

Multiple imputation for multilevel data with continuous and binary variables

November 28, 2017

Vincent Audigier^{1,2,3} ^{*}, Ian R. White^{4,5}, Shahab Jolani⁶, Thomas P. A. Debray^{7,8}, Matteo Quartagno⁹, James Carpenter⁹, Stef van Buuren¹⁰, Matthieu Resche-Rigon^{1,2,3}

1. Service de Biostatistique et Information Médicale, Hôpital Saint-Louis, AP-HP, Paris, France
2. Université Paris Diderot - Paris 7, Sorbonne Paris Cité, UMR-S 1153, Paris, France
3. INSERM, UMR 1153, Equipe ECSTRA, Hôpital Saint-Louis, Paris
4. MRC Biostatistics Unit, Cambridge Institute of Public Health, U.K.
5. MRC Clinical Trials Unit at UCL, London, U.K.
6. Department of Methodology and Statistics, School of Public Health and Primary Care, Maastricht University, Maastricht, The Netherlands
7. Julius Center for Health Sciences and Primary Care, University Medical Center Utrecht, Utrecht, The Netherlands
8. Cochrane Netherlands, University Medical Center Utrecht, Utrecht, The Netherlands
9. Department of Medical Statistics, London School of Hygiene & Tropical Medicine, London, U.K.
10. Department of Statistics, TNO Prevention and Health, Leiden, The Netherlands

^{*}vincent.audigier@univ-paris-diderot.fr

Abstract

We present and compare multiple imputation methods for multilevel continuous and binary data where variables are systematically and sporadically missing.

The methods are compared from a theoretical point of view and through an extensive simulation study motivated by a real dataset comprising multiple studies. Simulations are reproducible. The comparisons show why these multiple imputation methods are the most appropriate to handle missing values in a multilevel setting and why their relative performances can vary according to the missing data pattern, the multilevel structure and the type of missing variables.

This study shows that valid inferences can only be obtained if the dataset gathers a large number of clusters. In addition, it highlights that heteroscedastic MI methods provide more accurate inferences than homoscedastic methods, which should be reserved for data with few individuals per cluster. Finally, the method of [Quartagno and Carpenter \(2016a\)](#) appears generally accurate for binary variables, the method of [Resche-Rigon and White \(2016\)](#) with large clusters, and the approach of [Jolani et al. \(2015\)](#) with small clusters.

Keywords Missing data, Systematically missing values, Multilevel data, Mixed data, Multiple imputation, Joint modelling, Fully conditional specification.

1 Introduction

When individual observations are nested in clusters, statistical analyses generally need to reflect this structure; this is usually referred to as a multilevel structure, where individuals constitute the lower level, and clusters the higher level. This situation often arises in fields like survey research, educational science, sociology, geography, psychology and clinical studies among others.

Missing data affect nearly any dataset in most of these fields, and multilevel data present specific patterns of missing values. Some variables may be fully non-observed for some clusters because they were not measured, or because they were not defined consistently across clusters. [Resche-Rigon et al. \(2013\)](#) named such missing data patterns *systematically missing values*. Nowadays, systematically missing values are becoming increasingly common because of a greater availability of data coming from several sources, including different sets of variables ([Riley et al., 2016](#)). Since systematically missing values are related to the multilevel structure, they affect any field where multilevel data occur. Examples of multilevel data with systematically missing values include [Bos et al. \(2003\)](#), [Mullis et al. \(2003\)](#) or [Blossfeld et al. \(2011\)](#) in educational sciences, [GREAT Network \(2013\)](#), [Kunkel and Kaizar \(2017\)](#) in medicine, [Carrig et al. \(2015\)](#) in sociology. As opposed to systematically missing values, *sporadically missing values* are missing data specific to each individual observation, like in the common single-level framework. It is often the case that both types of missing data occur in a multilevel dataset. For instance, in educational research, questionnaires may be too long to be administered to all students and, therefore,

only a subset of items may be asked to each class, leading to systematically missing values. In addition, some students in each class may not answer some questions, leading to sporadically missing values.

Multiple imputation (MI) is a common strategy to deal with missing values in statistical analysis (Schafer, 1997; Rubin, 1987; Little and Rubin, 2002). It involves firstly specifying a distribution in accordance with the data, the *imputation model*, under which M imputations are drawn from their posterior predictive distribution given the observed values. Thus, M complete datasets are generated. Secondly, a standard statistical analysis is performed on each imputed dataset, leading to M estimates of the *analysis model's* parameters. Finally, the estimates are pooled according to Rubin's rules (Rubin, 1987). The standard assumption when using MI is ignorability (Schafer, 1997, p. 11), implying that missing values occur *at random* (Rubin, 1976), *i.e.* the probability of missingness depends solely on observed data. Several MI methods have been proposed, differing mainly in the assumed form of the imputation model; among methods assuming a parametric imputation model, two strategies are the most common. We refer to joint modelling (JM) imputation when a multivariate joint distribution is specified for all variables. Alternatively, we refer to fully conditional specification (FCS) when a conditional distribution is defined for each incomplete variable (van Buuren et al., 2006; Raghunathan et al., 2001).

In the standard statistical framework where data are complete, multilevel data induce non-independence between observations and require dedicated analysis models accounting for this source of dependency. The linear mixed effect model is one of such models. In the same way, with missing values, imputation models need to take into account dependency between observations, or otherwise the variance of prediction of missing values cannot be properly reflected. Indeed, when an incomplete variable is part of a linear mixed effect analysis model, but it is imputed ignoring the multilevel structure, the imputed values are too close to their expectation. Thus, biases can occur even when applying appropriate statistical methods on inappropriately imputed data.

To account for the multilevel structure with sporadically missing values, imputation models are generally based on regression models including a fixed or a random intercept for cluster (Drechsler, 2015). Methods using a fixed intercept treat the identifier of each cluster as a dummy variable. They are generally parametric, using normal regression for instance, but imputation according to semi-parametric method can also be relevant for complex datasets (Vink et al., 2015; Little, 1988). However, using a random intercept is generally preferable because fixed intercept inflates the true variability between clusters (Andridge, 2011; Graham, 2012). MI methods using a random intercept include JM-pan (Schafer and Yucel, 2002), JM-jomo (Quartagno and Carpenter, 2016a), JM-REALCOM (Goldstein et al., 2007, 2009; Carpenter and Kenward, 2013), JM-Mplus (Asparouhov and Muthén, 2010), JM-RCME (Yucel, 2011), FCS-pan (Schafer and Yucel, 2002), FCS-blimp (Enders et al., 2017), FCS-2lnorm (van Buuren, 2010), FCS-GLM (Jolani et al., 2015; Jolani, 2017), FCS-2stage (Resche-Rigon and White, 2016). These methods differ by the form of the imputation model (joint or not), but also by their ability to account for different types

of variables (continuous, binary, or others). In particular, JM-pan, JM-RCME, FCS-pan, FCS-2lnorm are not dedicated to binary variables.

Systematically missing values imply identifiability issues for imputation models which include a fixed intercept for cluster. Thus, only methods using random effects are appealing. Most of the first MI methods proposed to impute multilevel data were based on random intercept models, but systematically missing values were not considered. Therefore, not all of these methods are tailored for the imputation of such missing data. Table 1 summarizes the MI methods available for multilevel data.

Table 1: Summary of MI methods’ properties for multilevel data based on random intercept (JM-pan (Schafer and Yucel, 2002), JM-REALCOM (Goldstein et al., 2007, 2009; Carpenter and Kenward, 2013), JM-jomo (Quartagno and Carpenter, 2016a), JM-Mplus (Asparouhov and Muthén, 2010), JM-RCME (Yucel, 2011), FCS-pan (Schafer and Yucel, 2002), FCS-blimp (Enders et al., 2017) FCS-2lnorm (van Buuren, 2010), FCS-GLM (Jolani et al., 2015), FCS-2stage (Resche-Rigon and White, 2016))

Method (form - name)	Handles missing data:				Coded in R (R Core Team, 2016)
	Sporadic?	Systematic?	in continuous variable?	in binary variable?	
JM-pan	yes	yes	yes	no	yes, package <i>pan</i>
JM-REALCOM	yes	yes	yes	yes	no
JM-jomo	yes	yes	yes	yes	yes, package <i>jomo</i>
JM-Mplus	yes	yes	yes	yes	no
JM-RCME	yes	yes	yes	no	no
FCS-pan	yes	yes	yes	no	yes, package <i>mice</i>
FCS-blimp	yes	yes	yes	yes	no
FCS-2lnorm	yes	no	yes	no	yes, package <i>mice</i>
FCS-GLM	yes ¹	yes	yes	yes	yes, add-on for <i>mice</i>
FCS-2stage (REML or MM)	yes	yes	yes	yes ¹	yes, add-on for <i>mice</i>

¹ using variant reported in this paper

In this paper, we compare the most relevant MI methods for dealing with clustered datasets with systematically and sporadically missing variables, continuous and binary. Among JM methods for multilevel data, we focus on the JM-jomo method proposed in Quartagno and Carpenter (2016a), which can be

seen as a generalisation of JM-REALCOM and JM-Mplus, while among FCS methods, we focus on the FCS-GLM method presented in [Jolani et al. \(2015\)](#) and on the FCS-2stage method proposed in [Resche-Rigon and White \(2016\)](#). We do not focus on FCS-blimp which can be seen as a univariate version of JM-jomo.

The paper is organised as follows. Firstly, we present the three MI methods for handling multilevel data with systematically and sporadically missing values (Section 2). Both FCS methods have theoretical deficiencies in this general setting, so we propose improvements for them: accounting for binary variables in FCS-2stage, and accounting for continuous sporadically missing values in FCS-GLM. Secondly, these MI methods are compared through a simulation study (Section 3). Thirdly, MI methods are applied to a real data analysis (Section 4). Finally, practical recommendations are provided (Section 5).

