

A Bayesian Approach on the Effect of Different Covariance Structures on Repeated Measures Data

T. O. Yomi-Owojori¹, N. O. Afolabi², A. H. Ekong^{3*} and B. N. Okafor⁴

^{1,4} *National Horticultural Research Institute, Ibadan, Nigeria;* ² *Ogun State Institute of Technology, Igbesa, Nigeria;* ³ *Department of Statistics, Federal University of Agriculture Abeokuta, Ogun State, Nigeria*

Abstract. The correlation between pairs of repeated measures influences the estimation of the within-subject fixed factor effect. Bayesian approach considers the correlations as random in reality unlike in classical approach where they are assumed constants. Multivariate Wishart distribution is assumed prior for the covariance structure of the measurements, but this distribution is influenced by the scale matrix parameter. This study investigates the effect of three scale matrices namely the Identity matrix (ID), Variance Component matrix (VC) and Compound Symmetry matrix (SC) on the distribution and estimates of model under the Bayesian framework. Simulated data set for different sample sizes and a real repeated measures data set were both analysed for the three covariance structure models and comparisons of the models were made using deviance information criterion (DIC) and residual sum of squares (RSS). The results from the simulation study indicate that the data-based VC and SC prior specifications performed better compared to the ID, only in small sample size situations, as all three converge in large sample size. The SC scale matrix fits the real data set better than VC matrix and can be used as suitable prior for Wishart distribution on correlation of measurements in repeated measures data in Bayesian analysis.

Keywords: Covariance structure, Bayesian method, repeated measures, posterior, Wishart.

Published by: Department of Statistics, University of Benin, Nigeria

1. Introduction

It is very common in experiments in agricultural settings for multiple or repeat measurements to be made on an experimental units. For example to test the effect of fertilizer on the growth of specie of seedlings, after applying the treatment, variables such as plant height and girth are measured periodically, say once a week, for several weeks. It is important to account for individual differences of the seedlings in reaction to the fertilizer. To account for these individual differences the same seedling is tested in all conditions of the experiment (Field, 2016). In this kind of design, referred to as repeated measures designs, the assumptions of analysis of variance of experimental designs, particularly independence of response can no longer be guaranteed, and as a result there are correlations in the residual errors among time periods.

Significant steps in the analysis of repeated-measures data were made with the introduction of linear and nonlinear mixed-effects models (Olofen, et al., 2004), which distinguish within-subjects variance (from multiple measurements in each subject) versus between-subjects variance (from multiple subjects being measured). Distinguishing these types of variance can also be thought of as explicitly modeling random error in the data. This is useful in understanding how different individuals are from one another as compared to how different multiple measurements are for given individuals. The subject-factor is the random factor, while the treatment factor is fixed, the idea is to estimate the effect of some levels of the treatment on a sample of subjects; a mixed model analysis of variance is obtained.

Bayesian approach to statistical analysis especially in complex problems is gaining more acceptance with the availability of powerful statistical software and the increasing popularity of Markov Chain Monte Carlo (MCMC) methods. Bayesian ideas offer the benefits of prior information and hierarchical models that allow partial pooling of different data sources (Gelman (2014), Chaloner and Verdinelli (1995)). Bayesian method

*Corresponding author. Email: anieekong@outlook.com

provides an incorporated statistical analysis of an experiment, when good prior information has been selected. In experimental studies prior information can be available from historical control data (Campbell, 2010).

Kogan et al. (2016) studied the prediction accuracy in multivariate repeated-measures Bayesian forecasting models with examples drawn from research on Sleep and Circadian rhythms. They noted that in study designs with repeated measures for multiple subjects, population models capturing within- and between-subjects variances enable efficient individualized prediction of outcome measures (response variables) by incorporating individual's response data through Bayesian forecasting. When measurement constraints preclude reasonable levels of prediction accuracy, additional (secondary) response variables measured alongside the primary response may help to increase prediction accuracy.

Kogan et al. (2016) investigated this prediction accuracy for the case of substantial between-subjects correlation between primary and secondary response variables, assuming negligible within-subjects correlation and showed how to determine the accuracy of primary response predictions as a function of secondary response observations. They estimated the subject-specific sleep parameters in their medical study using polysomnography and wrist actigraphy. They also considered prediction accuracy in time-dependent, linear model and derive equations for the optimal timing of measurements to achieve, on average, the best prediction accuracy (Kogan et al., 2016). The study of Kogan et al. (2016) failed to consider the important within-subject correlation which is the distinguishing factor of a repeated measure design.

The work of Song et al. (2017) focused on a Bayesian approach to mixed-effects analysis of accuracy data in repeated-measures designs using mixed binomial regression models with application to human language, memory, and other cognitive processes. Their concern was response accuracy as the primary dependent measure. Song et al. (2017) presented logistic and probit mixed models that allow for random subject and item effects, as well as interactions between experimental conditions and both items and subjects in either one- or two-factor repeated-measures designs. They considered Bayesian model selection through Bayes factor via the Bayesian Information Criterion approximation and the Watanabe-Akaike Information Criterion. This they did with simulations to demonstrate the advantages of Bayesian approach over the more standard approach that consists of aggregating the accuracy data across trials within each condition and over the contemporary use of logistic and probit mixed models with model selection based on the Akaike Information Criterion (Song et al., 2017). Their work focused on categorical response variable only.

Luo and Wang (2015) studied Bayesian hierarchical model for multiple repeated measures and survival data with application to parkinson's disease in clinical studies. Their study was on a clinical site considering multisite studies and multiple patients recruited and treated by the same clinical site. Their study was to find out site-correlation and they saw that there was a significant site-correlation because of common environmental and socioeconomic status, and similar quality of care within site. Luo and Wang (2015) studied several hierarchical joint models with the hazard of terminal events dependent on shared random effects from various levels. They also conducted simulation study to evaluate the performance of various models under different scenarios of sites and patients. Their proposed hierarchical joint models were applied to the clinical trials data sets with binary outcomes.

In the case of repeated measures, the residual consists of a matrix of values consisting of diagonals of residual variances at each repeated measure, and covariance between successive measures as off-diagonals. There are various type of the covariance structures and finding the best covariance structure is much of the work in modeling repeated measures. Generally a subset of candidate structures is considered as in a repeated measures analysis. These include UN (Unstructured) as above, SC (Compound Symmetry), for longitudinal studies, AR(1) (Autoregressive lag 1) – if time intervals are evenly spaced or SP(POW) (Spatial Power) - if time intervals are unequally spaced. The dependence of the measures in the covariance matrix is usually examined using the correlation coefficients matrices (PSU EDU, 2018). We shall be examining the effects of various variance structures can on the relationship between within-subject observations for repeated measures data under the Bayesian statistical approach.

