mathbf{V}_o$. Since $\mathbf{V}_o = \mathbf{V}_d + \mathbf{V}_u$, it follows

$$\mathbf{V}_u = \frac{\kappa}{1-\kappa}\mathbf{V}_d,$$

which means that the covariance of the errors in the predictors has to be proportional to the covariance between the predictors. Notice that for $\mathbf{V}_u, \mathbf{V}_d$ to be covariance matrices we need $0 < \kappa < 1$ and since $\kappa/(1 - \kappa)$ is bijective on $(0, 1) \rightarrow \mathbb{R}_+$, in this setting the estimate of β will always be biased downwards. Notice that the multivariate formulae are a generalization of this.

7 Supplementary material B: Stochastic differential equations for the Ornstein–Uhlenbeck model (Paper II)

7.1 Convergence of stochastic process

One of the most important properties of the considered processes is whether the processes converge or not. Convergence of a process can be interpreted as adaptation. Below I review the definitions (see *e.g.* [22] or [55]) and the most basic convergence theorems.

Definition 7.1 (a.s. convergence). A stochastic process $\{X(t)\}$ converges *almost surely* to a random variable X if $P(\{\omega : X(t, \omega) \rightarrow X(\omega)\}) = 1$.

Definition 7.2 (convergence in probability). A stochastic process $\{X(t)\}$ converges *in probability* to a random variable X if for any $\epsilon > 0$ $P(\{\omega : |X(t, \omega) - X(\omega)| > \epsilon\}) \rightarrow 0$ as $t \rightarrow \infty$.

Definition 7.3 (L^p convergence). A stochastic process $\{X(t)\}$ converges in L^p to a random variable X if $E[|X(t)|^p] < \infty$ for all t and $E[|X(t) - X|^p] \rightarrow 0$ as $t \rightarrow \infty$.

Definition 7.4 (weak convergence). A stochastic process $\{X(t)\}$ converges *weakly* (or *in distribution*) to a random variable X if the distribution functions $F_{X(t)}(x)$ converge to the distribution function of X , $F_X(x)$ at every point x of continuity of F_X .

The convergences are related as almost sure convergence implies convergence in probability and convergence in probability implies weak convergence. L^p convergence implies convergence in probability. The implications do not hold the other way in general.

Theorem 7.5 ([54]). *For every t let $E[|\zeta_t|^2] < \infty$. It is sufficient and necessary for $\lim \zeta_t$ when $t \rightarrow t_0$ (or $t \rightarrow \infty$) to exist in L^2 that $E[\zeta_t \zeta_s]$ when $t, s \rightarrow t_0$ (or $t, s \rightarrow \infty$) exists and is finite.*

Theorem 7.6 ([54]). *It is necessary and sufficient for the limit of ζ_t to exist in probability for $t \rightarrow \infty$ that the two dimensional distributions $\Psi_{\zeta_t \zeta_s}$ converge weakly as $t, s \rightarrow \infty$.*

Theorem 7.7 (Lévy–Cramér continuity theorem [22] p.190). *The following are equivalent,*

- F_n converges weakly to F
- If $\phi_{F_n}(t)$ and $\phi_F(t)$ are respectively the characteristic functions of F_n and F and for every $t \in \mathbb{R}$ we have $\lim_{n \rightarrow \infty} \phi_{F_n}(t) = \phi_F(t)$.

i.e. *pointwise convergence of characteristic functions is equivalent to weak convergence.*

In the case when we have normal distributions they are characterized by their mean value and covariance structure. Therefore if we have convergence of mean vectors and covariance matrices this will result in convergence of their characteristic functions ($\phi_X(\vec{t}) = \exp(i\vec{\mu}\vec{t} - \vec{t}^T \Sigma \vec{t}/2)$ [55]) implying weak convergence. When we have a stochastic process with continuous time instead of discrete then if the mean vector and covariance converges this will imply convergence for every subsequence. In turn this will give us for every subsequence pointwise convergence of characteristic functions and so for every subsequence weak convergence of measures which by definition gives us weak convergence ($\lim_{s \rightarrow t} F_s(x) = F_t(x)$).

7.2 Matrix exponential and related integral

In most of the models to be able to calculate the covariance structure we need to be able to calculate the matrix integral function

$$\int_0^t e^{\mathbf{F}v} \mathbf{\Psi} e^{\mathbf{F}^T v} dv, \quad (7.1)$$

for a matrix $\mathbf{\Psi}$. Given an eigendecomposition $\mathbf{F} = \mathbf{P}\mathbf{\Lambda}\mathbf{P}^{-1}$, where $\mathbf{\Lambda}$ is the diagonal matrix of eigenvalues and $\mathbf{P}$ is the invertible matrix of eigenvectors, we have

$$e^{-t\mathbf{F}} = \mathbf{P} \begin{bmatrix} e^{-t\lambda_1} & \dots & 0 \\ \vdots & \ddots & \vdots \\ \dots & 0 & e^{-t\lambda_{k_Y}} \end{bmatrix} \mathbf{P}^{-1}$$

and therefore

$$\int_0^t e^{v\mathbf{F}} \mathbf{\Psi} e^{v\mathbf{F}^T} dv = \mathbf{P} \int_0^t [e^{v(\lambda_i + \lambda_j)}]_{1 \leq i, j \leq k} * \mathbf{\Upsilon} dv \mathbf{P}^T$$

where $\mathbf{\Upsilon} := \mathbf{P}^{-1} \mathbf{\Psi} \mathbf{P}^{-T}$. Suppose there are exactly $l \geq 0$ zero-valued eigenvalues and $(\lambda_1, \dots, \lambda_{k_Y}) = (\lambda_1, \dots, \lambda_{k_Y-l}, 0, \dots, 0)$. Then (7.1) can be obtained as

$$\mathbf{P} \left[\begin{array}{c|c} \left[\frac{1}{\lambda_i + \lambda_j} (e^{t(\lambda_i + \lambda_j)} - 1) \cdot \Upsilon_{ij} \right]_{1 \leq i, j \leq k-l} & \left[\frac{1}{\lambda_i} (e^{t\lambda_i} - 1) \cdot \Upsilon_{ij} \right]_{\substack{1 \leq i \leq k-l \\ (k-l+1) \leq j \leq k}} \\ \hline \left[\frac{1}{\lambda_j} (e^{t\lambda_j} - 1) \cdot \Upsilon_{ij} \right]_{\substack{(k-l+1) \leq i \leq k-l \\ 1 \leq j \leq k-l}} & t [\Upsilon_{ij}]_{(k-l+1) \leq i, j \leq k} \end{array} \right] \mathbf{P}^T.$$

This might not be the most effective way of calculating this integral as in hinges on the precision of calculating $\mathbf{F}$'s eigendecomposition. The papers [35], [52], [53] consider calculating the matrix exponential and related problems.

7.3 Ornstein–Uhlenbeck model

In [10] a special version of the Ornstein–Uhlenbeck model, equation (4.1), is presented (with software), where $\vec{\Psi}(t)$ is a step function over the phylogeny and $\mathbf{F}$ is a symmetric positive definite matrix (so a positive scalar in the one dimensional case). This restriction on $\mathbf{F}$ results in ease of interpretability and simple parametrization for numerical optimization of the likelihood function (see section 8). However this restriction is not necessary and using the eigenvalue decomposition of $\mathbf{F} = \mathbf{P}\mathbf{\Lambda}\mathbf{P}^{-1}$ ($\mathbf{P}\mathbf{\Lambda}\mathbf{P}^T$ for $\mathbf{F}$ symmetric positive definite) we can work with arbitrary $\mathbf{F}$ matrices.