2 Multiple imputation for multilevel continuous and binary data

2.1 Methods

2.1.1 Univariate missing data pattern

Random variables will be indicated in italics, while fixed values will be denoted in roman letters. Vectors will be in lower case, while matrices will be in upper case. Let $\mathbf{Y}_{n \times p} = (\mathbf{y}_1, \dots, \mathbf{y}_p)$ be an incomplete data matrix for n individuals in rows and p variables in columns. Let i be the index for the individuals ($1 \leq i \leq n$) and j for the columns ($1 \leq j \leq p$). $\mathbf{Y}$ is stratified into K clusters of size n_k where k denotes the index for the cluster ($1 \leq k \leq K$). $\mathbf{y}_{jk}$ denotes the n_k -vector corresponding to the vector $\mathbf{y}_j$ restricted to individuals within cluster k . Let $(\mathbf{y}_j^{\text{obs}}, \mathbf{y}_j^{\text{miss}})$ be the missing and observed parts of $\mathbf{y}_j$ and let $\mathbf{Y}^{\text{obs}} = (\mathbf{y}_1^{\text{obs}}, \dots, \mathbf{y}_p^{\text{obs}})$ and $\mathbf{Y}^{\text{miss}} = (\mathbf{y}_1^{\text{miss}}, \dots, \mathbf{y}_p^{\text{miss}})$.

In order to propose a unified presentation of the three MI methods, we assume in this section that the variable $y = y_p$ is the only incomplete variable and is continuous. Extension to several incomplete variables, continuous or binary, will be discussed in the next section.

The imputation step in MI aims to draw missing values from their predictive distribution $P(Y^{\text{miss}}|Y^{\text{obs}})$. To achieve this goal, an imputation model with parameter $\boldsymbol{\theta}$ is specified and realisations of the predictive distribution of missing values can be obtained by:

Step (1) drawing $\boldsymbol{\theta}$ from $P(\boldsymbol{\theta}|Y^{\text{obs}})$, its posterior distribution

Step (2) drawing missing data according to $P(Y^{\text{miss}}|Y^{\text{obs}}, \boldsymbol{\theta})$, their predictive distribution for a given $\boldsymbol{\theta}$.

For a single continuous incomplete variable, the posterior distribution can be specified by letting $\boldsymbol{\theta} = (\boldsymbol{\beta}, \boldsymbol{\Psi}, (\boldsymbol{\Sigma}_k)_{1 \leq k \leq K})$ be the parameters of a linear mixed effects model:

$$\begin{aligned} y_k &= \mathbf{Z}_k \boldsymbol{\beta} + \mathbf{W}_k b_k + \varepsilon_k \\ b_k &\sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Psi}) \\ \varepsilon_k &\sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_k) \end{aligned} \tag{1}$$

where y_k denotes the incomplete variable restricted to the cluster k , $\mathbf{Z}_k$ ($n_k \times q$) and $\mathbf{W}_k$ ($n_k \times q'$) are the known covariate matrices corresponding to two subsets of $(\mathbf{y}_{1k}, \dots, \mathbf{y}_{(p-1)k})$, $\boldsymbol{\beta}$ is the q vector of regression coefficients of fixed effects, b_k is the q' vector of random effects for cluster k , $\boldsymbol{\Psi}$ ($q' \times q'$) is the between cluster variance matrix, $\boldsymbol{\Sigma}_k = \sigma_k^2 \mathbb{I}_{n_k}$ ($n_k \times n_k$) is the variance matrix within cluster k . Model (1) is the imputation model used in FCS-GLM and FCS-2stage and potentially in JM-jomo for the case of a univariate missing data pattern.

Drawing parameters of the imputation model from their posterior distribution (Step (1)) can be achieved by several approaches (Little and Rubin, 2002, p200-222). A first approach uses explicit Bayesian modelling of (1), specifying a prior distribution for $\boldsymbol{\theta}$ and drawing from its posterior distribution. This approach is used in FCS-GLM and in JM-jomo: FCS-GLM uses a non-informative Jeffreys prior distribution, while JM-jomo uses a conjugate prior distribution.

A second approach uses the asymptotic distribution of a frequentist estimator of $\boldsymbol{\theta}$. More precisely, the parameters of this distribution are estimated from the data, and a value of $\boldsymbol{\theta}$ is drawn from this asymptotic distribution. FCS-2stage is based on this principle: the estimator used is called the *two-stage estimator* in IPD meta-analysis (Simmonds et al., 2005; Riley et al., 2008). It is also possible to use the Maximum Likelihood (ML) estimator (Resche-Rigon et al., 2013), but the two-stage estimator has the advantage of being easier and quicker to compute than the ML estimator for linear mixed effects models.

When the variable y is only sporadically missing, the posterior distribution of the parameters only involves individuals that are observed (Rubin, 1987, p. 165), so that both approaches easily handle missing data. However, systematically missing values complicate Step (1) for both approaches. Using Bayesian modelling, simulating the posterior distribution of $\boldsymbol{\theta}$ generally requires a Gibbs sampler (Geman and Geman, 1984), but the posterior distribution of $\boldsymbol{\Sigma}_k$ cannot be updated from the data at each iteration for systematically missing clusters. Similarly, $\boldsymbol{\Sigma}_k$ cannot be estimated from observed data by using the asymptotic method. Thus, MI methods developed for sporadically missing data cannot be directly used to impute systematically missing data. The problem is overcome by assuming a distribution across $(\boldsymbol{\Sigma}_k)_{1 \leq k \leq K}$, as proposed in JM-jomo and in FCS-2stage, or by assuming $\boldsymbol{\Sigma}_k = \boldsymbol{\Sigma}$ for all k , as proposed in FCS-GLM.

To draw missing values according to the parameters drawn at Step (1), missing values are predicted according to model (1) and Gaussian noise is added to

the prediction (Step (2)). However, the random coefficients are not strictly parameters of this model and, therefore, are not directly given by Step (1). Thus, to obtain realisations for $(b_k)_{1 \leq k \leq K}$, each random coefficient is drawn from its distribution conditional on y_k^{obs} and the parameters generated from Step (1). Imputation can then be performed. If data are sporadically missing, these conditional distributions are derived by classic calculation for Gaussian vectors; if data are systematically missing, random coefficients are drawn from their marginal distribution.

From this unified presentation, we now present the methods for several incomplete variables, which can also be binary.

2.1.2 Multivariate missing data pattern

JM-jomo (Quartagno and Carpenter, 2016a). To multiply impute multilevel data with several incomplete continuous variables, JM approaches are based on the multivariate version of model (1) where covariate matrices are matrices $(i_k \times p)$ of ones:

$$\begin{aligned} Y_k &= \mathbf{1}\beta + \mathbf{1}b_k + \varepsilon_k \\ b_k^V &\sim \mathcal{N}(0, \Psi) \\ \varepsilon_k^V &\sim \mathcal{N}(0, \Sigma_k) \end{aligned} \tag{2}$$

β is the $1 \times r$ matrix of regression coefficients of fixed effects, and b_k is the $1 \times r$ matrix of random effects. The superscript V indicates the vectorisation of a matrix by stacking its columns. Ψ ($r \times r$) is the between cluster variance matrix and Σ_k ($rn_k \times rn_k$) is the block diagonal variance matrix within cluster k .

Note that model (2) includes all variables on the left hand side of imputation model ($Y_k = (Y_k^{\text{miss}}, Y_k^{\text{obs}})$). Other modelling could be considered by including complete variables can be included on the left or right hand side. The proposed model has the advantage to limit overfitting when the number of variables is small compared to the number of individuals (Quartagno and Carpenter (2016a)).

To perform MI according to this imputation model, a Bayesian approach is used with the following independent prior distributions for $\theta = (\beta, \Psi, (\Sigma_k)_{1 \leq k \leq K})$:

$$\beta \propto 1 \tag{3}$$

$$\Psi^{-1} \sim \mathcal{W}(\nu_1, \Lambda_1) \tag{4}$$

$$\Sigma_k^{-1} | \nu_2, \Lambda_2 \sim \mathcal{W}(\nu_2, \Lambda_2) \text{ for all } 1 \leq k \leq K \tag{5}$$

$$\nu_2 \sim \chi^2(\eta), \Lambda_2^{-1} \sim \mathcal{W}(\nu_3, \Lambda_3) \tag{6}$$

where $\mathcal{W}(\nu, \Lambda)$ denotes the Wishart distribution with ν degrees of freedom and scale matrix Λ , and η denotes the degrees of freedom of the chi-squared distribution. The prior distributions for the covariance matrices are informative (Gelman, 2006); to make them as vague as possible, hyperparameters are set

as $\nu_1 = r$, $\mathbf{\Lambda}_1 = \mathbb{I}_r$, $\nu_3 = rK$, $\mathbf{\Lambda}_3 = \mathbb{I}_{rK}$ and $\eta = rK$. From this modelling, the posterior distribution of $\boldsymbol{\theta}$ can be derived. We report in Appendix A.1 the derived posterior distributions as well as the technical details to obtain realisations from them.

In summary, the parameters of the imputation model are drawn from their posterior distribution with a multivariate missing data pattern by using a Data-Augmentation (DA) algorithm (Tanner and Wong, 1987): given current values $\boldsymbol{\theta}^{(\ell)}$ for $\boldsymbol{\theta}$ and $Y^{\text{miss}^{(\ell)}}$ for Y^{miss} , the components of $\boldsymbol{\theta}$ are successively updated according to their posterior distribution (Equations 15-17 in Appendix A.1) given $(Y^{\text{obs}}, Y^{\text{miss}^{(\ell)}})$, providing $\boldsymbol{\theta}^{(\ell+1)}$. Then, $\boldsymbol{\theta}^{(\ell+1)}$ can be used to generate $Y^{\text{miss}^{(\ell+1)}}$ according to model (2). To obtain M independent realisations from the posterior distribution, the algorithm is run through a burn-in period (to reach the convergence to the posterior distribution) and then realisations are drawn by spacing them with several iterations (to ensure independence). Note that the number of iterations for the burn-in period and the number of iterations between to realisations need to be carefully checked (Schafer, 1997, p.160-169). Moreover, since generating $\boldsymbol{\theta}$ in its predictive distribution using the DA algorithm also requires imputation of missing data, Step (1) and Step (2) of MI (see Section 2.1.1) are not distinguished here.