2. Materials and method

2.1 Basic framework of Bayesian approach

Informally, to make inference about parameter β is to learn about the unknown β from data Y , i.e., based on the data, explore which values of β are probable and the extent of uncertainty associated with such estimates. In addition to having a model $f(y|\beta)$ and a likelihood function, Bayesian needs a distribution for

β . The distribution is called a prior distribution because it quantifies the uncertainty about β prior to seeing data. The prior may represent a blending of subjective belief and knowledge, in which case it would be a subjective prior. Alternatively, it could be a conventional prior supposed to represent small or no information (also called objective prior) (Ghosh et al., 2006).

The Bayesian framework is based on Bayes' Theorem, the conditional probability density of β given $Y = y$ which is given as

$$p(\beta|y) = \frac{p(\beta)f(y|\beta)}{\int p(\beta)f(y|\beta) d\beta} \quad (1)$$

where $p(\beta)$ is the prior density function and $f(y|\beta)$ is the density of Y , interpreted as the conditional density of Y given β . The numerator is the joint density of β and Y and the denominator is the marginal density of Y .

The conditional density $p(\beta|y)$ of β given $Y = y$ is called the posterior density, a quantification of our uncertainty about β in the light of data. The posterior distribution or a summary descriptive measures associated with the posterior distribution can be reported. For example, for a real valued parameter β , the posterior mean can be reported as

$$E(\beta|y) = \int_{-\infty}^{\infty} \beta p(\beta|y) d\beta \quad (2)$$

and for the posterior variance

$$\begin{aligned} Var(\beta|y) &= E(\beta - E(\beta|y))^2|y \\ &= \int_{-\infty}^{\infty} (\beta - E(\beta|y))^2 p(\beta|y) d\beta \end{aligned} \quad (3)$$

Finally, one could use the posterior distribution to answer more structured problems like estimation and inference. In the case of estimation of β , one would report the above summary measures which are the posterior odds of the relevant hypotheses (Ghosh et al., 2006).

2.2 Model parameters and priors specification

The one-way within-subject design may be modeled with an appropriate full model as

$$\mathbf{y} = \mu\mathbf{1} + \alpha + \beta + \epsilon \quad (4)$$

The subject-factor, α , is the random factor, while the treatment factor, β , is fixed, the idea is to estimate the effect of some levels of the treatment on a sample of subjects. The residual consists of a matrix of values called the variance-covariance matrix. The advantage of Bayesian approach is that there is a framework for such estimation of the covariance structure from the observed data by placing prior distribution on the covariance structure, which measures the dependence between the treatments.

The one-way model with single observation per treatment for each subject given as:

$$\mathbf{y} = \mu + \beta + \epsilon \quad (5)$$

where $\epsilon \sim N(0, \sigma^2)$ and the observations y are assumed to be correlated and follow a multivariate normal distribution with unknown population mean vector μ and covariance matrix Σ . That is $y \sim N(\mu, \Sigma)$. Non-informative independent normal priors are specified for the within-factor coefficients $\beta \sim N(\mu_0, k\Sigma)$, with $k\Sigma$ set to be large enough so as not to commit to any direction and difference on the factor effect. Since it is expected for realistic variation among the levels of a factor, the difference between the effect of one level from that of the other levels of the same factor cannot be inconsistent or unusual.

The prior distribution that is typically used for the covariance matrix of multivariate normally distributed variables, such as the covariance matrix Σ for the random effects, is the inverse-Wishart (IW) distribution (Gelman et al., 2013; Gelman and Hill, 2007). The IW distribution is a conjugate prior for the covariance matrix of multivariate normal distributed variables, which implies that when it is combined with the likelihood function, it will result in a posterior distribution that belongs to the same distributional family. Another important advantage of the IW distribution is that it ensures positive definiteness of the covariance matrix (Schuurman et al., 2016).

The IW distribution is specified with an $r \times r$ variance matrix V_0 , where r is equal to the number of random parameters, and with a number of degrees of freedom df , with the restriction that $df > r - 1$. V_0 is used to position the IW distribution in parameter space, and the df set the certainty about the prior information in the scale matrix. The larger the df , the higher the certainty about the information in V_0 , and the more informative is the distribution (Gelman et al., 2013; Gelman and Hill, 2007). The least informative specification results when $df = r$, which is the lowest possible number of df . Schuurman et al., (2016) noted that when the df increases, the variance V_0 will become smaller, which implies that the IW distribution will become more informative. Also the size of the variance is partly determined by Σ , the smaller the elements of Σ , the smaller the variance of the IW distribution, and hence the more informative the prior will be. However, setting the variance to large values also influences the position of the IW distribution in parameter space, in other words, specifying a IW prior distribution requires balancing the size of Σ and the df (Schuurman et al., 2016).

These priors specification can be summarized as,

- Likelihood:

$$N(y|\mu, \Sigma) \propto |\Sigma|^{\frac{N}{2}} \exp \left(-\frac{1}{2} \sum_{i=1}^N (y' \Sigma y - 2\mu' \Sigma y + \mu' \Sigma \mu) \right) \quad (6)$$

- Normal prior:

$$N(\beta|\mu_0, k\Sigma) \propto |\Sigma|^{\frac{1}{2}} \exp \left(-\frac{1}{2} (\beta' \Sigma \beta - 2\beta' k \Sigma \mu_0 + \mu_0' k \Sigma \mu_0) \right) \quad (7)$$

- Wishart prior:

$$W(\Sigma|v_0, V_0) \propto |\Sigma|^{\frac{v_0-D-1}{2}} \exp \left(-\frac{1}{2} \text{tr}(V_0^{-1} \Sigma) \right) \quad (8)$$

2.3 Bayesian posterior distribution

According to Bayes theorem in equation (1), the posterior distribution is proportional to the prior multiplied by the likelihood,

$$p(\beta|y) \propto p(\beta) \cdot p(y|\beta)$$

Therefore, taking the product of the likelihood and the priors in equations (6) to (8) gives

$$\begin{aligned} p(\beta) \cdot p(y|\beta) = & |\Sigma|^{\frac{N}{2}} \exp \left\{ \left(-\frac{1}{2} \sum_{i=1}^N (y' \Sigma y - N \bar{y}' \Sigma \mu - \mu' \Sigma N \bar{y} + \mu' \Sigma \mu) \right) \right\} \\ & \times |\Sigma|^{\frac{v_0-D-1}{2}} \exp \left(-\frac{1}{2} \text{tr}(V_0^{-1} \Sigma) \right) \\ & \times |\Sigma|^{\frac{1}{2}} \exp \left(-\frac{k}{2} (\beta' \Sigma \beta - \beta' \Sigma \mu_0 - \mu_0' \Sigma \beta + \mu_0' \Sigma \mu_0) \right) \end{aligned} \quad (9)$$

This can be rewritten as

$$p(\beta) \cdot p(\mathbf{y}|\beta) = |\Sigma|^{\frac{v_0+N-D-1}{2}} \times \exp \left\{ -\frac{1}{2} \left((k+N)\beta' \Sigma \beta - \mu' \Sigma (k\mu_0 + N\bar{\mathbf{y}})' - (k\mu_0 + N\bar{\mathbf{y}})' \Sigma \mu + k\mu_0 \Sigma \mu_0 + \sum_{i=1}^N \mathbf{y}' \Sigma \mathbf{y} + tr(\mathbf{V}_0^{-1} \Sigma) \right) \right\} \quad (10)$$