In Paper II we present a package that in fact allows for an arbitrary invertible drift matrix using its eigendecomposition. The traits evolve according to equation (4.1) whose solution is (see [12] or [17]),

$$\vec{Z}(t) = e^{-\mathbf{F}t} \vec{Z}(0) + \int_0^t e^{-\mathbf{F}(t-v)} \mathbf{F} \vec{\Psi} dv + \int_0^t e^{-\mathbf{F}(t-v)} \mathbf{\Sigma} d\mathbf{W}(v). \quad (7.2)$$

Using that $\frac{d}{dt} e^{\mathbf{F}t} = e^{\mathbf{F}t} \mathbf{F} = \mathbf{F} e^{\mathbf{F}t}$ and assuming $\vec{\Psi}(t) \equiv \vec{\Psi}$ is constant this can be written as,

$$\begin{aligned} \vec{Z}(t) &= e^{-\mathbf{F}t} \vec{Z}(0) + e^{-\mathbf{F}t} (e^{\mathbf{F}t} - \mathbf{I}) \vec{\Psi} + \int_0^t e^{-\mathbf{F}(t-v)} \mathbf{\Sigma} d\mathbf{W}(v) \\ &= \vec{\Psi} + e^{-\mathbf{F}t} (\vec{Z}(0) - \vec{\Psi}) + \int_0^t e^{-\mathbf{F}(t-v)} \mathbf{\Sigma} d\mathbf{W}(v), \end{aligned}$$

If $\vec{\Psi}(t)$ is a step function such that on the intervals $(t_{j-1}, t_j]$, $\vec{\Psi}(t)$ takes values $\vec{\Psi}_j$, then

$$\vec{Z}(t) = e^{-\mathbf{F}t} \left(\vec{Z}(0) + \sum_{j=1}^m (e^{t_j \mathbf{F}} - e^{t_{j-1} \mathbf{F}}) \vec{\Psi}_j \right) + \int_0^t e^{-\mathbf{F}(t-v)} \mathbf{\Sigma} d\mathbf{W}(v).$$

As this is a multivariate normal process we need to calculate the mean and covariance functions. They will be, conditional on the initial value $\vec{Z}(0)$,

$$\mathbb{E} [\vec{Z}(t)] = e^{-\mathbf{F}t} \left(\vec{Z}(0) + \sum_{j=1}^m (e^{t_j \mathbf{F}} - e^{t_{j-1} \mathbf{F}}) \vec{\Psi}_j \right)$$

and

$$\text{Var} [\vec{Z}(t)] = e^{-\mathbf{F}t} \left(\int_0^t e^{\mathbf{F}v} \mathbf{\Sigma} \mathbf{\Sigma}^T e^{\mathbf{F}^T v} dv \right) e^{-\mathbf{A}^T t}. \quad (7.3)$$

The second moment of the process can be calculated to be

$$\mathbf{P} \left(\left[\frac{1}{\lambda_i + \lambda_j} (1 - e^{-(\lambda_i + \lambda_j)t}) \right]_{1 \leq i, j \leq k} * (\mathbf{P}^{-1} \mathbf{\Sigma} \mathbf{\Sigma}^T \mathbf{P}^T) \right) \mathbf{P}^T,$$

where k denotes the dimension of $\mathbf{F}$.

We are also interested in the distribution of the data on the phylogeny. To do this we need to calculate $\text{Cov} [\vec{Z}_i, \vec{Z}_j]$ for all species i and j . We use the identity from [25]

$$\text{Cov} [\vec{Z}_i, \vec{Z}_j] = \text{Cov} \left[\mathbb{E} [\vec{Z}_i | \vec{Z}_{a_{ij}}], \mathbb{E} [\vec{Z}_j | \vec{Z}_{a_{ij}}] \right],$$

where $\vec{Z}_{a_{ij}}$ is the value of the trait vector of the most recent common ancestor of species i and j . This will result in,

$$\text{Cov} [\vec{Z}_i, \vec{Z}_j] = e^{-\mathbf{F}(t_i - t_{a_{ij}})} \text{Var} [\vec{Z}_a] e^{-\mathbf{F}^T(t_j - t_{a_{ij}})}. \quad (7.4)$$

Depending on whether $\mathbf{F}$ has positive real part eigenvalues or not we will get different asymptotic properties of the process.

7.3.1 Asymptotic properties of unitrait Ornstein–Uhlenbeck model

Consider the single trait case we have, assuming that $\Psi(t) \equiv \Psi$ is constant and let $\alpha > 0$.

$$\begin{aligned} dZ(t) &= -\alpha(Z(t) - \Psi)dt + \sigma dW \\ Z(t) &= \Psi + (Z(0) - \Psi)e^{-\alpha t} + \sigma \int_0^t e^{-\alpha(t-v)} dW(v). \end{aligned}$$

We see that

$$\begin{aligned} \mathbb{E}[Z(t)] &= \Psi + (Z(0) - \Psi)e^{-\alpha t} \rightarrow \Psi \\ \text{Var}[Z(t)] &= \frac{\sigma^2}{2\alpha}(1 - e^{-2\alpha t}) \rightarrow \frac{\sigma^2}{2\alpha}. \end{aligned}$$

It follows that $Z(t)$ converges weakly to a normal random variable $\mathcal{N}(\Psi, \frac{\sigma^2}{2\alpha})$. We can ask whether it will converge in probability. Observe that

$$\text{Cov}[Z(t), Z(s)] = \frac{\sigma^2}{2\alpha}(e^{-\alpha|t-s|} - e^{-\alpha(t+s)}) = \frac{\sigma^2}{2\alpha}e^{-\alpha|t-s|} + o(1),$$

by direct calculation or using equation (7.4). According to theorem 7.6 we do not have convergence for $t, s \rightarrow \infty$ as the limit depends on $|t - s|$ so we will be able to choose subsequences that will give different limits. Therefore the process does not converge in probability nor almost surely, only weakly. An intuition what this means can be that the process can have spikes but these will last for very short periods of time.

The limiting distribution is the stationary distribution of the process, this can be seen if we assume that $Z(0) \sim \mathcal{N}(\Psi, \frac{\sigma^2}{2\alpha})$. We then have

$$\begin{aligned} \mathbb{E}[Z(t)] &= \Psi + (\mathbb{E}[Z(0)] - \Psi)e^{-\alpha t} = \Psi + (\Psi - \Psi)e^{-\alpha t} = \Psi \\ \text{Var}[Z(t)] &= e^{-2\alpha t} \text{Var}[Z(0)] + \frac{\sigma^2}{2\alpha}(1 - e^{-2\alpha t}) \\ &= e^{-2\alpha t} \frac{\sigma^2}{2\alpha} + \frac{\sigma^2}{2\alpha}(1 - e^{-2\alpha t}) = \frac{\sigma^2}{2\alpha}. \end{aligned}$$

If $\alpha \leq 0$ then no stationary distribution exists, nor do we have weak convergence (and if $\alpha = 0$ the process becomes a Brownian motion).

7.3.2 Asymptotic properties of multitrait Ornstein–Uhlenbeck model

The distribution of the vector of traits $\vec{Z}$ at time t is always multivariate normal. We can ask what does this tend to as $t \rightarrow \infty$. The distribution of $\vec{Z}$ at $t = \infty$ (assuming $\mathbf{F}$ has positive real part eigenvalues) will be multivariate normal (as the characteristic functions will converge). The mean value of $\vec{Z}$ will tend to (assuming $\vec{\Psi}(t) \equiv \vec{\Psi}$ is constant from some time point),

$$\lim_{t \rightarrow \infty} \mathbf{E} [\vec{Z}(t)] = \vec{\Psi}$$

and the covariance structure

$$\begin{aligned} \lim_{t \rightarrow \infty} \text{Var} [\vec{Z}(t)] &= \int_0^\infty e^{-\mathbf{F}(t-v)} \mathbf{\Sigma} \mathbf{\Sigma}^T e^{-\mathbf{F}^T(t-v)} dv \\ &= \mathbf{P} \left(\left[\frac{1}{\lambda_i + \lambda_j} \right]_{1 \leq i, j \leq k_Y} * (\mathbf{P}^{-1} \mathbf{\Sigma} \mathbf{\Sigma}^T \mathbf{P}^{-T}) \right) \mathbf{P}^T. \end{aligned}$$

This gives us that the process describing the species' evolution on the tree tends, in distribution, to a multivariate normal with the above mean and covariance structure. What is obvious from equation (7.4) is that the covariance (and in turn correlation) between two traits in different species will be tending to 0 as $t \rightarrow \infty$ (as the covariance at $t_{a_{ij}}$ is fixed).