This method allows imputation of datasets with systematically and sporadically missing values. In particular, despite systematically missing values, the posterior distribution for $\boldsymbol{\Sigma}_k$ can be updated at each step of the DA algorithm by considering observed values from other clusters.

To deal with binary variables, a probit link and a latent variables framework have been proposed (Goldstein et al., 2009). Let L be the set of continuous variables joined with a set of latent variables corresponding to the binary variables, so that $L = (L^{\text{miss}}, L^{\text{obs}})$. At the end of each cycle of the DA algorithm, given current parameters $\boldsymbol{\theta}^{(\ell)}$ and random coefficients $b^{(\ell)}$, L^{miss} is drawn conditionally on L^{obs} .

Like $P(Y^{\text{miss}}, Y^{\text{obs}})$ for continuous incomplete variables, $P(L^{\text{miss}}, L^{\text{obs}})$ is a multivariate Gaussian distribution. Thus, drawing missing latent variables consists of drawing L from a Gaussian distribution under the positivity or negativity constraint imposed by observed binary values, which is straightforward (Carpenter and Kenward, 2013, p.96-98). Next binary data from Y^{miss} are derived from the previously drawn latent variables: the outcome 1 is drawn if the latent variable takes a positive value, and 0 otherwise.

The JM-jomo method is an extension of the JM-RCME (Yucel, 2011), JM-Mplus (Asparouhov and Muthén, 2010), JM-REALCOM (Goldstein et al., 2007, 2009; Carpenter and Kenward, 2013) and the JM-pan method (Schafer and Yucel, 2002); JM-jomo additionally allows for heteroscedasticity of the imputation model and imputation of binary (and more generally categorical) variables, while

JM-RCME only handles heteroscedasticity, and JM-REALCOM or JM-Mplus only handle categorical variables. Note that the Bayesian formulation of the imputation model is also very close to those used in FCS-blimp, FCS-pan and FCS-2lnorm

FCS-GLM (Jolani et al., 2015). Instead of using a JM approach, fully conditional specification can be used to multiply impute a dataset with several incomplete variables. The principle is to simulate the predictive distribution of the missing values by successively simulating the predictive distribution of the missing values of each incomplete variable conditionally on the other variables. Thus, instead of specifying a joint imputation model as (2), only the conditional distribution of each incomplete variable is required. Compared to JM approaches, FCS approaches make it easier to model complex dependence structures.

Jolani et al. (2015) use a FCS approach to perform multiple imputation of systematically missing variables only. For continuous incomplete variables, the conditional imputation model is model (1) assuming homoscedastic error terms, *i.e.* $\sigma_k = \sigma$ for all k . To draw missing values of y from their predictive distribution, a Bayesian formulation of the univariate linear mixed effects model based on non-informative independent priors is used. Details on the posterior distributions are provided in Appendix A.2, we only underline that they depend of the maximum likelihood estimates of the imputation model's parameters.

Imputation of a systematically missing variable y is performed as follows:

Step (1) θ is drawn according to the posterior (equations (18-20) in Appendix A.2)

Step (2) $P(y^{\text{miss}}|Y^{\text{obs}}, \theta)$ is simulated by:

- drawing b_k from $\mathcal{N}(0, \Psi)$ for all clusters in $1 \leq k \leq K$ where y_k is systematically missing
- drawing y_k^{miss} from $\mathcal{N}(\mathbf{Z}_k\beta + \mathbf{W}_kb_k, \sigma^2\mathbb{I}_{n_k})$ for all clusters in $1 \leq k \leq K$ where y_k is sporadically missing.

Binary variables are imputed in the same way by considering a GLMM model with a logit link.

FCS-GLM was originally developed to impute systematically missing variables only. We extend it to also impute sporadically missing continuous variable following the rationale of Resche-Rigon et al. (2013). To achieve this goal, Step (1) is essentially the same, the main difference lies in Step (2): each b_k is drawn conditionally on y_k^{obs} , instead of being drawn from its marginal distribution. However, when y is a binary variable, this conditional distribution is analytically intractable because of the logit link. Therefore, binary sporadically missing

variables are handled as binary systematically missing ones, which can potentially introduce bias and a lack of variability in the imputed values.

FCS-2lnorm (van Buuren, 2010), FCS-blimp (Enders et al., 2017) and FCS-pan (Schafer and Yucel, 2002) also proposed FCS approaches which use a conjugate prior to reflect the posterior distribution of the parameter of the imputation model. FCS-2lnorm is based on the model (1) as conditional imputation model and thus allows heteroscedasticity of errors for continuous variables. However, it cannot be directly applied to systematically missing clusters because of non-identifiability of (σ_k^2) for $1 \leq k \leq K$. On the contrary, both others assume homoscedasticity only.

FCS-2stage (Resche-Rigon and White, 2016). FCS-2stage is another FCS method drawing the parameters of the imputation model by using an asymptotic strategy: an estimator is evaluated from the observed data and the posterior distribution is then approximated (*cf* Section 2.1.1). This estimator is a two-stage estimator (Simmonds et al., 2005; Riley et al., 2008). Often used in IPD meta-analysis, it has the advantage of being quicker to compute than the usual one-stage estimator required for the previous method (through the expressions of posterior distributions).

More precisely, for a continuous incomplete variable y , the conditional imputation model (1) is re-written as follows:

$$\begin{aligned} y_k &= \mathbf{Z}_k (\boldsymbol{\beta} + b_k) + \varepsilon_k \\ b_k &\sim \mathcal{N}(0, \boldsymbol{\Psi}) \\ \varepsilon_k &\sim \mathcal{N}(0, \sigma_k^2 \mathbb{I}_{n_k}) \end{aligned} \tag{7}$$

Note that for clarity, the method is presented for the case $\mathbf{Z}_k = \mathbf{W}_k$. Extension to the more general imputation model is given in Resche-Rigon and White (2016). The parameter of this model is $\boldsymbol{\theta} = (\boldsymbol{\beta}, \boldsymbol{\Psi}, (\sigma_k)_{1 \leq k \leq K})$.

To fit the two-stage estimator, at stage one, the ML estimator of a linear model is computed on each available cluster:

$$\hat{\boldsymbol{\beta}}_k = \left(\mathbf{Z}_k^\top \mathbf{Z}_k \right)^{-1} \mathbf{Z}_k^\top \mathbf{y}_k \tag{8}$$

Then, at stage two, the following random effects model is used:

$$\hat{\boldsymbol{\beta}}_k = \boldsymbol{\beta} + b_k + \varepsilon'_k \tag{9}$$

with $b_k \sim \mathcal{N}(0, \boldsymbol{\Psi})$ and $\varepsilon'_k \sim \mathcal{N}\left(0, \sigma_k^2 \left(\mathbf{Z}_k \mathbf{Z}_k^\top \right)^{-1}\right)$. $\boldsymbol{\beta}$ and $\boldsymbol{\Psi}$ may be estimated by REML; alternatively, Resche-Rigon and White (2016) suggest using the method of moments (MM), which is even faster, especially with high dimensional $\boldsymbol{\beta}$ (DerSimonian and Laird, 1986; Jackson et al., 2013).

We explain in Appendix A.3 how the asymptotic distribution of such an estimator can be derived with incomplete data, as well as how realisations from this distribution can be obtained. Following such developments, imputation of variable y is performed as follows:

Step (1) θ is drawn according to the asymptotic posterior (Equations 23-26 in Appendix A.3).

Step (2) $y^{\text{miss}}|Y^{\text{obs}}, \theta$ is generated by:

- drawing b_k from $\mathcal{N}(0, \Psi)$ for $1 \leq k \leq K$ if y_k is systematically missing or conditionally on $\hat{\beta}_k$ if y_k is sporadically missing
- drawing y_{ki}^{miss} from $\mathcal{N}(\mathbf{z}_{ki}(\beta + b_k), \sigma_k^2)$ for all k ($1 \leq k \leq K$) and for all i ($1 \leq i \leq n_k$) such that y_{ki} , the observation i in cluster k , is missing.

Originally, this method was dedicated to handle incomplete continuous variables only. We extend it to handle binary variables with both sporadically and systematically missing values by applying a logit link in the imputation model (7). In this case, the two-stage estimator is based on logistic models at stage one. The missing values can be imputed according to a similar scheme as continuous variables.

Table 2 sums up the main modelling assumptions of each MI method. In the next section, the consequences of such differences are highlighted in terms of inference from incomplete data.

Table 2: Synthesis of the modelling assumptions of the MI methods JM-jomo, FCS-GLM and FCS-2stage

	heteroscedasticity assumption	link function for binary variables	strategy for proper MI
JM-jomo	yes	probit	Bayesian modelling based on conjugate prior
FCS-GLM	no	logit	Bayesian modelling based on Jeffrey prior
FCS-2stage	yes	logit	asymptotic method based on a two-stage estimator

2.2 Properties

2.2.1 FCS or JM

Comparisons between FCS and JM methods have been extensively studied (van Buuren, 2007; Lee and Carlin, 2010; Zhao and Yucel, 2009; Wagstaff and Harel, 2011; Kropko et al., 2014; Hughes et al., 2014; Resche-Rigon and White, 2016; Erler et al., 2016), particularly in settings without clustering. It is generally believed that FCS methods are less likely to yield biased imputations because they allow for more flexibility than JM methods. For multilevel data, the lack of flexibility for JM methods has been recently highlighted when the analysis model includes random slopes corresponding to incomplete variables (Enders et al., 2016). However, FCS-methods raised other issues like the variables selection for conditional models. Furthermore, the theoretical background of FCS is not well understood and constitutes a current topic of research (Zhu and Raghunathan, 2015; Liu et al., 2014; Bartlett et al., 2015). Indeed, contrary to a Gibbs sampler, convergence towards a joint posterior distribution cannot generally be proven (?). Nevertheless, simulation show that it does not affect the quality of imputation without clustering (van Buuren et al., 2006). In addition, estimation of conditional distributions is more computationally intensive than the estimation of a joint distribution.