Also the term in equation (10)

$$(k+N)\beta' \Sigma \beta - \mu' \Sigma (k\mu_0 + N\bar{\mathbf{y}})' - (k\mu_0 + N\bar{\mathbf{y}})' \Sigma \mu + k\mu_0 \Sigma \mu_0 + \sum_{i=1}^N \mathbf{y}' \Sigma \mathbf{y} + tr(\mathbf{V}_0^{-1} \Sigma)$$

can be rewritten as follows by adding and subtracting a term:

$$\begin{aligned} & (k+N)\beta' \Sigma \beta - \mu' \Sigma (k\mu_0 + N\bar{\mathbf{y}})' - (k\mu_0 + N\bar{\mathbf{y}})' \Sigma \mu \\ & + \frac{1}{k+N} (k\mu_0 + N\bar{\mathbf{y}})' \Sigma (k\mu_0 + N\bar{\mathbf{y}}) \\ & - \frac{1}{k+N} (k\mu_0 + N\bar{\mathbf{y}})' \Sigma (k\mu_0 + N\bar{\mathbf{y}}) \\ & + k\mu_0 \Sigma \mu_0 + \sum_{i=1}^N \mathbf{y}' \Sigma \mathbf{y} + tr(\mathbf{V}_0^{-1} \Sigma) \end{aligned}$$

The top two lines now factorise as

$$(k+N) \left(\beta - \frac{k\mu + N\bar{\mathbf{y}}}{k+N} \right)' \Sigma \left(\beta - \frac{k\mu + N\bar{\mathbf{y}}}{k+N} \right)$$

and subtracting $N\bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}}$, the following:

$$-\frac{1}{k+N} (k\mu_0 + N\bar{\mathbf{y}})' \Sigma (k\mu_0 + N\bar{\mathbf{y}}) + c + \sum_{i=1}^N \mathbf{y}' \Sigma \mathbf{y} + tr(\mathbf{V}_0^{-1} \Sigma)$$

can be written as

$$\begin{aligned} & \sum_{i=1}^N (\mathbf{y}' \Sigma \mathbf{y} - \mathbf{y}' \Sigma \bar{\mathbf{y}} - \bar{\mathbf{y}}' \Sigma \mathbf{y} + \bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}}) + N\bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}} + k\mu_0 \Sigma \mu_0 \\ & - \frac{1}{k+N} (k\mu_0 + N\bar{\mathbf{y}})' \Sigma (k\mu_0 + N\bar{\mathbf{y}}) + tr(\mathbf{V}_0^{-1} \Sigma) \end{aligned}$$

The sum term

$$\sum_{i=1}^N (\mathbf{y}' \Sigma \mathbf{y} - \mathbf{y}' \Sigma \bar{\mathbf{y}} - \bar{\mathbf{y}}' \Sigma \mathbf{y} + \bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}}) = \sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})' \Sigma (\mathbf{y} - \bar{\mathbf{y}}).$$

Now

$$N\bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}} + k\mu_0 \Sigma \mu_0 - \frac{1}{k+N} (k\mu_0 + N\bar{\mathbf{y}})' \Sigma (k\mu_0 + N\bar{\mathbf{y}}) + tr(\mathbf{V}_0^{-1} \Sigma)$$

can be expanded as

$$\begin{aligned} & N\bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}} + k\mu_0 \Sigma \mu_0 - \frac{1}{k+N} (k^2 \mu_0' \Sigma \mu_0 + Nk\mu_0 \Sigma \bar{\mathbf{y}} + Nk\bar{\mathbf{y}} \Sigma \mu_0 + N^2 \bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}}) \\ & = \frac{Nk}{k+N} (\bar{\mathbf{y}}' \Sigma \bar{\mathbf{y}} - \bar{\mathbf{y}}' \Sigma \mu_0 - \mu_0 \Sigma \bar{\mathbf{y}} + \mu_0' \Sigma \mu_0) = \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)' \Sigma (\bar{\mathbf{y}} - \mu_0). \end{aligned}$$

The following terms are scalars:

$$\sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})' \boldsymbol{\Sigma} (\mathbf{y} - \bar{\mathbf{y}}) \text{ and } \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)' \boldsymbol{\Sigma} (\bar{\mathbf{y}} - \mu_0)$$

and any scalar is equal to its trace, so

$$\text{tr}(\mathbf{V}_0^{-1} \boldsymbol{\Sigma}) + \sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})' \boldsymbol{\Sigma} (\mathbf{y} - \bar{\mathbf{y}}) + \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)' \boldsymbol{\Sigma} (\bar{\mathbf{y}} - \mu_0)$$

can be rewritten as

$$\text{tr}(\mathbf{V}_0^{-1} \boldsymbol{\Sigma}) + \text{tr} \left(\sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})' \boldsymbol{\Sigma} (\mathbf{y} - \bar{\mathbf{y}}) \right) + \text{tr} \left(\frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)' \boldsymbol{\Sigma} (\bar{\mathbf{y}} - \mu_0) \right)$$

Since $\text{tr}(\mathbf{ABC}) = \text{tr}(\mathbf{CAB})$, the above sum equals

$$\text{tr}(\mathbf{V}_0^{-1} \boldsymbol{\Sigma}) + \text{tr} \left(\sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})(\mathbf{y} - \bar{\mathbf{y}})' \boldsymbol{\Sigma} \right) + \text{tr} \left(\frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)(\bar{\mathbf{y}} - \mu_0)' \boldsymbol{\Sigma} \right).$$

Using the fact that $\text{tr}(\mathbf{A} + \mathbf{B}) = \text{tr}(\mathbf{A}) + \text{tr}(\mathbf{B})$, we can rewrite the sum as

$$\begin{aligned} & \text{tr} \left(\mathbf{V}_0^{-1} \boldsymbol{\Sigma} + \sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})(\mathbf{y} - \bar{\mathbf{y}})' \boldsymbol{\Sigma} + \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)(\bar{\mathbf{y}} - \mu_0)' \boldsymbol{\Sigma} \right) \\ &= \text{tr} \left(\left(\mathbf{V}_0^{-1} + \sum_{i=1}^N (\mathbf{y} - \bar{\mathbf{y}})(\mathbf{y} - \bar{\mathbf{y}})' + \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)(\bar{\mathbf{y}} - \mu_0)' \right) \boldsymbol{\Sigma} \right). \end{aligned}$$