7.4 Ornstein–Uhlenbeck Brownian Motion model

In [26] (with further developments in [33], [42]) a very special type of the Ornstein–Uhlenbeck model (introduced in [34]) was presented. It assumed that the drift matrix has only non-zero values on the first row and the diffusion matrix is diagonal, this implies that the variables “assigned” to the other rows evolve marginally as independent Brownian motions. We denote by $y(t)$ the first variable and the remaining ones by $x_1(t), \dots, x_{k_x}(t)$ which results in the equation,

$$\begin{aligned} dy(t) &= -\alpha[y(t) - b_0 - \sum_{l=1}^{k_x} b_l x_l(t)]dt + \sigma_y dW_y, \\ dx_l(t) &= \sigma_{x_l} dW_{x_l}, \quad 1 \leq l \leq k_x, \end{aligned}$$

Due to the layering structure this model is termed the Ornstein–Uhlenbeck Brownian Motion (OUBM) model. In this particular case the solution equation (7.2) takes the form,

$$\begin{aligned} y(t) &= e^{-\alpha t} y(0) + (1 - e^{-\alpha t})(b_0 + \sum_{l=1}^{k_x} b_l x_l(0)) \\ &\quad + \int_0^t e^{-\alpha(t-s)} \sigma_y dW_y(s) + \sum_{l=1}^{k_x} \int_0^t (1 - e^{-\alpha(t-s)}) b_l \sigma_{x_l} dW_{x_l}(s), \\ x_l(t) &= x_l(0) + \int_0^t \sigma_{x_l} dW_{x_l}(s), \quad 1 \leq l \leq k_x. \end{aligned}$$

The corresponding moments are,

$$\begin{aligned}
\mathbb{E}[y(t)] &= e^{-\alpha t} y(0) + (1 - e^{-\alpha t})(b_0 + \sum_{l=1}^{k_x} b_l x_l(0)) \\
\mathbb{E}[x_l(t)] &= x_l(0) \quad 1 \leq l \leq k_x \\
\text{Var}[y(t)] &= \frac{1-e^{-2\alpha t}}{2\alpha} \sigma_y^2 + \left(t + \frac{1-e^{-2\alpha t}}{2\alpha} - \frac{2}{\alpha}(1 - e^{-\alpha t}) \right) \sum_{l=1}^{k_x} \sigma_{x_l}^2 b_l^2 \\
\text{Var}[x_l(t)] &= \sigma_{x_l}^2 t \quad 1 \leq l \leq k_x \\
\text{Cov}[y, x_l(t)] &= (t - \frac{1-e^{-\alpha t}}{\alpha}) \sigma_{x_l}^2 b_l \quad 1 \leq l \leq k_x \\
\text{Cov}[x_r, x_l(t)] &= 0 \quad 1 \leq r, l \leq k_x, r \neq l.
\end{aligned}$$

We calculate the correlation structure between species i and j stemming from the phylogeny,

$$\begin{aligned}
\text{Cov}[y^i, y^j] &= e^{-\alpha(t_i+t_j-2t_{a_{ij}})} \left(\frac{1-e^{-2\alpha t_{a_{ij}}}}{2\alpha} (\sigma_y^2 + \sum_{l=1}^{k_x} \sigma_{x_l}^2 b_l^2) + \right. \\
&\quad \left. + (t_{a_{ij}} - \frac{2(1-e^{-\alpha t_{a_{ij}}})}{\alpha}) \sum_{l=1}^{k_x} \sigma_{x_l}^2 b_l^2 \right) \\
\text{Cov}[y^i, x_l^j] &= e^{-\alpha(t_i-t_{a_{ij}})} (t_{a_{ij}} - \frac{1-e^{-\alpha t_{a_{ij}}}}{\alpha}) \sigma_{x_l}^2 b_l \\
&\quad + (1 - e^{-\alpha(t_i-t_{a_{ij}})}) t_{a_{ij}} \sigma_{x_l}^2 b_l \\
\text{Cov}[x_r^i, x_l^j] &= 0 \quad r \neq l \\
\text{Cov}[x_l^i, x_l^j] &= t_{a_{ij}} \sigma_{x_l}^2.
\end{aligned}$$

We can immediately see that, as long as not all b_l are 0 there will be no stationary distribution and no convergence of the process even in distribution as the Brownian motion type variables causes the variance to escape to infinity at a linear rate.

7.5 Multivariate Ornstein–Uhlenbeck Brownian Motion model

In Paper II we generalize the above Ornstein–Uhlenbeck Brownian Motion model to the situation where y can be a vector $\vec{Y}(t)$ of length k_Y and the x variables can evolve in a correlated fashion. We term it multivariate Ornstein–Uhlenbeck Brownian Motion (mvOUBM). The stochastic differential equation formulation for this model is in equation (4.2) in section 4.2 and we called it the multivariate Ornstein–Uhlenbeck Brownian Motion (mvOUBM). This model has a very special decomposition of the drift matrix,

$$\mathbf{F} = \left[\begin{array}{c|c} \mathbf{A} & \mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right].$$

The easiness of computing any moments of the model hinges directly on the easiness of computing

$$e^{-\mathbf{F}t} = \sum_{n=0}^{\infty} \frac{(-t)^n}{n!} \mathbf{F}^n.$$

We have

$$\left[\begin{array}{c|c} \mathbf{A} & \mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right]^n = \begin{cases} \mathbf{I} & n = 0 \\ \left[\begin{array}{c|c} \mathbf{A}^n & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] + \left[\begin{array}{c|c} \mathbf{A}^{n-1} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \left[\begin{array}{c|c} \mathbf{0} & \mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] & n \neq 0 \end{cases}$$

which can be easily shown by induction. Assuming $\mathbf{A}^{-1}$ exists we get

$$\begin{aligned} e^{-\mathbf{F}t} &= \mathbf{I} + \sum_{n=1}^{\infty} \frac{(-t)^n}{n!} \left(\left[\begin{array}{c|c} \mathbf{A}^n & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] + \left[\begin{array}{c|c} \mathbf{A}^{n-1} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \left[\begin{array}{c|c} \mathbf{0} & \mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \right) \\ &= \mathbf{I} + \left[\begin{array}{c|c} \sum_{n=1}^{\infty} \frac{(-t)^n}{n!} \mathbf{A}^n & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \\ &\quad + \left(\mathbf{I} + \sum_{n=1}^{\infty} \frac{(-t)^n}{n!} \left[\begin{array}{c|c} \mathbf{A}^n & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] - \mathbf{I} \right) \left[\begin{array}{c|c} \mathbf{0} & \mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \\ &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] + \left(\left[\begin{array}{c|c} e^{-t\mathbf{A}} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] - \mathbf{I} \right) \left[\begin{array}{c|c} \mathbf{0} & \mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \\ &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & \mathbf{0} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \left[\begin{array}{c|c} \mathbf{I} & \mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] - \left[\begin{array}{c|c} \mathbf{0} & \mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{0} \end{array} \right] \\ &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right]. \end{aligned}$$