2.2.2 One-stage or two-stage estimator

The two FCS approaches use different estimators of the imputation model: the FCS-GLM method uses the one-stage estimator of parameters of model (1), while the two-stage estimator uses the re-written model (7). The one-stage estimator has the drawback to be computationally intensive and slow to converge (Schafer and Yucel, 2002), particularly with binary variables (Noh and Lee, 2007). The two-stage estimator solves this computational time issue, but tends to have a larger variance (Mathew and Nordstrom, 2010) and requires large clusters with binary outcome to avoid separability problems (Albert and Anderson, 1984) and to reduce the small-sample bias of the ML estimator (Firth, 1993). Furthermore, by using a limited number of observations at stage one, the FCS-2stage method is more prone to suffer of overfitting if the number of covariates or the number of missing values is large.

2.2.3 Heteroscedasticity

JM-jomo and FCS-2stage allow for heteroscedasticity of the imputation model, whereas FCS-GLM assumes homoscedastic error variances. It has previously been demonstrated that data generated from a joint homoscedastic model (similar to model (2)) will yield heteroscedastic conditional distributions (Resche-Rigon and White, 2016). As a result, imputation models allowing for heteroscedasticity tend to yield more reliable imputations. Previous simulation studies seem to support this point (van Buuren, 2010; Resche-Rigon and White, 2016). However, homoscedasticity can be relevant when studies are very small,

since it overcomes overfitting issues by shrinking cluster-specific parameter estimates towards their weighted average.

2.2.4 Bayesian modelling or asymptotic strategy for **Step (1)**

JM-jomo and FCS-GLM consider an explicit Bayesian specification of the imputation model, which implies that uncertainty of θ is fully propagated. Conversely, FCS-2stage only propagates the asymptotic uncertainty, and may therefore be problematic in small samples. Regardless, in large samples, both approaches should yield similar results (Little and Rubin, 2002, p. 216).

As a direct result of Bayesian modelling, JM-jomo and FCS-GLM require the specification of a prior distribution for θ . Various priors have been proposed for hierarchical models such as model (1) (Robert, 2007, p. 456-506). In general, it is recommended to use proper prior distributions when working with multivariate linear mixed effect models (Schafer and Yucel, 2002), as this helps to avoid convergence issues of the Gibbs sampler. To this purpose, JM-jomo considers conjugate prior distributions. A major advantage of using conjugate prior distributions is that drawing from the posterior distribution avoids the systematic recourse to MCMC methods. In particular, when using univariate linear mixed effect model with fully observed covariates, the posterior distribution becomes analytically tractable.

In contrast to JM-jomo, FCS-GLM considers the Jeffreys prior for drawing imputations. This prior is derived from the sampling distribution and can therefore be regarded as non-informative. The counterpart of using Jeffreys prior in GLMM is that the (joint) posterior distribution for θ is generally improper (Natarajan and Kass, 2000) even if each marginal distribution is proper. Note that FCS-GLM overcomes this potential issue by assuming independent posterior distributions for each components of θ .

In conclusion, all methods are likely to yield different posterior distributions, particularly in the presence of small sample sizes.

2.2.5 Binary variables

JM-jomo considers a probit link to model binary variables, while both FCS approaches consider a logit link. Although both link functions tend to yield similar predictions, the probit link is more convenient for imputation purposes in multilevel data. The underlying reason is that conditional distributions of random coefficients can be easily simulated with a probit link, because it is based on latent normal variables for which these conditional distributions are well known, but not for mixed models with a logistic link. As a result, imputation of sporadically missing values is achieved in the same way as systematically missing variables for FCS-GLM, i.e. by ignoring the relationship between the random effects and the observed values on the imputed variables. It implies that this method is not relevant with binary sporadically missing variables including few missing values. Conversely, for FCS-2stage it is still possible to draw random coefficients from the conditional distribution, by considering the distribution

of the random coefficients conditionally on the ML estimates given at stage one. Nevertheless, because of the asymptotic unbiasedness property of the ML estimator for logistic regression models (used at stage one), performances of the FCS-2stage method for binary variables are deteriorated when all clusters contain few observed individuals.

3 Simulations

3.1 Simulation design

We consider a simulation study to assess the relative performance of the MI methods described in Section 2. Based on aforementioned properties, we anticipate that FCS-GLM is problematic when imputing binary sporadically missing variables, that FCS-2stage is problematic in datasets with few participants and/or clusters, and that JM-jomo is sensitive to the proportion of missing values because of the influence of the prior distribution. For this reason, we here evaluate the proportion of systematically and sporadically missing values, the data type of imputed variables, the size of included clusters and the size of the total dataset. Other settings will be also investigated to cover a large range of practical cases. For all investigated configurations, we generated $T = 500$ complete datasets, after which we introduced missing values. Afterwards, we applied the MI methods on each incomplete dataset by considering $M = 5$ complete datasets, and obtained parameter estimates from the multiply imputed datasets.

3.1.1 Data generation

For each simulation, we generated a dataset with four variables (y, x_1, x_2, x_3) ; x_1 and x_3 are continuous variables, and x_2 is a binary variable. The outcome variable y (continuous or binary) is defined according to a GLMM where x_1 and x_2 are the covariates. We use x_3 as an auxiliary variable explaining the missing data mechanism. More precisely, data are simulated as follows:

1. Draw K realisations of the triplet of variables (v_1, v_2, v_3) so that

$$(v_1, v_2, v_3) \sim \mathcal{N}(\mathbf{0}, \mathbf{\Sigma}_v) \quad (10)$$

where $\mathbf{\Sigma}_v$ is a (3×3) covariance matrix.

2. Draw two continuous variables x_1 and x_3 so that:

$$\begin{pmatrix} x_{1ki}, x_{3ki} \end{pmatrix} \sim \mathcal{N}\left(\begin{pmatrix} \alpha_1 + v_{1k}, \alpha_3 + v_{3k} \end{pmatrix}, \mathbf{\Sigma}_x\right) \quad (11)$$

where α_1 and α_3 are the fixed intercepts, v_{1k} and v_{3k} are the random intercepts for cluster k drawn at previous step and $\mathbf{\Sigma}_x$ is a (2×2) covariance matrix.

3. Draw a binary variable x_2 according to the model:

$$\text{logit}(P(x_{2ki} = 1)) = \alpha_2 + v_{2k} \quad (12)$$

where α_2 is a fixed intercept and v_{2k} is the random intercept for cluster k , drawn from (10)

4. Draw a response variable y :

- for a continuous variable y :

$$y_{ki} = \beta^{(0)} + \beta^{(1)}x_{1ki} + \beta^{(2)}x_{2ki} + u_k^{(0)} + u_k^{(1)}x_{1ki} + \varepsilon_{ki} \quad (13)$$

where $\varepsilon_{ki} \sim \mathcal{N}(0, \sigma_y^2)$ and $(u_k^{(0)}, u_k^{(1)})$ is the random effects for cluster k so that $(u_k^{(0)}, u_k^{(1)}) \sim \mathcal{N}(0, \Psi)$ with $\Psi = \begin{pmatrix} \psi_{00} & \psi_{01} \\ \psi_{01} & \psi_{11} \end{pmatrix}$

- for a binary variable y :

$$\text{logit}(P(y_{ki} = 1)) = \beta^{(0)} + \beta^{(1)}x_{1ki} + \beta^{(2)}x_{2ki} + u_k^{(0)} + u_k^{(1)}x_{1ki} \quad (14)$$

with the same assumption for $(u_k^{(0)}, u_k^{(1)})$.

Parameters are tuned to mimic the structure of GREAT data. This data set is an individual patient data meta-analysis provided by the GREAT Network (GREAT Network, 2013) which explores risk factors associated with short-term mortality in acute heart failure. It consists of 28 observational cohorts gathering characteristics, potential risk factors and outcomes of 11685 patients, corresponding respectively to the second and the first level of a multilevel structure. One challenge consists in explaining the left ventricular ejection fraction (LVEF), which is realised with an ultrasound, from biomarkers easier to measure, such as BNP, which is a blood biomarker, and the atrial fibrillation (AFIB). More detail on GREAT data are given in Appendix B.

Data on the variable BNP is used to generate the continuous covariates x_1 and x_3 , while the variable AFIB is used to generate the binary covariate x_2 . We used complete-case analysis to estimate parameters of the analysis model (13). Conversely, we tuned the covariate distribution according to the posterior distribution estimated by the fully Bayesian approach (Section 2.1.2) implemented in the R package *jomo* (Quartagno and Carpenter, 2016b). Thus, unless otherwise specified, default parameters are specified as follows: $K = 28, 18 \leq n_k \leq 1093$,

$$\Sigma_v = \begin{pmatrix} 0.12 & 0.001 & 0.001 \\ 0.001 & 0.12 & 0.001 \\ 0.001 & 0.001 & 0.12 \end{pmatrix}, (\alpha_1, \alpha_3) = (2.9, 2.9), \Sigma_x = \begin{pmatrix} 0.36 & 0.108 \\ 0.108 & 0.36 \end{pmatrix},$$

$$\alpha_2 = 0.42, (\beta^{(0)}, \beta^{(1)}, \beta^{(2)}) = (0.72, -0.11, 0.03), \Psi = \begin{pmatrix} 0.0077 & -0.0015 \\ -0.0015 & 0.0004 \end{pmatrix}, \sigma_Y = 0.15.$$

They define the base-case configuration. Then, these parameters will be varied one-by-one to investigate more specific configurations that can affect the methods. Details about the parameters used for each case is provided in Appendix C.1.1 in Table 9.