Putting all of that together, let $\mathbf{S} = \sum_{i=1}^N (\bar{\mathbf{y}} - \bar{\mathbf{y}})(\bar{\mathbf{y}} - \bar{\mathbf{y}})'$, then the result of the product of the likelihood and the priors gives

$$\begin{aligned} p(\beta) \cdot p(\mathbf{y}|\beta) &= |\boldsymbol{\Sigma}|^{\frac{1}{2}} \times \exp \left\{ -\frac{k+N}{2} \left(\beta - \frac{k\mu+N\bar{\mathbf{y}}}{k+N} \right)' \boldsymbol{\Sigma} \left(\beta - \frac{k\mu+N\bar{\mathbf{y}}}{k+N} \right) \right\} \\ &\times |\boldsymbol{\Sigma}|^{\frac{v_0+N-D-1}{2}} \exp \left\{ -\frac{1}{2} \text{tr} \left(\left(\mathbf{V}_0^{-1} + \mathbf{S} + \frac{Nk}{k+N} (\bar{\mathbf{y}} - \mu_0)(\bar{\mathbf{y}} - \mu_0)' \right) \boldsymbol{\Sigma} \right) \right\} \end{aligned} \quad (11)$$

From Bayes' Theorem, the conditional probability density of β given $\mathbf{Y} = \mathbf{y}$ is given as

$$p(\beta|\mathbf{y}) = \frac{p(\beta) \cdot p(\mathbf{y}|\beta)}{\int p(\beta) \cdot p(\mathbf{y}|\beta) d\beta}$$

To obtain the posterior distribution, the integral of equation (11) has to be evaluated and used in the formula but this integral is a high dimensional integral and intractable in evaluating, hence there is no closed form expression for this posterior. The posterior distribution for one dimension can be approximated using numerical integration methods like Gaussian Quadrature technique and different Monte Carlo integration techniques, but these processes are very difficult manually, and for high-dimensional integrals, these processes can be evaluated using algorithms implemented in statistical software such as CRAN.

2.4 Inverse-Wishart prior specifications for different covariance structures

The objective here is to examine how to specify an uninformative prior distribution for parameter of interest, so that the influence of the prior specification on the estimates of Σ random effects is minimal, under the specific circumstance that the true sizes of some of these variances are small, as would be the case for the coefficients β . The Inverse-Wishart distribution is a conjugate prior for the covariance matrix of multivariate normal distributed variables, which implies that when it is combined with the likelihood function, it will result in a posterior distribution that belongs to the same distributional family. Another important advantage of the Inverse-Wishart distribution is that it ensures positive definiteness of the covariance matrix (Schuurman et al., 2016).

The first Inverse-Wishart prior specification for the variance V_0 to be examined is the one that is commonly used as an uninformative prior specification, and which shall referred to as the Identity Matrix (ID) specification. In this specification the diagonal elements of variance structure are set to 1 and the off diagonal elements are set to zero (Schuurman et al., 2016). The second specification for the variance V_0 is based on prior estimates of the variances of the random parameters. Using estimates of the variances as input for the IW prior specification ensures that the prior specification will be close to the data, and therefore should limit bias. This covariance structure is denoted as Variance Component (VC).

The third specification for the variance V_0 is referred to as the Compound Symmetry (SC) variance matrix. Here, the constant variance from the data is a scalar factor for the variance matrix having the correlations between treatments within subjects assumed to be the same for each set of treatments, regardless of the difference in the levels of the treatment factor.

The simulation study consists in examining the performance of the Wishart priors for different sizes of (small) variances in Σ , that is the three different covariance structures and for different sample sizes of 15, 25, 50 and 100, based on prior estimates of the variances of the random parameters from the simulated data. The SC with constant variance is obtained from the simulated data as the scalar factor for the scale matrix having the correlations between treatments within subjects assumed to be the same for each set of treatments as set as 0.5, regardless of the difference in the levels of the treatment factor. We compared the three prior specifications for Σ and estimates of all models using WinBUGS in conjunction with the R-package R2WinBUGS.

A real data from an experiment on the production of citrus fruits carried out in the screen house nursery of the National Horticultural Research Institute (NIHORT) in the savanna transition zone of Nigeria was analysed under this framework. The experiment investigates the effect of different treatments (growth media) on citrus seedling production to attain buddable size within a shorter time frame. The growth media include cured sawdust(CS), poultry manure(PM), Sharp Sand(SS), Top Soil(TS), Cured Sawdust:Sharp Sand(CS:SS)(1:1), Top Soil:Sharp Sand(TS:SS)(1:1), Top Soil:Sawdust:Poultry Manure:Sharp Sand(T:S:P:S), Sawdust:Topsoil (S:T)(1:1), Sharp Sand:Poultry Manure(S:P)(1:2), Sharp Sand:Poultry Manure(S:P)(1:1), Poultry Manure:Top Soil(P:T), Poultry Manure:Sharp Sand:Top Soil(P:S:T) and Cured Sawdust:Top Soil:Sharp Sand:Poultry Manure(C:T:S:P). The effect of these growth media was observed on the number of leaves for which repeated measurements were taken in 4 weeks intervals. (See the real dataset in Appendix A).

2.5 Markov Chain Monte Carlo (MCMC) numerical approximations

The posterior distribution $p(\beta|y)$ is not analytically tractable manually and so the marginal posterior $p(\beta_j|y)$ to be obtained is gotten using Markov Chain Monte Carlo (MCMC) method by sampling joint posterior distribution $p(\beta|y)$. Independent sampling is difficult but dependent sampling from Markov chain with $p(\beta|y)$ at its equilibrium distribution is easier. A sequence of random variables $\beta^{(0)}, \beta^{(1)}, \beta^{(2)}, \dots$, forms a Markov chain if $\beta^{(i+1)} \sim p(\beta|\beta^{(i)})$, that is, the next value, $\beta^{(i+1)}$, is conditioned on the present value, $\beta^{(i)}$, and independent of past values $\beta^{(i-1)}, \dots, \beta^{(0)}$.

The Gibbs sampling is popularly used in sampling from posterior distribution using Gibbs algorithm which generates a Markov chain sampling from full conditional distributions since they are conditioned on all other parameters (Best et al., 2011). The Gibbs sampling proceeds by choosing starting values $\beta_j^{(0)}, j = 1, \dots, t$ and from the distribution $p(\beta_1|\beta_2^{(0)}, \beta_3^{(0)}, \dots, \beta_t^{(0)}, y)$, a new value $\beta_1^{(1)}$ is sampled. Next sample $\beta_2^{(1)}$ from $p(\beta_2|\beta_1^{(1)}, \beta_3^{(0)}, \dots, \beta_t^{(0)}, y)$ and so on, till $\beta_t^{(1)}$ is sampled from $p(\beta_t|\beta_1^{(1)}, \beta_2^{(1)}, \dots, \beta_{t-1}^{(1)}, y)$. This process is repeated many thousands of times to eventually obtain sample from $p(\beta|y)$, which gives the posterior

estimates of the parameter.