If we further assume that the “optimum” vector $\vec{\psi}(t)$ is a step function we get that $E \left[\vec{Y}^T \vec{X}^T \right] (t)$ equals,

$$\begin{aligned} &\left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \left[\begin{array}{c} \vec{Y} \\ \vec{X} \end{array} \right] (0) \\ &+ \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \sum_{j=1}^m \int_{t_{j-1}}^{t_j} e^{\mathbf{F}v} \mathbf{F} dv \left[\begin{array}{c} \vec{\psi}_j \\ \vec{0} \end{array} \right] \\ &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \left[\begin{array}{c} \vec{Y} \\ \vec{X} \end{array} \right] (0) \\ &+ \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \sum_{j=1}^m (e^{\mathbf{F}t_j} - e^{\mathbf{F}t_{j-1}}) \left[\begin{array}{c} \vec{\psi}_j \\ \vec{0} \end{array} \right] \\ &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \left[\begin{array}{c} \vec{Y} \\ \vec{X} \end{array} \right] (0) \\ &+ \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \sum_{j=1}^m \left(\left[\begin{array}{c|c} e^{t_j\mathbf{A}} & (e^{t_j\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \right. \\ &\quad \left. - \left[\begin{array}{c|c} e^{t_{j-1}\mathbf{A}} & (e^{t_{j-1}\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \right) \left[\begin{array}{c} \vec{\psi}_j \\ \vec{0} \end{array} \right] \\ &= \left[\begin{array}{c} -\mathbf{A}^{-1}\mathbf{B}\vec{X}(0) + e^{-t\mathbf{A}}(\vec{Y}(0) + \sum_{j=1}^m (e^{t_j\mathbf{A}} - e^{t_{j-1}\mathbf{A}})\vec{\psi}_j + \mathbf{A}^{-1}\mathbf{B}\vec{X}(0)) \\ \hline \vec{X}(0) \end{array} \right] \end{aligned}$$

Let us denote $\left[\begin{array}{c|c} \Sigma^{yy} & \Sigma^{yx} \\ \hline \Sigma^{xy} & \Sigma^{xx} \end{array} \right] \left[\begin{array}{c|c} \Sigma^{yy} & \Sigma^{yx} \\ \hline \Sigma^{xy} & \Sigma^{xx} \end{array} \right]^T =: \left[\begin{array}{c|c} \Sigma_{11} & \Sigma_{12} \\ \hline \Sigma_{21} & \Sigma_{22} \end{array} \right]$, and by equation (7.3) we get the covariance equaling,

$$\begin{aligned} &= \left[\begin{array}{c|c} e^{-t\mathbf{A}} & (e^{-t\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right] \\ &\times \int_0^t \left[\begin{array}{c|c} e^{v\mathbf{A}}\Sigma_{11} + (e^{v\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B}\Sigma_{21} & e^{v\mathbf{A}}\Sigma_{12} + (e^{v\mathbf{A}} - \mathbf{I})\mathbf{A}^{-1}\mathbf{B}\Sigma_{22} \\ \hline \Sigma_{21} & \Sigma_{22} \end{array} \right] \\ &\times \left[\begin{array}{c|c} e^{v\mathbf{A}^T} & \mathbf{0} \\ \hline \mathbf{B}^T\mathbf{A}^{-T}(e^{v\mathbf{A}^T} - \mathbf{I}) & \mathbf{I} \end{array} \right] dv \left[\begin{array}{c|c} e^{-t\mathbf{A}^T} & \mathbf{0} \\ \hline \mathbf{B}^T\mathbf{A}^{-T}(e^{-t\mathbf{A}^T} - \mathbf{I}) & \mathbf{I} \end{array} \right]. \end{aligned}$$

This then becomes

$$\begin{aligned} \text{Var} \left[\vec{Y}(t) \right] &= \int_0^t e^{-\mathbf{A}v} \Sigma_{11} e^{-\mathbf{A}^T v} dv \\ &\quad + \int_0^t e^{-\mathbf{A}v} \mathbf{A}^{-1} \mathbf{B} \Sigma_{21} e^{-\mathbf{A}^T v} dv \\ &\quad + \int_0^t e^{-\mathbf{A}v} \Sigma_{12} \mathbf{B}^T \mathbf{A}^{-T} e^{-\mathbf{A}^T v} dv \\ &\quad + \int_0^t e^{-\mathbf{A}v} \mathbf{A}^{-1} \mathbf{B} \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T} e^{-\mathbf{A}^T v} dv \\ &\quad - \mathbf{A}^{-1} \mathbf{B} (\Sigma_{21} + \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T}) \mathbf{A}^{-T} (\mathbf{I} - e^{-\mathbf{A}^T t}) \\ &\quad - (\mathbf{I} - e^{-\mathbf{A}t}) \mathbf{A}^{-1} (\Sigma_{12} + \mathbf{A}^{-1} \mathbf{B} \Sigma_{22}) \mathbf{B}^T \mathbf{A}^{-T} \\ &\quad + t \mathbf{A}^{-1} \mathbf{B} \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T} \\ \text{Cov} \left[\vec{Y}(t), \vec{X}(t) \right] &= (\mathbf{I} - e^{-\mathbf{A}t}) \mathbf{A}^{-1} (\Sigma_{12} + \mathbf{A}^{-1} \mathbf{B} \Sigma_{22}) \\ &\quad - t \mathbf{A}^{-1} \mathbf{B} t \Sigma_{22} \\ \text{Var} \left[\vec{X}(t) \right] &= t \Sigma_{22}. \end{aligned}$$

We calculate the covariance structure coming from the phylogeny between species i and j ,

$$\begin{aligned} \text{Cov} \left[\vec{Y}_i, \vec{Y}_j \right] &= e^{-\mathbf{A}(t_i - t_{a_{ij}})} \left(\int_0^{t_{a_{ij}}} e^{v\mathbf{A}} \Sigma_{11} e^{v\mathbf{A}^T} dv \right. \\ &\quad + \int_0^{t_{a_{ij}}} e^{v\mathbf{A}} \mathbf{A}^{-1} \mathbf{B} \Sigma_{21} e^{v\mathbf{A}^T} dv + \int_0^{t_{a_{ij}}} e^{v\mathbf{A}} \Sigma_{12} \mathbf{B}^T \mathbf{A}^{-T} e^{v\mathbf{A}^T} dv \\ &\quad + \left. \int_0^{t_{a_{ij}}} e^{v\mathbf{A}} \mathbf{A}^{-1} \mathbf{B} \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T} e^{v\mathbf{A}^T} dv \right) e^{-\mathbf{A}^T (t_j - t_{a_{ij}})} \\ &\quad - e^{-\mathbf{A}(t_i - t_{a_{ij}})} (\mathbf{I} - e^{-t_{a_{ij}} \mathbf{A}}) \mathbf{A}^{-1} (\Sigma_{12} + \mathbf{A}^{-1} \mathbf{B} \Sigma_{22}) \mathbf{B}^T \mathbf{A}^{-T} \\ &\quad + \mathbf{A}^{-1} \mathbf{B} (\Sigma_{21} + \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T}) \mathbf{A}^{-T} (\mathbf{I} - e^{-t_{a_{ij}} \mathbf{A}^T}) e^{-\mathbf{A}^T (t_j - t_{a_{ij}})} \\ &\quad + t_{a_{ij}} \mathbf{A}^{-1} \mathbf{B} \Sigma_{22} \mathbf{B}^T \mathbf{A}^{-T} \\ \text{Cov} \left[\vec{Y}_i, \vec{X}_j \right] &= e^{-\mathbf{A}(t_i - t_{a_{ij}})} (\mathbf{I} - e^{-t_{a_{ij}} \mathbf{A}}) (\Sigma_{12} + \mathbf{A}^{-1} \mathbf{B} \Sigma_{22}) - t_{a_{ij}} \mathbf{A}^{-1} \mathbf{B} \Sigma_{22} \\ \text{Cov} \left[\vec{X}_i, \vec{X}_j \right] &= t_{a_{ij}} \Sigma_{22}. \end{aligned}$$

In Figure 4 we present a number of simulations of the trajectories of the multivariate Ornstein–Uhlenbeck Brownian Motion model for different parameter values. More examples of simulations can be found in the Paper II. In all of these cases $\Sigma^{xx} = 1$, $\vec{Y}(0) = \vec{0}$, $\vec{\psi} = \vec{0}$ and $\vec{X}(0) = 0$.