3.1.2 Missing data mechanisms

Variables are independently systematically missing on (x_1, x_2) with probability π_{sys} . In addition, for any clusters where each covariate is not systematically missing, sporadically missing values are generated with probability π_{spor} . Unless otherwise specified $\pi_{sys} = 0.25$ and $\pi_{spor} = 0.25$. Thus, the proportion of missing values on x_1 and x_2 is roughly 0.44. Two missing data mechanisms are considered: a MCAR and a MAR. For the MCAR mechanism, sporadically missing data is generated independently to the data, while for the MAR mechanism, sporadically missing values are added according to the observed values of the auxiliary variable x_3 . In both cases systematically missing values remain MCAR.

3.1.3 Methods

The simulation study evaluates a total of 10 methods. The reference methods are as follows:

- Full - Analysis of original dataset, before introduction of missing values
- CC - Case-wise deletion of individuals with incomplete data

Further, we consider three methods that allow imputation of sporadically and systematically missing data in multilevel data by adopting random effects distributions. The performance of these methods is of primary interest in the current simulation study:

- JM-jomo ([Quartagno and Carpenter, 2016a](#))
- FCS-GLM ([Jolani et al., 2015](#))
- FCS-2stage ([Resche-Rigon and White, 2016](#)) (estimation using REML and MM)

Finally, we consider five ad-hoc methods that were not designed to be used in multilevel data with a combination of sporadically and systematically missing values. Nevertheless, these methods are evaluated to highlight the relative merits of aforementioned, dedicated, methods:

- JM-pan ([Schafer and Yucel, 2002](#)): JM imputation by linear mixed effect models assuming homoscedasticity
- FCS-2lnorm ([van Buuren, 2010](#)): FCS imputation by linear mixed effect models assuming heteroscedasticity

- FCS-noclust (Schafer, 1997): FCS imputation by normal or logistic regression
- FCS-fixclust: FCS imputation by normal or logistic regression with fixed intercept to account for the second-level
- FCS-fixclustPMM (Little, 1988): FCS imputation by predictive mean matching with fixed intercept to account for the second-level

As explained in the introduction, these later methods are not dedicated to multilevel data with continuous and binary variables, as well as sporadically and systematically missing values. JM-pan and FCS-2lnorm only allow imputation of continuous data (Section 2.1.2). For this reason, binary variables are treated as continuous, without applying any rounding strategy (Allison, 2002). Furthermore, because of the heteroscedasticity assumption, parameters of FCS-2lnorm are not identifiable in the presence of systematically missing values. We address the issue by imputing sporadically missing values and systematically missing values separately from each other. In particular, clusters without systematically missing data are used to fit the imputation model and to impute clusters with sporadically missing values. Afterwards, parameters obtained from the first clusters without systematically missing data are used to impute the remaining clusters with systematically missing data. As regard FCS-noclust, FCS-fixclust and FCS-fixclustPMM, these methods are based on fixed intercept, implying non-identifiability of the intercept with systematically missing variables. The issue is addressed by centring dummy variables such that clusters with systematically missing values are imputed using the observed average across the remaining clusters.

3.1.4 Assessment of the inference

The primary parameters of interest are $\beta^{(1)}$, $\beta^{(2)}$, ψ_{00} and ψ_{11} in model (13) or (14). The performance of estimating these parameters will be assessed according to the bias, the root mean squared error (RMSE), the root mean square of estimated standard error (Model SE), the empirical Monte Carlo standard error (Emp SE) and the coverage of the associated confidence interval. The average time required to multiply impute one dataset is also reported.

3.1.5 Implementation

Simulations are performed with the R software (R Core Team, 2016). Multiple imputation with the JM-jomo method is performed with the R package *jomo* (Quartagno and Carpenter, 2016b). The number of iterations for the burn-in step is set to 2000, and 1000 iterations are run between each imputed dataset. Convergence has been checked from an incomplete dataset simulated from the base-case configuration by following the stationarity of the parameters of the imputation model.

Multiple imputation with FCS methods is performed using the R package *mice* (van Buuren and Groothuis-Oudshoorn, 2011) with 5 cycles. Convergence

has been checked from an incomplete dataset simulated from the base-case configuration by following the stationary of marginal quantities (means and standard deviations). For both FCS approaches, conditional imputation models consist of all available covariates, which are included in the fixed and random design matrices ($\mathbf{Z}_k = \mathbf{W}_k$).

In both cases, MI is performed using 5 imputed datasets. Each imputed dataset is analysed using the R package *nlme* (Pinheiro et al., 2016) for a continuous outcome and using the *glmer* package for a binary outcome (Bates et al., 2015). Calculation was performed on an Intel®Xeon®CPU E7530 1.87GHz. The R code used to perform simulations is available from <https://arxiv.org/src/1702.00971v1/anc/>

3.2 Results

3.2.1 Base-case configuration

Table 3 describes the simulation study results for the base-case configuration. Overall, we found that all methods yielded satisfactory estimates for fixed effects parameters of the binary variable $\beta^{(2)}$. In particular, biases were smaller than 2%, and coverage of the confidence intervals were close to their nominal level. Performance differences mainly occurred for estimation of fixed effects parameters for the continuous variable $\beta^{(1)}$ and for estimation of the variance of random effects ψ_{00} and ψ_{11} . In particular, the variance of $\widehat{\beta^{(1)}}$ is generally underestimated, leading to confidence intervals that do not reach their nominal level.

As expected, we found that ad-hoc methods suffered from several deficiencies. First of all, they underestimated the variance of $\widehat{\beta^{(1)}}$. This is likely related to the use of imputation models with homoscedastic error terms (FCS-noclust, FCS-fixclust, FCS-fixclustPMM and JM-pan), and to not properly modelling heterogeneity between clusters (FCS-noclust, FCS-fixclust and FCS-fixclustPMM). For FCS-2lnorm, underestimation of the variance is likely also caused by ignoring uncertainty on the random coefficients for systematically missing values. A second problem of the ad-hoc methods is that they yield severely biased estimates for the variance of random effects (ψ_{00} and ψ_{11}). Finally, FCS-noclust also introduces bias on point estimates for the fixed effects coefficients. In particular, by ignoring the multilevel data structure, imputed values of FCS-noclust are biased towards the overall mean and thereby affect corresponding covariate-outcome associations.

The primary methods of interest, JM-jomo, FCS-GLM and FCS-2stage provide inferences that are more satisfying as compared to the ad-hoc methods. In particular, biases are smaller and confidence intervals are closer to their nominal level. Nevertheless, some important differences were identified. First of all, the JM-jomo method tends to overestimate the variance of the estimators $\widehat{\beta^{(1)}}$ and $\widehat{\beta^{(2)}}$. Conversely, FCS-GLM tends to underestimate this variance for $\widehat{\beta^{(1)}}$, similar to ad-hoc methods assuming homoscedasticity. Another drawback of FCS-GLM is that its implementation required substantially more time to

generate an imputed dataset. Finally, the FCS-2stage method provided satisfactory inferences with both versions (REML or MM). In particular, it is the only method to provide unbiased estimates for variance components ψ_{00} and ψ_{11} .

3.2.2 Robustness to the proportion of systematically missing values

To assess the influence of the proportion of systematically missing values, we modified this proportion to 0.1, 0.25 and 0.4 (configurations 2, 1, 3 in Table 9 in Appendix C.1.1 respectively). The proportion of sporadically missing values is modified accordingly to keep the same proportion of missing values in expectation.

We found that for the three methods of primary interest, the bias on the parameters of the model remains stable regardless the proportion of systematically missing values (see Appendix C.2.1). Figure 1 reports the relative bias for the standard error estimates. For JM-jomo, the standard error estimate for $\beta^{(1)}$ tends to deviate from the empirical standard error when the proportion of systematically missing values increases, becoming upwardly biased as the relative extent of systematically missing data increases. This issue is likely related to the use of (informative) conjugate prior distributions, as their influence on the posterior is substantial when the proportion of systematically missing variables is large. Because the FCS methods use prior distributions that are derived from the data, they were less sensitive to overestimation of standard errors.

3.2.3 Robustness to the number of clusters

Influence of the number of clusters is assessed by restricting the generated datasets to their K first clusters, and varying K in $\{7, 14, 28\}$. Note that as consequence the total sample size also increases with K (2139, 4256 and 11685 respectively). Figure 2 describes the impact of the number of clusters on the bias. The impact on the variance estimate is reported in Appendix C.2.2.

For all estimands, the bias obtained from the JM-jomo method is substantial when the number of clusters is small, but decreases when as the number of clusters increases. On the contrary, the FCS methods provide more robust estimates for the variance of the random effects when the number of clusters is small. This behaviour is again likely related to the choice of the prior distributions.

3.2.4 Robustness to the cluster size

To explore the robustness of the inferences provided by the MI methods with respect to the size of the clusters, we extended the simulation study to generate clusters with equal sizes varying in $\{15, 25, 50, 100, 200, 400\}$. Relative biases are reported in Figure 3.

The biases obtained by the FCS-2stage methods are large for small clusters, but decrease when the cluster size increases. This behaviour was expected since the posterior distribution for the conditional imputation models's parameters

Table 3: Simulation study results from the base-case configuration. Point estimate, relative bias, model standard error, empirical standard error, 95% coverage and RMSE for analysis model's parameters and for several methods (Full data, Complete-case analysis, FCS-noclust, FCS-fixclust, FCS-fixclustPMM, JM-pan, FCS-2lnorm, JM-jomo , FCS-GLM, FCS-2stage with REML estimator, FCS-2stage with moment estimator). Criteria are based on 500 incomplete datasets. Average time to multiply impute one dataset is indicated in minutes. Criteria related to the continuous (resp. binary) covariate are in light (resp. dark) grey. True values are $\beta^{(1)} = -0.11$, $\beta^{(2)} = 0.03$, $\sqrt{\psi_{00}} = 0.088$, $\sqrt{\psi_{11}} = 0.02$.