In summary, suppose the Gibbs sampler is at iteration s , then:

- (1) From sample $\beta^{(s+1)}$ from its full conditional distribution, β_n and Σ_n are computed from $\mathbf{Y}$ and Σ^{s+1} , with sample $\beta^{(s+1)} \sim N(\mu, \Sigma)$.
- (2) From sample Σ^{s+1} from its full conditional distribution and $\mathbf{Y}$, $\mathbf{V}_0$ is computed, with $\Sigma^{s+1} \sim \text{inverseWishart}(v_0, \mathbf{V}_0)$.

The Gibbs algorithms is implemented particularly in Windows Bayesian Analysis Using Gibbs Sampling (WinBUGS) software (Lunn et al., 2000) in conjunction with the R-package R2WinBUGS (Sturtz et al., 2005), which both shall be used in data analysis in the next chapter, to evaluate the posteriors estimates as discussed in this section for Bayesian methods for repeated measures of the crossover designs.

3. Results and discussion

3.1 Simulation results on prior for covariance structures

Table 1 shows the results of the model estimates and statistics for the simulation data with the Wishart prior set on the three variance structures of ID, VC and SC. It can be observed that with sample size of 15, examination of the differences in DIC show a higher magnitude of difference between the ID model and VC model and also between ID model and SC model.

Table 1: Simulation results of variance matrices for sample size 15

Variable	Sample size 15		
	ID Model	VC Model	SC Model
β_0	-26.34	24.27	21.19
β_1	5.68	0.2541	0.692
μ_1	19.10	26.30	26.73
μ_2	21.94	26.43	27.57
μ_3	24.78	26.57	27.42
μ_4	27.62	26.69	27.77
μ_5	30.46	26.81	28.11
RSS	1397	878.40	907.30
Deviance	79.97	97.26	99.74
DIC	-296.710	103.48	105.16
pD	-376.68	6.22	5.42

These differences are substantial and show that the ID model is significantly different from the other variance structures. The VC variance structure is supported as best by the DIC and RSS which clearly shows that the VC model is superior (smallest RSS) fit than the SC model for small sample size.

For the small sample size of 15, the parameter estimates of the group means from Table 1 can be observed to be quite different for the three covariance structures defined for the Wishart distribution. This difference results in marked differences too in the regression coefficients for the repeated measures models defined by the covariance structures of ID, VC and SC.

From Table 2 it can be observed that with sample size of 25, the ID model gets closer to the VC and SC models in its RSS and DIC values. Although the DIC of ID model is lower than the DICs of VC and SC models, its RSS is higher than that of VC and SC. For an increase of sample size to 25, the parameter estimates of the group means from Table 2 can be observed to be similar for covariance structures of VC and SC, unlike ID structure and this difference results in marked differences too in the regression coefficients for the repeated measures models defined for each covariance structure.

From Table 3 it can be observed that with sample size of 50, the differences in DIC of all three models can be seen as practically negligible and hence ID, VC and SC models are virtually indistinguishable at this sample size. However, the RSS of ID model is larger than those of VC and SC models and hence shows ID as not a best fit for the data values.

For an increase of sample size to 50, the parameter estimates of the group means from Table 3 can be observed to be approximately the same for covariance structures of VC and SC, although for ID structure,

Table 2: Simulation results of variance matrices for sample size 25

Variable	Sample size 25		
	ID Model	VC Model	SC Model
β_0	-1.557	6.104	8.935
β_1	3.081	2.256	1.973
μ_1	23.09	24.16	24.72
μ_2	24.63	25.28	25.71
μ_3	26.17	26.41	26.69
μ_4	27.71	27.54	27.68
μ_5	29.25	28.67	28.66
RSS	1802	1796	1798
Deviance	148.1	163.3	167.5
DIC	158.95	174.29	178.56
pD	10.87	11.00	11.02

Table 3: Simulation results of variance matrices for sample size 50

Variable	Sample size 50		
	ID Model	VC Model	SC Model
β_0	2.349	5.204	6.03
β_1	2.745	2.373	2.53
μ_1	24.30	25.79	26.24
μ_2	25.68	27.07	27.51
μ_3	27.05	28.36	28.77
μ_4	28.42	29.65	30.04
μ_5	29.79	30.93	31.30
RSS	5508	5166	5081
Deviance	371.1	369.8	369.9
DIC	383.23	382.38	382.58
pD	12.17	12.54	12.71

the difference begins to narrow down closely. This can be seen too, in the regression coefficients for the repeated measures models defined for each covariance structure.

From Table 4 it can be observed that with sample size of 100, the differences in DIC of all three models is seen as practically negligible and their respective RSS are approximately the same, hence ID, VC and SC models are virtually indistinguishable at this sample size of 100. From Table 4, for sample size 100, the

Table 4: Simulation results of variance matrices for sample size 100

Variable	Sample size 100		
	ID Model	VC Model	SC Model
β_0	28.29	28.29	28.22
β_1	-0.238	-0.239	-0.231
μ_1	26.38	26.38	26.37
μ_2	26.26	26.26	26.25
μ_3	26.14	26.14	26.14
μ_4	26.02	26.02	26.02
μ_5	25.90	25.91	25.91
RSS	7288	7293	7295
Deviance	704.5	703.4	703.3
DIC	719.42	718.32	718.19
pD	14.95	14.94	14.94

parameter estimates of the group means and the model coefficients for all three ID, VC and SC can be seen to be approximately the same particularly for the model coefficients for ID and VC covariance structures. The estimates of the model coefficients for SC are only different by about three decimal places of approximation.

Figures 1 to 3 show the interaction plots of the three model comparison statistics for the three variance matrices ID, VC and SC. This gives a visual summary of the effect of the different variance structure for the

Wishart distribution set as prior.

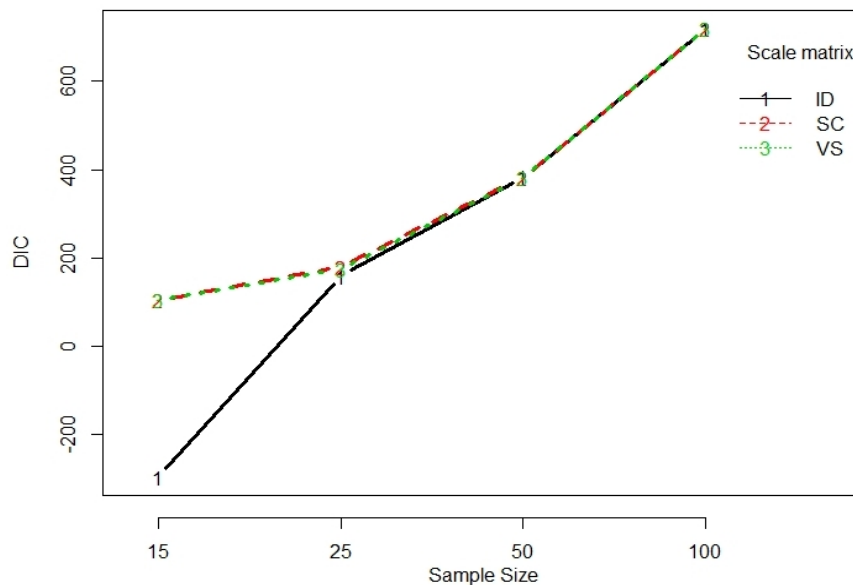


Figure 1: Interaction plots of DIC for variance structures by sample sizes

As can be observed in Figure 1, the DIC of model VC and model SC overlap, while DIC of model ID overlaps with the others as sample size increases but has a higher magnitude of DIC with small sample size. From Figure 2, the residuals sum of squares of model VC and model SC are very similar and that of model ID only gets close to those of VC and SC as sample size increases. Although at some sample size, residual sum of squares of model ID was slightly higher than those of model VC and model SC. From Figure 3,