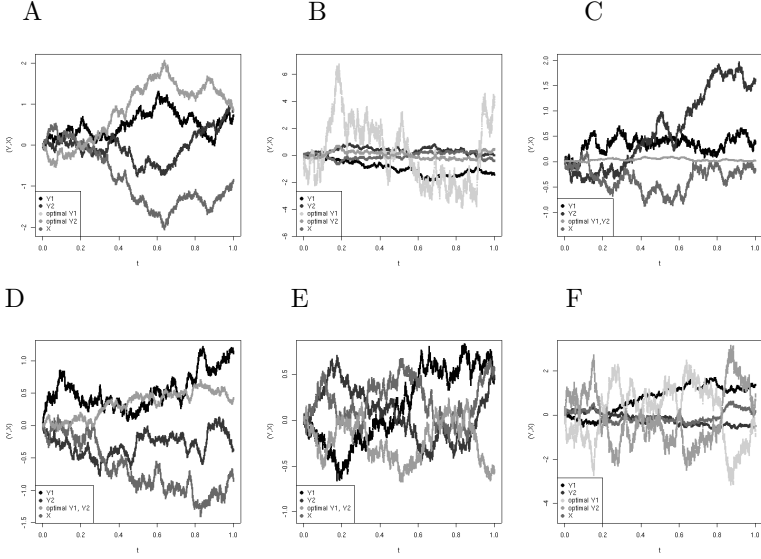


Figure 4: Simulations for different parameters

$$\text{A: } \mathbf{A} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 1 & -0.5 \\ 0 & 1 \end{bmatrix}$$

$$\text{B: } \mathbf{A} = \begin{bmatrix} 0.1 & -0.001 \\ -0.01 & 1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 1 & -0.5 \\ 0 & 1 \end{bmatrix}$$

$$\text{C: } \mathbf{A} = \begin{bmatrix} 0.01 & -0.001 \\ -0.001 & 0.1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 0.001 \\ 0.001 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$\text{D: } \mathbf{A} = \begin{bmatrix} 1 & 0.999 \\ 0.999 & 1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 1 \\ 1 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 0.894 & -0.447 \\ 0 & 0.775 \end{bmatrix}$$

$$\text{E: } \mathbf{A} = \begin{bmatrix} 1 & -0.999 \\ -0.999 & 1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 0.001 \\ 0.001 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 0.894 & -0.447 \\ 0 & 0.775 \end{bmatrix}$$

$$\text{F: } \mathbf{A} = \begin{bmatrix} 1 & 0.999 \\ 0.999 & 1 \end{bmatrix} \mathbf{B} = \begin{bmatrix} 0.005 \\ -0.005 \end{bmatrix} \Sigma^{yy} = \begin{bmatrix} 0.894 & -0.447 \\ 0 & 0.775 \end{bmatrix}$$

Other simulation examples can be found in Paper II.

7.5.1 Stationary regime of the mvOUBM model

All of the discussion about the asymptotic properties of the multivariate Ornstein–Uhlenbeck Brownian Motion model naturally also apply to the univariate one. In particular, there is no stationary regime as the covariance of the process escapes to infinity linearly in time (due to the “ $\vec{X}$ ”

variables evolving marginally as Brownian motion). However, given that $\mathbf{A}$ has positive real part eigenvalues, the process $\vec{Y} - \mathbf{A}^{-1}\mathbf{B}\vec{X} - \vec{\psi}$ (or $y - \vec{b}\vec{X} - \psi$ in the univariate case)

tends (in distribution) to a mean $\vec{0}$ normal random variable with covariance

$$\mathbf{P} \left(\left[\frac{1}{\lambda_i + \lambda_j} \right]_{1 \leq i, j \leq k_Y} * (\mathbf{P}^{-1} (\boldsymbol{\Sigma}_{11} + \boldsymbol{\Sigma}_{12}\mathbf{B}^T\mathbf{A}^{-T} + \mathbf{A}^{-1}\mathbf{B}\boldsymbol{\Sigma}_{21} + \mathbf{A}^{-1}\mathbf{B}\boldsymbol{\Sigma}_{22}\mathbf{B}^T\mathbf{A}^{-T}) \mathbf{P}^{-T}) \right) \mathbf{P}^T.$$

Unlike in the Ornstein–Uhlenbeck type model the covariance between traits in different species after divergence does not disappear but tends to a constant, due to the covariance remaining from the driving Brownian motion,

$$\begin{aligned} \text{Cov} \left[\vec{Y}_i, \vec{Y}_j \right] (t) &\rightarrow t_{a_{ij}} \mathbf{A}^{-1} \mathbf{B} \boldsymbol{\Sigma}_{22} \mathbf{B}^T \mathbf{A}^{-T}, \\ \text{Cov} \left[\vec{Y}_i, \vec{X}_j \right] (t) &\rightarrow -t_{a_{ij}} \mathbf{A}^{-1} \mathbf{B} \boldsymbol{\Sigma}_{22}, \\ \text{Cov} \left[\vec{X}_i, \vec{X}_j \right] (t) &= t_{a_{ij}} \boldsymbol{\Sigma}_{22}. \end{aligned}$$

7.6 Relations between the models

We can discuss a hierarchy amongst the models. In all of the presented cases the distribution of the observed data is multivariate normal. The simplest model is independent observations of predictors with equal mean $\vec{\mu}$ and (co)variance $\boldsymbol{\Sigma}$. The distribution of the data is then multivariate normal $\mathcal{N}(\mathbf{1} \otimes \vec{\mu}, \mathbf{I} \otimes \boldsymbol{\Sigma})$.

We can see that this is a special case of the Brownian motion model with $\mathbf{T} = \mathbf{I}$ as the distribution of the data under the Brownian motion model is $\mathcal{N}(\mathbf{1} \otimes \vec{X}(0), \mathbf{T} \otimes \boldsymbol{\Sigma})$, where $\mathbf{T}$ is defined by the phylogeny as the matrix of times of the most recent common ancestor for each pair of species.

In turn the Brownian motion model can be considered as a submodel of the multivariate Ornstein–Uhlenbeck Brownian Motion model. If in the multivariate Ornstein–Uhlenbeck Brownian Motion model we will have $\mathbf{A} = \mathbf{0}$ and $\mathbf{B} = \mathbf{0}$ then the stochastic differential equation describing the process will be

$$\begin{bmatrix} d\vec{Y}(t) \\ d\vec{X}(t) \end{bmatrix} = \mathbf{0}dt + \begin{bmatrix} \boldsymbol{\Sigma}^{yy} & \boldsymbol{\Sigma}^{yx} \\ \boldsymbol{\Sigma}^{xy} & \boldsymbol{\Sigma}^{xx} \end{bmatrix} d\mathbf{W}(t)$$

and the data distribution is $\mathcal{N}(\mathbf{1} \otimes [\vec{Y}(0) \ \vec{X}(0)]^T, \mathbf{T} \otimes \begin{bmatrix} \boldsymbol{\Sigma}^{yy} & \boldsymbol{\Sigma}^{yx} \\ \boldsymbol{\Sigma}^{xy} & \boldsymbol{\Sigma}^{xx} \end{bmatrix})$.

Notice that we have to derive this from the stochastic differential equation formulation and not from the covariance structure of the process as to calculate it we assumed that $\mathbf{A}$ is invertible.

The multivariate Ornstein–Uhlenbeck Brownian Motion model is in turn a special case of the generalized (to arbitrary drift matrix) Ornstein–Uhlenbeck model, by setting the appropriate “bottom” (if we order the variables) rows to zero. This model is in turn a special case of the general model presented in [34] where one does not assume that $\mathbf{F}$ needs to be invertible.