	Full	C	FCS-noclust	FCS-fixclust	FCS-fixclustPMM	JM-pan	FCS-2lnorm	JM-jomo	FCS-GLM	FCS-2stageREML	FCS-2stageMM
$\beta^{(1)}$ est	-0.1101	-0.1104	-0.1102	-0.1102	-0.1039	-0.1088	-0.1098	-0.1087	-0.1100	-0.1089	-0.1090
$\beta^{(1)}$ rbias (%)	0.1	0.3	0.2	0.2	-5.5	-1.1	-0.2	-1.2	-0.0	-1.0	-0.9
$\beta^{(1)}$ model se	0.0047	0.0070	0.0043	0.0043	0.0042	0.0044	0.0056	0.0068	0.0047	0.0059	0.0059
$\beta^{(1)}$ emp se	0.0048	0.0071	0.0057	0.0058	0.0066	0.0056	0.0063	0.0057	0.0057	0.0058	0.0058
$\beta^{(1)}$ 95% cover	93.8	92.2	86.0	85.6	60.4	87.6	92.2	97.6	91.1	95.0	95.0
$\beta^{(1)}$ rmse	0.0048	0.0071	0.0057	0.0058	0.0090	0.0058	0.0063	0.0058	0.0057	0.0059	0.0059
$\beta^{(2)}$ est	0.0301	0.0299	0.0300	0.0300	0.0290	0.0295	0.0301	0.0297	0.0297	0.0295	0.0297
$\beta^{(2)}$ rbias (%)	0.2	-0.4	0.2	-0.0	-3.2	-1.7	0.2	-1.1	-1.1	-1.6	-1.0
$\beta^{(2)}$ model se	0.0029	0.0053	0.0044	0.0043	0.0043	0.0044	0.0064	0.0069	0.0046	0.0054	0.0049
$\beta^{(2)}$ emp se	0.0030	0.0053	0.0043	0.0042	0.0044	0.0042	0.0055	0.0049	0.0042	0.0044	0.0044
$\beta^{(2)}$ 95% cover	94.2	94.4	95.2	95.4	93.6	95.6	95.4	97.2	94.2	97.0	96.2
$\beta^{(2)}$ rmse	0.0030	0.0053	0.0043	0.0042	0.0045	0.0042	0.0055	0.0049	0.0043	0.0045	0.0044
$\sqrt{\psi_0}$ est	0.0859	0.0835	0.0747	0.0745	0.0712	0.0806	0.0816	0.0941	0.0783	0.0869	0.0863
$\sqrt{\psi_0}$ rbias (%)	-2.1	-4.8	-14.9	-15.1	-18.8	-8.2	-7.0	7.3	-10.8	-0.9	-1.6
$\sqrt{\psi_0}$ rmse	0.0154	0.0240	0.0189	0.0191	0.0213	0.0146	0.0166	0.0150	0.0171	0.0148	0.0150
$\sqrt{\psi_1}$ est	0.0193	0.0184	0.0138	0.0138	0.0135	0.0137	0.0174	0.0224	0.0152	0.0197	0.0194
$\sqrt{\psi_1}$ rbias (%)	-3.7	-8.1	-30.9	-31.1	-32.3	-31.4	-13.2	12.0	-24.2	-1.3	-2.9
$\sqrt{\psi_1}$ rmse	0.0041	0.0073	0.0073	0.0074	0.0075	0.0074	0.0053	0.0043	0.0066	0.0046	0.0048
time			1.8	1.2	0.9	1.4	3.3	8.0	114.7	2.5	1.0

Figure 1: Robustness to the proportion of systematically missing values: estimate of the relative bias for the SE estimate for $\widehat{\beta^{(1)}}$ (left), $\widehat{\beta^{(2)}}$ (right) according to π_{syst} for each MI method. The estimated relative bias is calculated by the difference between the model SE and the empirical SE, divided by the empirical SE. The proportion of sporadically missing values is accordingly modified to keep a constant proportion of missing values (in expectation).

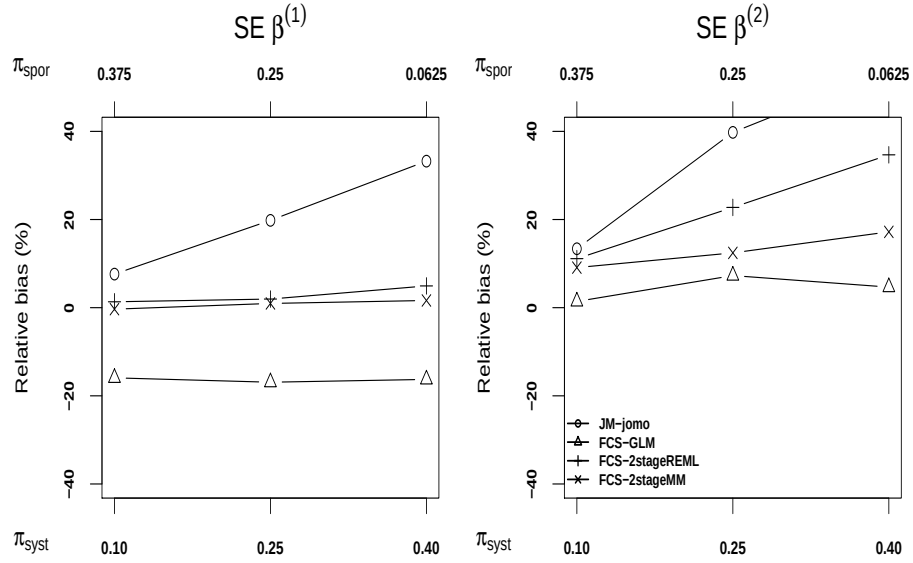


Figure 2: Robustness to the number of clusters: relative bias for the estimate of $\beta^{(1)}$ (left), $\beta^{(2)}$ (middle) and ψ_{11} (right) according to K for each MI method.

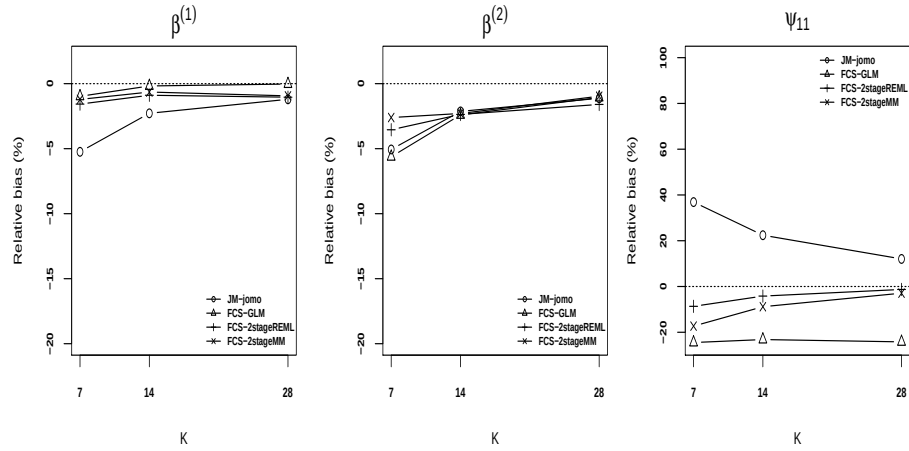
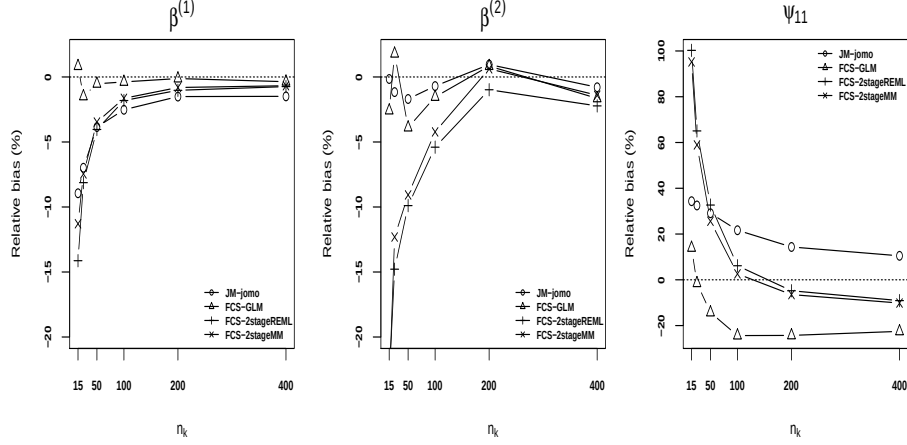


Figure 3: Robustness to the cluster size: relative bias for the estimate of $\beta^{(1)}$ (left), $\beta^{(2)}$ (middle) and ψ_{11} (right) according to n_k for each MI method. Criteria are based on 500 incomplete datasets.



are based on asymptotic properties (*cf* Section 2.2). For the JM-jomo method, the bias mainly depends on the sample size for $\beta^{(1)}$ only. On the contrary, the FCS-GLM method is fairly stable for $\beta^{(1)}$ and $\beta^{(2)}$ across the sample size. Note that the bias on ψ_{11} observed on the base-case configuration disappears with small clusters as well as the undercoverage issue because of a better estimate of the standard error (see Figure 6 in appendix C.2.3) making the method relevant in such a case.

3.2.5 Robustness to the type of imputed variables

Results in Table 4 demonstrate that also the type of imputed variable affects the performance of the imputation methods. In general, we found that JM-jomo provides smaller bias for $\beta^{(1)}$, $\beta^{(2)}$ and ψ_{11} as compared to FCS-GLM and FCS-2stage. This is likely related to the fact that FCS-GLM is not tailored for imputing sporadically missing data of binary variables, and that the two-stage estimator used in FCS-2stage is known to be biased in the presence of small clusters.

3.2.6 Robustness to the variance of random effects

Table 5 provides inference results when the covariance matrix of random effects is multiplied by a factor 2. The biases reported for the variance of random effects are less than 2% for JM-jomo, while they reached 11% in the base-case configuration. Such behaviour can be explained by the smaller influence of the prior distribution for random effects when the effect of random effects is stronger. On the contrary, the bias increases for the FCS-2stage methods.