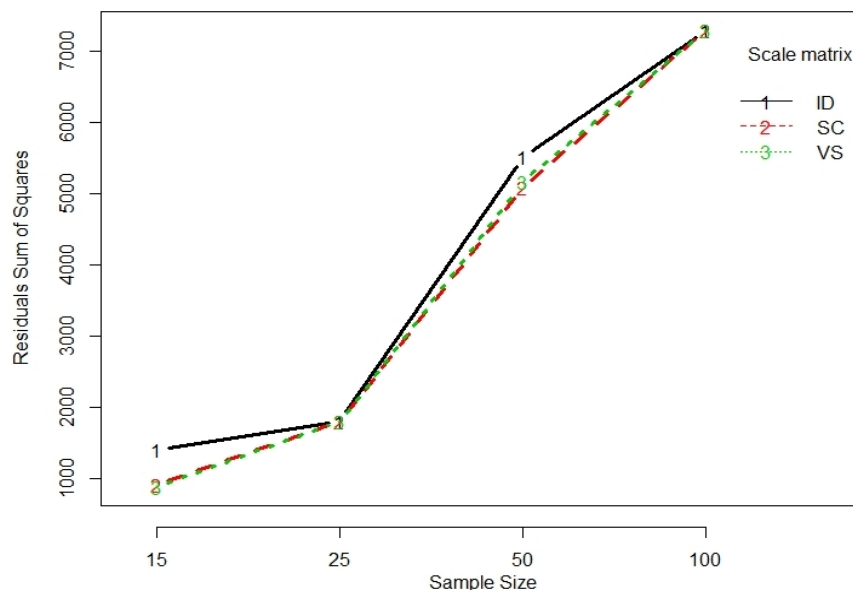


Figure 2: Interaction plots of RSS for variance structures by sample sizes

the deviance of model with ID variance structure is lower than the deviances for models with VC and SC covariance structures. But as been the case for residuals sum of squares and DICs of the different variance structures in large sample sizes, the deviance of model ID, model VC and model SC all approximately converge to same values.

From the foregoing, the simulation output showed that the results with ID matrix was substantially different from those of VC and SC matrices with small sample sizes and as sample size increased, the result with ID matrix coincided with those of VC and SC matrices, which both showed only slight variations in their results for any given sample size. Hence, it is safe to say that the ID matrix is very much influenced by sample size.

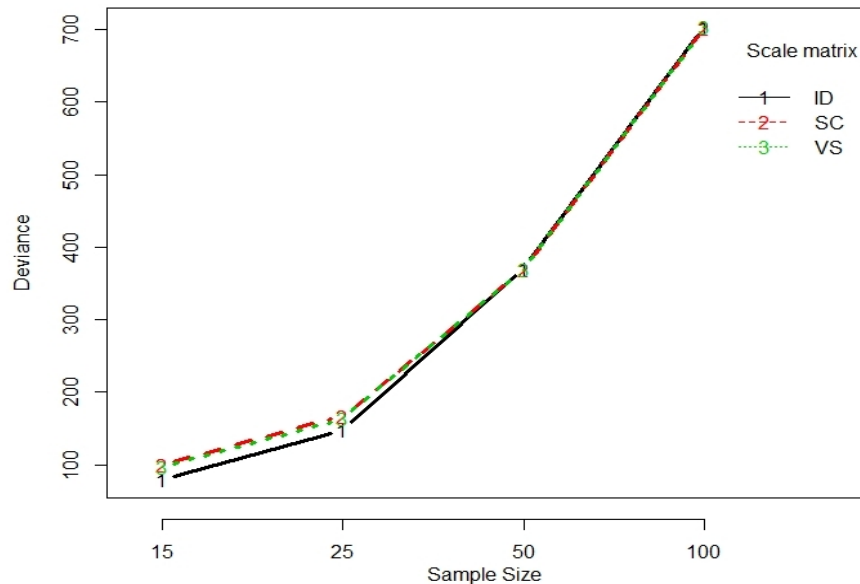


Figure 3: Interaction plots of deviance for variance structures by sample sizes

3.2 Analysis of real data set

Figure 4 shows the interaction plots of the data on the number of leaves over the time period for the different growth media. The object is to see if the difference in the number of leaves produced over time is significantly different for the different growth media.

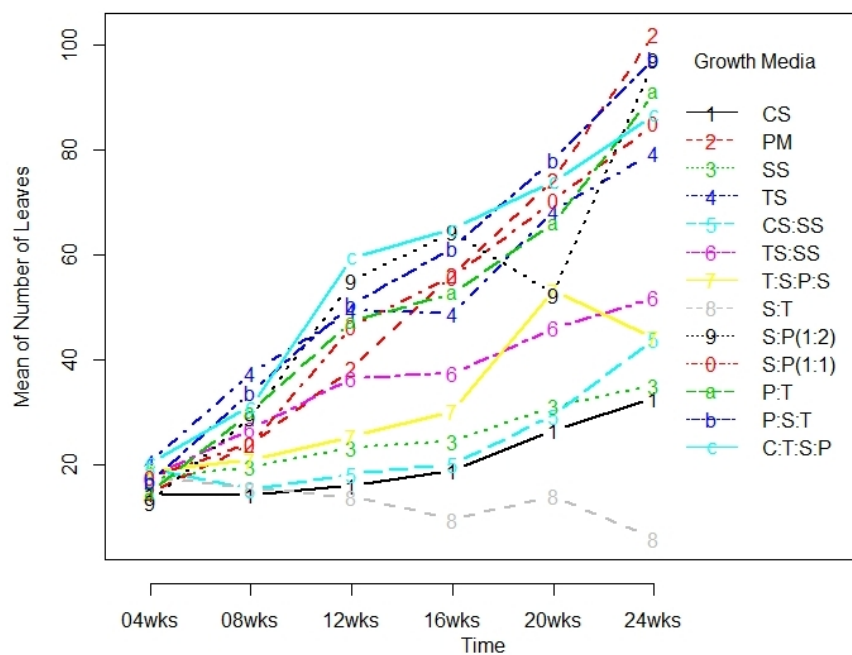


Figure 4: Interaction plots of number of leaves by growth media

Table 5 gives the classical ANOVA output of the estimates of the treatment effects and p-values and the corresponding effect size. It can be seen that at 5% significance level, the growth media SS, CS:SS and S:T were considered not to be significant, as their p-values are greater than 0.05 and their effect sizes are very small.