Then the most general class is the general linear model where we assume any data covariance matrix $\mathbf{V}$ and some superstructure (on the one from the previous models) of the mean values.

One can also consider relation between the models by looking at the different classes of $\mathbf{F}$ (or respectively $\mathbf{A}$). The simplest one is when we assume the matrix has a single value on its diagonal. Then we can consider any diagonal matrix. If we assumed that the diagonal values have to be positive then the next more general class is the class of symmetric positive definite matrices, then the class of matrices with positive eigenvalues and then the general family of invertible matrices. This sort of hierarchy of models allows one to see which model fits better to the data via *e.g* the log-likelihood ratio test. There are of course other model selection criteria like the AIC, AIC_c, BIC *e.t.c.* One can connect biological hypothesis with the different models (or matrix classes) and the model comparison criteria tell us how each hypothesis is favoured.

If we define by k the number of parameters in the model, by n the number of observed values (number of species multiplied by number of variables minus the number of missing values) and by $\mathcal{L}$ the log-likelihood then the different information criteria will have the following formulae,

$$\begin{aligned} \text{AIC} &= 2 \cdot \mathcal{L} + 2 \cdot k \\ \text{AIC}_c &= 2 \cdot \mathcal{L} + 2 \cdot k + \frac{2 \cdot k \cdot (k+1)}{n-k-1} = \text{AIC} + \frac{2 \cdot k \cdot (k+1)}{n-k-1} \\ \text{BIC} &= 2 \cdot \mathcal{L} + k \cdot \log n. \end{aligned}$$

8 Supplementary material C: Matrix parametrizations used in mvSLOUCH package (Paper II)

To be able to work with estimating the parameters we have to be able to move around the space of decomposable matrices which entails the parametrization of the $\mathbf{F}/\mathbf{A}$ matrix. There are two possibilities that all the eigenvalues are real or that some are complex. Since the matrix has k_Y^2 elements we aim to parametrize it by at most k_Y^2 variables. Its worth noting the paper [43] as it considers parametrizations of covariance matrices (symmetric, positive-definite real).

8.1 Real eigenvalues case

Since we assumed that none of eigenvalues are 0, $\mathbf{P}$ is invertible and we can do a unique (if the elements of $\mathbf{R}$'s diagonal are forced to be positive) QR decomposition of $\mathbf{P}$ [28] (page 112). We have that $\mathbf{P} = \mathbf{Q}\mathbf{R}$ where $\mathbf{Q}$ is an orthogonal matrix and $\mathbf{R}$ is upper triangular. Since the R software returns each eigenvector to be of unit length ($\langle x, x \rangle = x^T x$) it was noticed that it returns the $\mathbf{R}$ matrix as such that its columns are also of unit length.

Lemma 8.1. *If $\mathbf{P}$ is invertible and such that its columns are of unit length ($\langle x, x \rangle = x^T x = 1$) and if in its QR decomposition $\mathbf{P} = \mathbf{Q}\mathbf{R}$, $\mathbf{R}$'s diagonal elements are forced to be positive then $\mathbf{R}$'s columns will also be of unit length.*

Proof Denote the elements of the matrices as, $\mathbf{P}_{ij} = p_{ij}$, $\mathbf{Q}_{ik} = q_{ik}$ and $\mathbf{R}_{kj} = r_{kj}$, notice that from the properties of QR decomposition, $\sum_{i=1}^{k_Y} q_{ik} q_{il} = 1$ iff $l = k$, otherwise 0 and for $k > j$ $r_{kj} = 0$. Now consider the length of the j -th column of $\mathbf{P}$,

$$\begin{aligned} 1 &= \mathbf{P}_{\cdot j}^T \mathbf{P}_{\cdot j} = \sum_{i=1}^{k_Y} p_{ij}^2 = \sum_{i=1}^{k_Y} \left(\sum_{k=1}^j q_{ik} r_{kj} \right)^2 \\ &= \sum_{i=1}^{k_Y} \sum_{k_1=1}^j \sum_{k_2=1}^j q_{ik_1} q_{ik_2} r_{k_1 j} r_{k_2 j} \\ &= \sum_{k_1=1}^j \sum_{k_2=1}^j r_{k_1 j} r_{k_2 j} \sum_{i=1}^{k_Y} q_{ik_1} q_{ik_2} = \sum_{k=1}^j r_{kj}^2 = \mathbf{R}_{\cdot j}^T \mathbf{R}_{\cdot j}, \end{aligned}$$

which shows that all columns of $\mathbf{R}$ are also of unit length.

Q.E.D.

We also know that an orthogonal matrix $\mathbf{Q}$ of size $k_Y \times k_Y$ can be represented by $k_Y(k_Y - 1)/2$ Givens rotations [2] (another approach to parametrization is in [49]). Each Givens rotation can be represented by one value [20], since we assume that $c^2 + s^2 = 1$ (see [20] for meaning of notation), and the position pair (i, j) in $\mathbf{P}$ which it is meant to zero. We

therefore have a representation by n eigenvalues, $k_Y(k_Y - 1)/2$ numbers representing the Givens rotations of $\mathbf{Q}$ and $k_Y(k_Y - 1)/2$ values from above $\mathbf{R}$'s diagonal (since we can calculate the diagonal values as each column is of unit length). When generating $\mathbf{R}$ its j th column has to be generated from the j dimensional unit sphere with the diagonal value being positive. And so we have a parametrization with k_Y^2 parameters.

8.2 Complex eigenvalues case

In the complex case the situation is not so simple as the QR decomposition as above would require $2k_Y^2$ parameters if applied directly. We have to try to take advantage of two facts; that if a complex number is an eigenvalue then its conjugate will also be (so we have k_Y parameters describing the eigenvalues) and that if we have a complex eigenvector its conjugate will also be an eigenvector (and so $\mathbf{P}$ will have k_Y^2 distinct elements). A problem with the QR approach is that the $\mathbf{R}$ implementation `qr()` does not work perfectly in the complex case. When $\mathbf{P}^T$ was decomposed by $\mathbf{R}$ and then the product $\mathbf{Q}^T \mathbf{R}^T$, of the calculated $\mathbf{Q}^T$ and $\mathbf{R}^T$ matrices was done, we did not receive $\mathbf{P}^T$ but a matrix where $\mathbf{P}^T$'s rows were permuted. Another approach is to use singular value decomposition, decomposing $\mathbf{P}$ as $\mathbf{P} = \mathbf{W} \mathbf{D} \mathbf{V}^*$, where $\mathbf{W}$ and $\mathbf{V}$ are unitary, $\mathbf{D}$ is a diagonal matrix with positive real values and $*$ is the conjugate transpose.

It was noticed using $\mathbf{R}$'s `svd()` function that if $\mathbf{P}$'s columns are of unit length and for each complex column it has, it also has its conjugate then all values of $\mathbf{W}$ are real (so we have an orthogonal matrix and using Givens rotations we can parametrize it by $k_Y(k_Y - 1)/2$ numbers) and that for every complex row of $\mathbf{V}$, $\mathbf{V}$ also contains its conjugate. It is obvious that with these conditions we will get a matrix $\mathbf{P}$ of the desired properties, but it is not so obvious that this is a situation which will exist for every $\mathbf{P}$. For this we would need to prove a statement as below.

Proposition 8.2. *If $\mathbf{P}$ is invertible and such that its columns are of unit length ($\langle x, x \rangle = x^T x = 1$) and for every complex column it also contains its conjugate then there exists a singular value decomposition $\mathbf{P} = \mathbf{W} \mathbf{D} \mathbf{V}^*$ of it such that $\mathbf{W}$ contains only real entries and for each complex row of $\mathbf{V}$, $\mathbf{V}$ also contains its conjugate.*

Since $\mathbf{V}$ is unitary it is parametrizable by $k_Y(k_Y - 1)/2$ complex Givens rotations (a discussion of the appropriate definition of a complex Givens rotation is in [7]). But this would give us $k_Y(k_Y - 1)$ parameters. One should then utilize the fact that for each complex row we have its complex conjugate, and parametrize $\mathbf{V}^T$.