Table 4: Binary covariates. Point estimate, relative bias, model standard error, empirical standard error, 95% coverage and RMSE for analysis model's parameters and for several methods (Full data, JM-jomo, FCS-GLM, FCS-2stage with REML estimator, FCS-2stage with moment estimator). Criteria are based on 500 incomplete datasets. Average time to multiply impute one dataset is indicated in minutes. True values are $\beta^{(1)} = -0.11$, $\beta^{(2)} = 0.03$, $\sqrt{\psi_{00}} = 0.088$, $\sqrt{\psi_{11}} = 0.02$.

	Full	JM-jomo	FCS-GLM	FCS-2stageREML	FCS-2stageMM
$\beta^{(1)}$ est	-0.1098	-0.1091	-0.1081	-0.1080	-0.1083
$\beta^{(1)}$ rbias (%)	-0.2	-0.8	-1.7	-1.8	-1.5
$\beta^{(1)}$ model se	0.0050	0.0074	0.0057	0.0063	0.0056
$\beta^{(1)}$ emp se	0.0049	0.0064	0.0059	0.0060	0.0061
$\beta^{(1)}$ 95% cover	95.0	97.0	92.0	94.0	90.4
$\beta^{(1)}$ rmse	0.0049	0.0064	0.0062	0.0063	0.0063
$\beta^{(2)}$ est	0.0303	0.0298	0.0295	0.0294	0.0295
$\beta^{(2)}$ rbias (%)	1.0	-0.6	-1.8	-2.1	-1.6
$\beta^{(2)}$ model se	0.0029	0.0072	0.0044	0.0051	0.0045
$\beta^{(2)}$ emp se	0.0028	0.0047	0.0043	0.0044	0.0043
$\beta^{(2)}$ 95% cover	95.0	98.6	95.2	96.2	95.0
$\beta^{(2)}$ rmse	0.0028	0.0047	0.0044	0.0044	0.0044
$\sqrt{\psi_{00}}$ est	0.0876	0.0861	0.0807	0.0834	0.0827
$\sqrt{\psi_{00}}$ rbias (%)	-0.2	-1.9	-8.0	-4.9	-5.7
$\sqrt{\psi_{00}}$ rmse	0.0123	0.0122	0.0137	0.0129	0.0131
$\sqrt{\psi_{11}}$ est	0.0198	0.0232	0.0151	0.0185	0.0164
$\sqrt{\psi_{11}}$ rbias (%)	-0.8	16.2	-24.3	-7.7	-17.8
$\sqrt{\psi_{11}}$ rmse	0.0046	0.0053	0.0069	0.0057	0.0068
time		5.6	95.1	1.4	0.6

Table 5: Higher variance for Ψ . Point estimate, relative bias, model standard error, empirical standard error, 95% coverage and RMSE for analysis model's parameters and for several methods (Full data, JM-jomo , FCS-GLM, FCS-2stage with REML estimator, FCS-2stage with moment estimator). Criteria are based on 500 incomplete datasets. Average time to multiply impute one dataset is indicated in minutes. Criteria related to the continuous (resp. binary) covariate are in light (resp. dark) grey. True values are $\beta^{(1)} = -0.11$, $\beta^{(2)} = 0.03$, $\sqrt{\psi_{00}} = 0.124$, $\sqrt{\psi_{11}} = 0.028$.

	Full	JM-jomo	FCS-GLM	FCS- 2stageREML	FCS- 2stageMM
$\beta^{(1)}$ est	-0.1102	-0.1087	-0.1099	-0.1089	-0.1090
$\beta^{(1)}$ rbias (%)	0.2	-1.1	-0.1	-1.0	-0.9
$\beta^{(1)}$ model se	0.0061	0.0075	0.0059	0.0070	0.0070
$\beta^{(1)}$ emp se	0.0063	0.0071	0.0072	0.0073	0.0073
$\beta^{(1)}$ 95% cover	93.8	95.0	90.1	93.0	93.8
$\beta^{(1)}$ rmse	0.0063	0.0073	0.0072	0.0073	0.0073
$\beta^{(2)}$ est	0.0301	0.0297	0.0295	0.0294	0.0296
$\beta^{(2)}$ rbias (%)	0.2	-0.8	-1.6	-2.0	-1.3
$\beta^{(2)}$ model se	0.0029	0.0070	0.0046	0.0055	0.0050
$\beta^{(2)}$ emp se	0.0030	0.0048	0.0042	0.0044	0.0044
$\beta^{(2)}$ 95% cover	94.2	98.2	94.2	97.4	96.0
$\beta^{(2)}$ rmse	0.0030	0.0048	0.0043	0.0044	0.0044
$\sqrt{\psi_0}$ est	0.1220	0.1225	0.1089	0.1172	0.1166
$\sqrt{\psi_0}$ rbias (%)	-1.7	-1.3	-12.2	-5.6	-6.0
$\sqrt{\psi_0}$ rmse	0.0198	0.0175	0.0240	0.0203	0.0205
$\sqrt{\psi_1}$ est	0.0275	0.0279	0.0220	0.0262	0.0260
$\sqrt{\psi_1}$ rbias (%)	-2.8	-1.5	-22.1	-7.3	-8.2
$\sqrt{\psi_1}$ rmse	0.0048	0.0043	0.0083	0.0057	0.0059
time		7.7	102.5	2.2	0.9

3.2.7 Other configurations

Other configurations that have been investigated are presented in Appendix C.1.1. These configurations consider the nature of the outcome (configuration 5), the missing data mechanism for sporadically missing values (configurations 6, 7, 8), the complexity of the analysis model (configuration 9), the number of individuals with unequal cluster sizes (configuration 11, 12), the intra-class correlation (configurations 13, 14), the correlation between random intercepts generating variables x_1, x_2, x_3 (configuration 15), the correlation between continuous variables in each cluster (configuration 16), the covariance matrix of the random effects (configuration 17, 18), the use of a probit link for generating binary covariates (configuration 19). Figure 7 in Appendix C.2.4 reports the distribution of the relative bias over the several configurations, while tables gathering inference results are available in the supplementary materials. We found similar results as compared to the base-case configuration.

4 Application to GREAT data

The MI methods are applied to the GREAT data (Appendix B). We considered model (13) for analysis model, with x_1 representing the variable BNP, x_2 the variable AFIB and y the variable LVEF. Although only three variables are included in the analysis model, the imputation models are based on nine variables to render the MAR assumption more credible (Enders, 2010; Schafer, 1997).

Table 6: GREAT data: Point estimate, model standard error for parameters of a linear mixed effects model for several methods (Complete-case analysis, JM-jomo, FCS-GLM, FCS-2stage with REML estimator, FCS-2stage with moment estimator). 20 imputed arrays are considered for MI methods. 10 iterations are used for FCS methods. Time to multiply impute the dataset is indicated in minutes. Criteria related to the continuous (resp. binary) covariate are in light (resp. dark) grey.

		CC	JM-jomo	FCS-GLM	FCS-2stageREML	FCS-2stageMM
β_{BNP}	est	-0.1132	-0.0891	-0.1002	-0.0854	-0.1009
	model se	0.0108	0.0078	0.0163	0.0099	0.0112
β_{AFIB}	est	0.0268	0.0216	0.0218	0.0215	0.0273
	model se	0.0071	0.0046	0.0066	0.0040	0.0045
ψ_{00}	est	0.1112	0.1075	0.1232	0.1220	0.1189
ψ_{BNP}	est	0.0290	0.0306	0.0348	0.0351	0.0332
time (min)			94.0	30819.5	361.3	31.8

Results in Table 6 indicate that the missing data mechanism in the GREAT application is not likely to be missing completely at random. In particular, complete-case analysis yielded estimates for the fixed coefficient β_{BNP} farther away from null as compared to the MI methods.

As expected, standard errors obtained by CC were larger than those obtained from the base-case configuration of the simulation study. In general, standard errors for fixed effects estimates were smaller for the MI methods as compared to CC. An exception occurred for the variable β_{BNP} when using FCS-2stageMM or FCS-GLM. Possibly, this is related to convergence issues resulting from limited number of iterations and relatively small number of imputed data sets. For instance, when we allowed for 50 iterations (rather than 10) and generated 50 imputed data sets (instead of 20), FCS-2stageMM yielded a standard error of 0.0091 for β_{BNP} .

Remarkably, we found that JM-jomo yielded the smallest standard errors for fixed effect estimates. This situation did not arise in the simulation studies, and is likely related to over-parametrisation of the FCS methods. In particular, the conditional imputation models assume random effects for all covariates, which

leads to a substantial increase of the number of parameters as the number of covariates increases, hence inflating the variance around imputed values.

Finally, in agreement with the simulation studies, we found that the FCS-GLM method required substantially more computation time as compared to the other MI methods. This limitation is somewhat problematic, as the number of incomplete variables was rather limited in the GREAT application. As a result, checking convergence of the distribution of missing values to their posterior distribution becomes very difficult.

5 Discussion

As international collaboration becomes more common and access to large shared datasets increases, researchers increasingly often face incomplete multilevel data. Thus, handling systematically missing values becomes inevitable (Debray et al., 2015b,a). In this work, we compared three recent multiple imputation methods for addressing this issue: JM-jomo, FCS-GLM and FCS-2stage. We also considered several extensions to better handle continuous and binary data in the presence of sporadically and systematically missing values. We highlighted the relevance of using these methods, and demonstrated their superiority over ad-hoc strategies through extensive simulation studies. Although the differences between the three imputation models are mainly technical, their properties may substantially differ according to the considered dataset.