Table 6 shows the Markov Chain Monte Carlo (MCMC) simulation means from the posterior distribution given the data set obtained. It can be observed that the result from this prior specification of Gamma prior on the conditional error variance is similar to the results of the classical approach given in Table 5.

Table 7 shows the MCMC Gibbs sampling posterior estimates of the parameters from WinBUGS on the real data with the Wishart distribution on the scale matrix as defined for the three covariance structures ID,

Table 5: Repeated measures ANOVA of data set

	Estimate	Std. error	df	t value	p-value	Effect size
Intercept	20.504	7.704	6.327	2.662	0.03559	
PM	31.148	3.780	294	8.239	< 0.0001	Large
SS	4.629	3.780	294	1.224	0.22175	Very small
TS	30.154	3.780	294	7.976	< 0.0001	Large
CS:SS	3.900	3.780	294	1.032	0.30310	Very small
TS:SS	15.513	3.780	294	4.104	< 0.0001	Medium
T:S:P:S	11.598	3.780	294	3.068	0.00236	Small
S:T	-7.790	3.780	294	-2.061	0.04022	Small
S:P(1:2)	31.327	3.780	294	8.286	< 0.0001	Large
S:P(1:1)	29.373	3.780	294	7.770	< 0.0001	Large
P:T	29.788	3.780	294	7.879	< 0.0001	Large
P:S:T	35.924	3.780	294	9.502	< 0.0001	Large
C:T:S:P	35.549	3.780	294	9.403	< 0.0001	Large

VC and SC.

Table 6: MCMC posterior samples of data set

Treatment	Mean	SD	SE	2.5	97.5
Intercept	20.546	4.281	0.02708	23.404	29.029
PM	31.126	6.109	0.03864	35.248	43.049
SS	4.541	6.099	0.03857	8.617	16.497
TS	30.098	6.073	0.03841	34.163	42.117
CS:SS	3.888	6.096	0.03856	8.007	15.811
TS:SS	15.461	6.040	0.03820	19.549	27.268
T:S:P:S	11.544	6.047	0.03824	15.637	23.411
S:T	-7.811	6.058	0.03831	-3.710	4.019
S:P(1:2)	31.295	6.065	0.03836	35.372	43.129
S:P(1:1)	29.348	6.108	0.03863	33.420	41.312
P:T	29.726	6.063	0.03835	33.830	41.543
P:S:T	35.884	6.017	0.03805	39.948	47.755
C:T:S:P	35.537	6.052	0.03827	39.622	47.425
σ^2	446.824	39.936	0.23360	470.194	524.954

Table 7 shows the MCMC Gibbs sampling posterior estimates of the parameters from WinBUGS on the real data with the Wishart distribution on the scale matrix as defined for the three covariance structures ID, VC and SC.

From Table 7, the mean values of the regression parameters and the posterior means are different for the three covariance structures of the scale matrix. The SC covariance structure on the scale matrix fit the data best than the ID and VC covariance structures as the residual sum of squares, RSS and deviance information criterion, DIC for SC is lower than those of ID and VC. Also it can be observed from Table 7 that the effective number of parameters, $pD = 37.73$, for the SC model is smaller than those of ID and VC models.

From Table 7, it can be observed that the posterior means of treatments from the SC model were in reverse order of values unlike those from ID and VC models. Parameter estimates from ID and VC models showed that the mean of four growth media were different from the rest growth media, hence these four growth media that produced less number of seedling leaves.

4. Conclusion

The effects of the various variance structures on the Wishart distribution for the covariance matrix for parameters of repeated measures design were compared under three types of variance structures, namely the Variance Component (VC), Compound Symmetry (SC) and Identity Matrix (ID) using the Bayesian approach. The results from the simulation study indicated that the data-based VC and SC prior specifications for the variance matrix of the fixed parameters performed better compared to the ID specification only in

Table 7: Bayesian Posterior Means with Wishart Prior for Data Set

Variable	Node Statistics		
	ID Model	VC Model	SC Model
β_0	254.7	77.18	15.63
β_1	-26.04	-7.19	0.58
μ_1	228.60	70.0	16.20
μ_2	202.60	62.81	16.78
μ_3	176.60	55.62	17.36
μ_4	150.50	48.43	17.93
μ_5	124.50	41.25	18.51
μ_6	98.43	34.06	19.09
μ_7	72.39	26.87	19.66
μ_8	46.34	19.66	20.24
μ_9	20.30	12.50	20.82
μ_{10}	-5.74	5.31	21.39
μ_{11}	-31.78	-1.88	21.97
μ_{12}	-57.82	-9.06	22.55
μ_{13}	-83.86	-16.25	23.12
RSS	2482000	368900	76460
Deviance	830.20	551.50	541.80
DIC	-575188	-4106.66	579.53
pD	-575518	-4528.18	37.73

small sample size situations, which is the case in most repeated measures designs. The VC and SC specifications performed well because they are based on estimates of the variances from the data, though it should be noted that the difference in the covariance structures of VC and SC is that VC is the diagonal matrix of variance estimates of the groups in the data, while SC is triangular matrix of 1s leading diagonals and off-diagonals being a constant correlation. For the three models it was observed that for most parameters, the estimates were quite similar across the different covariance structure prior specifications. This was seen by the values of the diagnostics statistics of each model, the deviance, deviance information criterion (DIC), and residual sum of squares (RSS). As would be expected, the largest differences were found when the sample size was small, but when the sample sizes were large the different covariance structures converged to same results. For small sample sizes, models with VC and SC priors were better than the common ID based prior and hence, the ID prior should be seldomly used or carefully evaluated before usage in small sample cases. While the choice of VC or SC in small sample size is a matter of the degree of standard error one is willing to make in the experiment. The results from the analysis of the real data, the SC covariance structure on the scale matrix fits the data best than the ID and VC covariance structures as the residual sum of squares, RSS and deviance information criterion, DIC for SC is lower than those of ID and VC and showed that the mean number of leaves produced by growth media CS, SS, CS:SS and S:T were lower than the other growth media. The study showed that the Bayesian approach proved suitable in small sample size experiments which are common for repeated measure designs since the effects of the different variance structures were significantly different in small sample size situations. Also it can be noted from the study that without assuming constant correlation between pairs of measurements within each random factor in a repeated measures design (as in the case for classical approach), Bayesian approach gives a platform for which different appropriate correlations based on data actually collected can be incorporated in estimating the true effects of the parameters of interest.

References

- Bolstad W.M. (2007) Introduction to Bayesian statistics, Wiley and Sons, USA.
- Campbell, G. (2010). Guidance for the use of Bayesian statistics in medical device clinical, trials. Guidance for Industry and FDA Staff, 2010, pp. 15-20.
- Chaloner, K. and Verdinelli, I. (1995). Bayesian experimental design. A review source: Statistical Science, vol. 10, pp. 273-304, April, 1995.
- Field, A. (2016). Repeated Measures ANOVA: issues with Repeated Measures Designs. www.discoveringstatistics.com
- Gelman, A., and Hill, L.S. (2007). Data analysis using regression and multilevel/hierarchical models. New York, NY: Cambridge University Press.