The poor man's approach is to work on writing out the formulae for the series of Givens rotations and then using this system of linear equations (it is linear in $\mathbf{V}$'s elements, since columns are of unit length,

there are $k_Y^2 - k_Y$ of them) to represent the $\mathbf{V}$ matrix. Other options of parametrizing unitary matrices are in *e.g.* [30], [31], [50], [51] but they are of limited interest for now since they consider an arbitrary unitary matrix.

It remains to deal with the singular value matrix $\mathbf{D} = \text{diag}(\mathbf{d}_1, \mathbf{k}_Y, \mathbf{k}_Y)$. Once we know $\mathbf{W} = [w_{ij}]$ and $\mathbf{V} = [v_{ij}]$ and use the fact that $\mathbf{P}$'s columns are of unit length and polar decomposition I believe that we can discard $\mathbf{D}$ as I believe it should be calculable from the other matrices.

In polar decomposition $\mathbf{P} = \mathbf{U}\mathbf{F}$, where $\mathbf{U}$ is unitary while $\mathbf{F}$ is positive-definite Hermitian, which makes it symmetric with positive real values on the diagonal. We have the relations $\mathbf{F} = \mathbf{V}\mathbf{D}\mathbf{V}^*$ and $\mathbf{U} = \mathbf{W}\mathbf{V}^*$. Doing the multiplications $\mathbf{P} = \mathbf{W}\mathbf{D}\mathbf{V}^*$ and $\mathbf{F} = \mathbf{V}\mathbf{D}\mathbf{V}^*$ we get that for each i ,

$$\sum_{k=1}^n \sum_{l=1}^n v_{ik} v_{il} d_k d_l \sum_{r=1}^n \bar{v}_{kr} \bar{v}_{lr} = 1$$

subject to the constraints

$$\begin{aligned} d_i &\in \mathbb{R} \\ d_i &> 0 \\ \sum_{k=1}^n v_{ik} \bar{v}_{ki} d_k &> 0, \end{aligned}$$

the last constraint coming from that $\mathbf{F}$'s diagonal elements must be real. It remains to be studied how this equation behaves, how many solutions does it have, do the constraints force one solution and does a solution exist for arbitrary orthogonal $\mathbf{U}$ and unitary $\mathbf{V}$ pairs.

8.3 Matrix parametrization implemented in package

An important part of the program is to allow the user to choose a desired type of $\mathbf{F}/\mathbf{A}$ matrix (size k_Y by k_Y), the following ones are currently implemented,

SingleValueDiagonal

A diagonal matrix with just a single value along the diagonal is represented by this number.

Diagonal

A diagonal matrix is represented by a vector which contains this diagonal.

Symmetric

A symmetric matrix is represented by a vector which contains its upper triangular part, in R code

`A[upper.tri(A,diag=T)]<-v.`

SymmetricPositiveDefinite

A symmetric positive definite is represented exactly as in `ouch` (using `ouch's sym.par()`), that is by a vector which contains its Cholesky decomposition. In R code `X[lower.tri(X,diag=T)]<-v; A=X%*%t(X).`

TwoByTwo

An arbitrary 2×2 matrix is stored as `A<-matrix(v,2,2)`.

UpperTri

An upper triangular matrix is stored as `A[upper.tri(A,diag=T)]<-v`.

LowerTri

A lower triangular is stored as `A[lower.tri(A,diag=T)]<-v`.

DecomposablePositive

We now consider storing a decomposable matrix with real positive eigenvalues. It is stored in a vector of length k_Y^2 . The first k_Y values store the logarithm of the eigenvalues. The remaining $k_Y^2 - k_Y$ numbers code the eigenvector matrix $\mathbf{P}$. The matrix is coded by $\mathbf{P}$'s QR decomposition which is unique as $\mathbf{P}$ is invertible and $\mathbf{R}$'s diagonal is assumed to be positive. The coding is done in the following manner, $\mathbf{P} = \mathbf{Q}\mathbf{R}$, where $\mathbf{Q}$ is an orthogonal matrix and $\mathbf{R}$ is upper triangular. We know that $\mathbf{R}$'s columns are of unit length. This is because the eigenvectors ($\mathbf{P}$'s columns) are of unit length and $\mathbf{Q}$ is orthogonal. Therefore since $\mathbf{R}$'s diagonal is positive it is determined by the values of $\mathbf{R}$ above the diagonal. $\mathbf{Q}$ is coded by $k_Y(k_Y - 1)/2$ Givens rotations. $\mathbf{R}$'s columns are coded in the following manner, we know every value has to be between $[-1, 1]$. We can see that

$\mathbf{R}_{ij}$ from v

if $i == 1$ and $j == 1$ then

$\mathbf{R}_{11} = 1$

else if $j > i$ then

if $i == 1$ then

$\mathbf{R}_{ij} = \frac{2}{\pi} \text{atan}(v)$

else

$a = \sqrt{1 - \sum_{k=1}^{i-1} \mathbf{R}_{kj}^2}$

$\mathbf{R}_{ij} = \frac{2a}{\pi} \text{atan}(v)$

end if

else if $j == i$ then

$\mathbf{R}_{ii} = \sqrt{1 - \sum_{k=1}^{i-1} \mathbf{R}_{ki}^2}$

else

$\mathbf{R}_{ij} = 0$

end if

return $\mathbf{R}_{ij}$

this gives, that the matrix is parametrized by k_Y^2 values.

DecomposableNegative

The procedure is exactly the same as in the decomposable positive case, except that the eigenvalues are stored as the logarithm of them negated.

DecomposableReal

The procedure is exactly the same as in the decomposable positive case,

except that the eigenvalues are stored directly.

Invertible

The storing of the matrix as decomposable real does not guarantee it to be invertible (it could have an eigenvalue of 0) and more importantly does not allow for a complex eigendecomposition. Storing an invertible matrix is done by QR decomposition, $\mathbf{F}/\mathbf{A} = \mathbf{Q}\mathbf{R}$. $\mathbf{Q}$ as before is stored by $k_Y(k_Y - 1)/2$ Givens rotations. $\mathbf{R}$'s diagonal is stored as its logarithm (to force it be non-zero, for it to be invertible, and positive as assumed) and the rest are stored as, `R[upper.tri(R,diag=F)]<-v`.

In addition the user can force all elements of the parametrized matrix's diagonal to be positive/negative. This is done after transforming the parameter vector into the matrix by exponentiating the diagonal of $\mathbf{A}$. The matrix Σ^{yy} can be parametrized as either `SingleValueDiagonal`, `Diagonal`, `Symmetric`, `UpperTrim` `LowerTri` or `Any`,
 `$\Sigma^{yy} <- \text{matrix}(v, \text{nrow}=k_Y, \text{ncol}=k_Y)$` .