- Gelman, A., Carlin, J.B., Stern, H.S., Dunson, D.B., Vehtari, A., and Rubin, D.B. (2013). *Bayesian data analysis* (3d ed.). Boca Raton, FL: Chapman and Hall/CRC.
- Gelman, A. (2014). Analysis of variance: why it is more important than ever. *Annals of Statistics*, vol. 33, pp 1-53, February. 2014.
- Ghosh, J., Delampady, M. and Samanta, T. (2006). *An Introduction to Bayesian Analysis: Theory and Methods*. Springer Science+Business Media, NYC 10013, USA.
- Grzenda, W (2015). The Advantages of Bayesian Methods over Classical Methods in the context of credible intervals. *Information Systems Management* (2015), Vol. 4 (1), pages 53-63
- Kogan, C., Kalachev, L., and Van Dongen, H. P. A. (2016). Prediction Accuracy in Multivariate Repeated-Measures Bayesian Forecasting Models with Examples Drawn from Research on Sleep and Circadian Rhythms. *Computational and Mathematical Methods in Medicine* Volume 2016, Article ID 4724395, 23 pages. Hindawi Publishing Corporation <http://dx.doi.org/10.1155/2016/4724395>
- Lunn, D.J., Thomas, A., Best, N., and Spiegelhalter, D.J. (2000). WinBUGS – a Bayesian modelling framework: Concepts, structure, and extensibility. *Statistics and Computing*, 10, 325–337. doi: doi:10.1023/A:1008929526011
- Luo, S. and Wang, J. (2015). Bayesian Hierarchical Model for Multiple Repeated Measures and Survival Data: an application to Parkinson's Disease. *Stat Med*. 2015 October 30; 33(24): 4279–4291. doi:10.1002/sim.6228
- Olofsen, E., Dinges, D. F. and Van Dongen, H. P. A. (2004). Nonlinear Mixed-effects Modeling: Individualization and Prediction. *Aviation, Space, and Environmental Medicine*, vol. 75, no. 3, pp. A134–A140, 2004.
- PSU EDU (2018). STAT 502, Introduction to Repeated Measures. Retrieved from <http://Onlinecourses.science.psu.edu/statprogram/stat502>
- Schuurman, N. K., Grasman, R. P. P. P., and Hamaker, E. L. (2016). A Comparison of Inverse-Wishart Prior Specifications for Covariance Matrices in Multilevel Autoregressive Models. *Multivariate Behavioral Research*, 51(2-3), 185-206. <https://doi.org/10.1080/00273171.2015.1065398>
- Song, Y., Nathoo, F. S. and Masson, M. E. J. (2017). A Bayesian approach to the mixed-effects analysis of accuracy data in repeated-measures designs. *Journal of Memory and Language* Vol. 96 (2017) 78–92
- Sturtz, S., Ligges, U., and Gelman, A. (2005). R2WinBUGS: A package for running WinBUGS from R. *Journal of Statistical Software*, 12, 1–16.

Appendix

REAL DATA SET

Treatment	Rep	NL4	NL8	NL12	NL16	NL20	NL24
CS	1	13	13.1	12.6	13	20.4	27
PM	1	19.6	25	29.8	57.6	73.6	100.2
SS	1	20.2	21.8	22.2	20.6	31.6	33
TS	1	19.8	42	45.2	43.4	63.8	72
CS:SS	1	22.4	9.8	16	17.3	23	28.3
TS:SS	1	22.4	30.8	36	31	48	44.6
T:S:P:S	1	19.8	18	27	33	79	40.3
S:T	1	16.4	18.2	16.2	14.8	38.5	8.8
S:P(1:2)	1	13.5	25.6	60	82	32	89.3
S:P(1:1)	1	22.4	33.8	51.4	68	87.7	74.8
P:T	1	14.4	39.2	57.4	71.8	77	97
P:S:T	1	16.8	39	53	46.4	88.4	94.8
C:T:S:P	1	21.6	31.2	57	71.2	73.2	71
CS	2	15	16.2	14.4	26	38.8	40.4
PM	2	12.6	23.2	45	54.3	81.5	98.8
SS	2	17	16.6	25.2	28.6	30.8	37
TS	2	20.8	41.8	55.6	26	83	77.4
CS:SS	2	20.2	14.8	17.2	21	31.2	38.8
TS:SS	2	18	20.4	30.4	31.6	43.4	41.8
T:S:P:S	2	20.8	27.6	28	29	40	52.2
S:T	2	17.2	15.8	11.6	9	7.2	7.2
S:P(1:2)	2	11	31.6	48.6	49.8	60	97.6
S:P(1:1)	2	18	20	48.2	65.6	91	101.6
P:T	2	13.6	28.2	70.3	42.4	65.2	92.8
P:S:T	2	17.4	32.4	11.2	57.8	72.6	96.6
C:T:S:P	2	19.4	31.8	52.4	64.2	74.8	92.2
CS	3	14.2	15.2	19.4	10.6	13.8	17
PM	3	9.76	29	44.8	53.4	69.8	102.6
SS	3	16	23	28.6	31.8	32.4	38
TS	3	20.8	27	9.6	62	58	85.4
CS:SS	3	17.8	20.2	57.6	20	31.2	67.8
TS:SS	3	16.2	29.2	44.8	47.8	46.6	53.2
T:S:P:S	3	17.2	20.4	28.6	30	30	39.4
S:T	3	19	11.2	9.6	7.2	5	3.8
S:P(1:2)	3	10.9	32	57.6	58.4	58.8	97.2
S:P(1:1)	3	16	18.4	35.8	39.6	51.2	81.8
P:T	3	16.6	25.2	53.6	49.8	65.2	83.2
P:S:T	3	16.6	27.2	56.8	74	76.5	94.6
C:T:S:P	3	18.2	29.2	57.6	66	82.4	94.6
CS	4	16	12.6	25.8	26	33.4	46.4
PM	4	16.8	18.5	26.8	60	72.8	106.8
SS	4	19	16.6	17.4	17	29.8	32.2
TS	4	20.6	39	43.6	64	68.6	82.6
CS:SS	4	16.6	16.6	20.4	22.4	32	41.4
TS:SS	4	13.2	27	35	39.2	46	67.8
T:S:P:S	4	17.2	17.6	18.2	28.6	63.8	44.8
S:T	4	17.2	17	18	7.4	5.2	3.7
S:P(1:2)	4	16	27	54.4	67.4	58.8	104.4
S:P(1:1)	4	14	23.8	50	50.2	51.6	82.2
P:T	4	13.8	28	32	47	56.6	91.2
P:S:T	4	16.6	37	45	67.4	75.2	104
C:T:S:P	4	22.2	32.4	53	58.2	65	88.6