References

- [1] M. Ackin and C. Ritenbaugh. Analysis of multivariate reliability structures and the induced bias in linear model estimation. *Statistics in Medicine*, 16:1647–1661, 1996.
- [2] T. W. Anderson, I. Olkin, and L. G. Underhill. Generation of random orthogonal matrices. *SIAM Journal on Scientific and Statistical Computing*, 8(4):625–629, 1987.
- [3] C. S. Ballantine. Products of positive definite matrices I. *Pacific Journal of Mathematics*, 23(3):427–433, 1967.
- [4] C. S. Ballantine. Products of positive definite matrices II. *Pacific Journal of Mathematics*, 24(1):7–17, 1968.
- [5] C. S. Ballantine. Products of positive definite matrices III. *Journal of Algebra*, 10(1):174–182, 1968.
- [6] C. S. Ballantine. Products of positive definite matrices IV. *Linear Algebra and Its Applications*, 3(1):79–114, 1970.
- [7] D. Bindel, J. Demmel, and W. Kahan. On computing Givens rotations reliably and efficiently. *ACM Transactions on Mathematical Software*, 28(2):206–238, 2002.
- [8] G. Box and G. M. Jenkins. *Time Series Analysis*. Holden-Day, San Francisco, 1970.
- [9] J. P. Buonaccorsi. *Measurement Error Models, Methods and Applications*. CRC Press, 2010.
- [10] M. A. Butler and A. A. King. Phylogenetic comparative analysis: a modelling approach for adaptive evolution. *The American Naturalist*, 164(6):683–695, 2004.
- [11] N. R. Drapper and H. Smith. *Applied Regression Analysis*. Wiley, New York, 1998.
- [12] L. C. Evans. *An Introduction to Stochastic Differential Equations Version 1.2*. Department of Mathematics UC Berkley.
- [13] J. Felsenstein. Phylogenies and the comparative method. *The American Naturalist*, 125(1):1–15, 1985.
- [14] R. P. Freckleton, P. H. Harvey, and M. Pagel. Phylogenetic analysis and comparative data: A test and review of evidence. *The American Naturalist*, 160(6), 2002.
- [15] W. A. Fuller. *Measurement Error Models*. Wiley, 1987.

- [16] W. A. Fuller. Estimation in the presence of measurement error. *International Statistical Review*, 63(2):121–147, 1995.
- [17] C. W. Gardiner. *Handbook of Stochastic Methods*. Springer, Berlin, 2004.
- [18] T. Garland, A. W. Dickerman, C. M. Janis, and J. A. Jones. Phylogenetic analysis of covariance by computer simulation. *Systematic Biology*, 42(3):265–292, 1993.
- [19] L. J. Gleser. The importance of assessing measurement reliability in multivariate regression. *Journal of the American Statistical Association*, 87(419):696–707, 1992.
- [20] G. H. Golub and C. F. Van Loan. *Matrix Computations*. The Johns Hopkins University Press, Baltimore, 1989.
- [21] A. Grafen. The phylogenetic regression. *Philosophical Transactions of the Royal Society of London B*, 326(1233):119–157, 1989.
- [22] G. R. Grimmett and D. R. Stirzaker. *Probability and Random Processes*. Oxford University Press, Oxford, 2001.
- [23] J. Gurland. Distribution of quadratic forms and ratios of quadratic forms. *The Annals of Mathematical Statistics*, 24(3):416–427, 1953.
- [24] T. F. Hansen. Stabilizing selection and the comparative analysis of adaptation. *Evolution*, 51(5):1341–1351, 1997.
- [25] T. F. Hansen and E. P. Martins. Translating between microevolutionary process and macroevolutionary patterns: the correlation structure of interspecific data. *Evolution*, 50(4):1404–1417, 1996.
- [26] T. F. Hansen, J. Pienaar, and S. H. Orzack. A comparative method for studying adaptation to randomly evolving environment. *Evolution*, 62:1965–1977, 2008.
- [27] P. H. Harvey and M. Pagel. *The Comparative Method in Evolutionary Biology*. Oxford University Press, Oxford, 1991.
- [28] R. A. Horn and C. R. Johnson. *Matrix Analysis*. Cambridge University Press, Cambridge, 1985.
- [29] J. P. Huelsenbeck and B. Rannala. Detecting correlation between characters in a comparative analysis with uncertain phylogeny. *Evolution*, 57(6):1237–1247, 2003.
- [30] C. Jarslkog. Recursive parametrization and invariant phases of unitary matrices. *Journal of Mathematical Physics*, 46(1), 2005.

- [31] C. Jarslkog. A recursive parametrization of unitary matrices. *Journal of Mathematical Physics*, 46, 2005.
- [32] C. R. Johnson. Positive definite matrices. *The American Mathematical Monthly*, 77(3):259–264, 1970.
- [33] A. Labra, J. Pienaar, and T. F. Hansen. Evolution of thermal physiology in *Liolaemus* lizards: Adaptation, phylogenetic inertia, and niche tracking. *The American Naturalist*, 174(2):204–220, 2009.
- [34] E. P. Martins and T. F. Hansen. Phylogeny and the comparative method: a general approach to incorporating phylogenetic information into the analysis of interspecific data. *The American Naturalist*, 149(4):646–667, 1997.
- [35] C. Moler and C. F. Van Loan. Nineteen dubious ways to compute the exponential of a matrix, twenty-five years later. *SIAM Review*, 45(1):1–37, 2003.
- [36] M. Pagel. Seeking the evolutionary regression coefficient: An analysis of what comparative methods measure. *Journal of Theoretical Biology*, 164:191–205, 1993.
- [37] M. Pagel. Detecting correlated evolution on phylogenies: A general method for the comparative analysis of discrete characters. *Proceedings of the Royal Society B*, 255:37–45, 1994.
- [38] M. Pagel. Inferring the historical patterns of biological evolution. *Nature*, 401:877–884, 1999.
- [39] M. Pagel. The maximum likelihood approach to reconstructing ancestral character states of discrete characters on phylogenies. *Systematic Biology*, 48(3):612–622, 1999.
- [40] E. Paradis. *Analysis of Phylogenetics and Evolution with R*. Springer, New York, 2006.
- [41] E. Paradis and J. Claude. Analysis of comparative data using generalized estimating equations. *Journal of Theoretical Biology*, 218:175–185, 2002.
- [42] J. Pienaar, K. Bartoszek, T. F. Hansen, and K. Voje. Overview of comparative methods for studying adaptation on adaptive landscapes. in prep.
- [43] J. C. Pinheiro and D. M. Bates. Unconstrained parametrizations for variance–covariance matrices. *Statistics and Computing*, 6:289–296, 1996.

- [44] F. Plard, C. Bonenfant, and J. M. Gaillard. Revisiting the allometry of antlers among deer species: male–male sexual competition as a driver. *Oikos*, 120:601–606, 2011.
- [45] R Development Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria, 2010. ISBN 3-900051-07-0.
- [46] M. D. Smith. On the expectation of a ratio of quadratic forms in normal variables. *Journal of Multivariate Analysis*, 31:244–257, 1989.
- [47] M. D. Smith. Expectations of a ratios of quadratic forms in normal variables: evaluating some top–order invariant polynomials. *Australian Journal of Statistics*, 35(3):271–282, 1993.
- [48] M. D. Smith. Comparing approximations to the expectations of a ratio of quadratic forms in normal variables. *Econometric Reviews*, 15(1):81–95, 1996.
- [49] G. W. Stewart. The efficient generation of random orthogonal matrices with an application to condition estimators. *SIAM Journal of Numerical Analysis*, 17(3):403–409, 1980.
- [50] P. Diță. On the parametrization of unitary matrices by the moduli of their elements. *Communications in Mathematical Physics*, 159:581–591, 1994.
- [51] P. Diță. Factorization of unitary matrices. *Journal of Physics A: Mathematical and General*, 36:2781–2789, 2003.
- [52] C. F. Van Loan. The sensitivity of the matrix exponential. *SIAM Journal on Numerical Analysis*, 14(6):971–981, 1977.
- [53] C. F. Van Loan. Computing integrals involving the matrix exponential. *IEEE Transactions on Automatic Control*, 23(3):395–404, 1978.
- [54] A. Wentzell. *Course in the Theory of Stochastic Processes*. McGraw–Hill Inc, 1981.
- [55] D. Williams. *Probability with Martingales*. Cambridge University Press, Cambridge, 1991.
- [56] P. Y. Wu. Products of positive semidefinite matrices. *Linear Algebra and its Applications*, 111:53–61, 1988.