In general, we found that JM-jomo tends to be conservative. This behaviour is in line with simulation study presented in Quartagno and Carpenter (2016a), and is related to the use of inverse-Wishart prior distributions for modelling the covariance matrices. Although this distribution avoids convergence issues of the Gibbs sampler (Schafer and Yucel, 2002), its use is not necessarily supported by the data. Furthermore, the prior distributions for the parameters of the inverse Wishart distribution appear to be rather influential. In particular, by sampling the covariance matrices using few degrees of freedom, too much variability is introduced for the within-cluster covariance matrices. As a result, fixed effects estimates vary too much across imputed datasets, leading to over-estimation of the variance components. Note that the influence of the prior distributions of imputation model’s parameters have also been recently exhibited in the context of continuous data Kunkel and Kaizar (2017).

Bias is observed for JM-jomo when the number of individuals and/or clusters is small and/or variance of random effects is small. In such situations, the Inverse-Wishart prior distributions become very informative (Gelman, 2006). Furthermore, because the Inverse-Wishart distribution tends to generate too much variability, its use may lead to shrinkage of regression coefficients when imputing continuous covariates ($\beta^{(1)}$ in the simulation study). This issue is less problematic for binary covariates ($\beta^{(2)}$) because the diagonal terms of the within covariance matrices are constraint to be equal to one. However, inference for GLMM models with few clusters is challenging, even without missing data

(McNeish and Stapleton, 2016). When the number of clusters is large, substantial improvements can be obtained by increasing the degrees of freedom of the chi-squared distribution.

For FCS-GLM, we found that imputations were quite accurate for continuous variables. However, FCS-GLM is limited by the homoscedasticity assumption (van Buuren, 2010; Resche-Rigon and White, 2016). In particular, by fitting a homoscedastic model on heteroscedastic data, standard errors tend to be underestimated, even in the absence of missing data. As a result, the FCS-GLM method cannot fully propagate the sampling variability, leading to an underestimation of the variance of the parameters of the analysis model. This results in confidence intervals that are too narrow. However, the homoscedastic assumption can become a main advantage with small clusters since it avoids overfitting issues. As a result, the standard errors become well estimated also for continuous covariates. For this reason, FCS-GLM is a relevant method to use with small clusters. Another current problem of FCS-GLM is the time required for generating imputed datasets, particularly in large datasets. These results are in line with the simulation study of Resche-Rigon and White (2016) comparing FCS-GLM and FCS-2stage.

FCS-2stage does not present any recurrent trend. Simulation study results suggest that using the log-normal distribution as approximation for the posterior distribution of the error variance outperformed modelling through Inverse-Wishart distributions (as in JM-jomo and FCS-GLM). Although the MM estimator of FCS-2stage is known to underestimate relevant variance components, similar inferences were obtained when adopting REML (Langan et al., 2016). However, FCS-2stage may be problematic when imputing binary covariates in small clusters, as the maximum likelihood estimator used in stage one is known to yield biased estimates in such circumstances. For this reason, we investigated the use of Firth’s correction (Firth, 1993), but its implementation did not yield substantial improvements. Further research is warranted to investigate how FCS-2stage may be improved when applied to datasets with small clusters. In the GREAT application, we found that JM-jomo and FCS-2stage produce similar point estimates, but that the former yields smaller standard errors. This may reflect the ability of JM-jomo to borrow information about the study-specific covariance matrix across studies; and/or the appropriateness of using the Inverse Wishart model in the GREAT data.

A key issue in all MI methods is the use of *congenial* imputation models (Meng, 1994). Congeniality means that there is a joint model which implies the imputation model and the analysis model as submodels. Some results have been obtained for continuous variables when the analysis model does not include a random slope (Quartagno and Carpenter, 2016a; Resche-Rigon and White, 2016). However, as raised in Grund et al. (2016), with a random slope, these imputation models are uncongenial. Indeed, considering model (1), the outcome depends on a product of two random variables: the random effect (b_k)

and the associated covariate ($\mathbf{W}_k$). Consequently, the marginal distribution of the outcome becomes highly complex, whereas a joint imputation model like the one used in JM-jomo (without covariate in the right part of model (2)), assumes simpler Gaussian marginal distributions in each cluster. In the same way, for FCS methods, the conditional distribution of one covariate is no longer analytically tractable. This implies that the distribution of the covariates given the outcome cannot be written according to a GLMM model. Thus, imputation models are misspecified whatever the imputation method used. Nevertheless, our simulation study highlights that this is a minor practical issue since the biases remain very small for fixed coefficients and variance of random effects (see also Appendix C.2.5). Recent progress to provide imputation models ensuring congeniality even in complex settings (Bartlett et al., 2015) seems promising for the multilevel setting.

Another source of mis-specification is the choice of random and fixed effects in the imputation models. In particular, tuning each conditional imputation model is tricky in practice for FCS approaches, particularly with a lot of variables, although models can become over-parametrised otherwise. More generally, finding conditional imputation models with few parameters in FCS approaches is a current topic of research (Zhao and Long, 2016) in the one-level case, and appears even more challenging in the two-level approach. In this paper we used the default model of the methods: for JM-jomo, all variables are in the response part of model (2), which corresponds to normal marginal distributions with random intercept, whereas FCS approaches include all covariates in fixed and random effects, making such marginal distribution more complex. Consequently, the imputation models are not strictly the same according to this aspect, that could also explain some differences between JM and FCS approaches.

An additional difficulty for all MI methods is the imputation of binary variables. Whereas JM-jomo overcomes it quite well by considering a fully Bayesian multilevel modelling approach with a probit link function, FCS-GLM and FCS-2stage are less tailored for such variables: FCS-GLM because it considers a logit link making difficult to handle sporadically missing values, and FCS-2stage because it infers separately on each clusters, making samples too small to provide accurate inferences. For these reasons, both FCS methods could be improved by adopting: a probit link function for FCS-GLM, and applying bias correction for variances and point estimates for FCS-2stage. Note that in contrast to FCS-2stage, the methods FCS-GLM and JM-jomo can also handle nominal and count variables. FCS-2stage could further be extended by considering additional link functions in the regression models used at stage 1. Finally, although FCS-2lnorm provides encouraging performance for imputing missing continuous and binary multilevel data, it does not *properly* reflect (Schafer, 1997, p. 105) the variability of random coefficients. For this reason, its usefulness remains limited in the presence of systematically missing values.

Although we only considered a 2-level setting in this study, extensions of the presented MI methods to a higher hierarchical structure are relatively straightforward. At this moment, only JM-jomo proposes an implemented solution to

address such situations. Note that missing values may also occur at level-2. JM-jomo naturally handles this setting, but FCS approaches have also been developed to achieve this goal ([van Buuren and Groothuis-Oudshoorn, 2011](#)). In addition, we did not focus the topic on longitudinal data, or more generally on data with very few observations per cluster, like in education research field. However, systematically missing values are also frequent in such cases and raise additional overfitting issues of the imputation models.

Furthermore, our study focuses on the use of GLMM models to analyse multilevel data, which facilitates the use of MI. Indeed, direct maximum likelihood inference is another possible strategy to address missing data, but its implementation becomes difficult when dealing with multilevel variables ([Longford, 2008](#); [Schafer, 1997](#)). However, many other statistics than the parameters of a GLMM model can be of interest. For instance, [Curran and Hussong \(2009\)](#); [Curran et al. \(2008\)](#) proposed using item response theory to fit measurement models.

From a general point of view, whatever the imputation method used, accurate inferences for a GLMM model can be expected only with a high (or moderate) number of clusters. Heteroscedastic MI methods perform better than homoscedastic methods, which should be reserved with few individuals only. Methods based on conjugate prior distributions should be used with caution when the proportion of missing values is very high. Multiple imputation of binary variable is challenging, and methods can have drawbacks in such a case.

In this regards, JM-jomo could be recommended when the number of incomplete binary variables is large and when the number of observed clusters is large. FCS-2stage performs quite well, but should be avoided when clusters are small, or equivalently, when the proportion of sporadically missing values is large. This method is particularly relevant compared to the others when the number clusters with systematically missing variables is large. The MM version offers a quick solution to have an initial overview of the inference results. Finally, FCS-GLM appears advantageous when clusters are small.

From our point of view, the topic of inference for multilevel incomplete data needs to strengthen the theoretical underpinnings to improve the fit of imputation models, as well as some developments to broaden the scope of the evaluated methods. Among them, the congeniality has been recently discussed but need more efforts for analysis model with random slope. Machine-learning methods offering more flexibility could be considered for this purpose. In addition, solutions to handle missing data without assuming the ignorability of the missing data mechanism need to be investigated. Furthermore, in the big data era, MI methods handling a large number of variables need also to be studied. Finally, proposing imputation models for nominal or ordered variables avoiding informative prior distributions constitutes a main line of research.

Acknowledgements

We thank the Global Research on Acute Conditions Team ([GREAT Network, 2013](#)) for allowing the use of their data.

Funding

TPAD was supported by the Netherlands Organization for Scientific Research (91617050 and 91215058). IRW was funded by the Medical Research Council [Unit Programme number U105260558 and MC_UU_12023/21]. VA and MRR was supported by the Agence nationale de sécurité du médicament et des produits de santé (ANSM).

References

- Albert, A. and Anderson, J. (1984). On the existence of maximum likelihood estimates in logistic regression models. *Biometrika*, 71(1):1–10.
- Allison, P. (2002). *Missing Data*. Thousand Oaks, CA: Sage.
- Andridge, R. R. (2011). Quantifying the impact of fixed effects modeling of clusters in multiple imputation for cluster randomized trials. *Biometrical Journal*, 53(1):57–74.
- Asparouhov, T. and Muthén, B. (2010). Multiple imputation with Mplus. Technical report. Retrieved from <http://www.statmodel.com/download/Imputations7.pdf>.
- Bartlett, J., Seaman, S., White, I., Carpenter, J., and for the Alzheimer’s Disease Neuroimaging Initiative (2015). Multiple imputation of covariates by fully conditional specification: Accommodating the substantive model. *Statistical Methods in Medical Research*, 24(4):462–487.
- Bates, D., Mächler, M., Bolker, B., and Walker, S. (2015). Fitting linear mixed-effects models using *lme4*. *Journal of Statistical Software*, 67(1):1–48.
- Blossfeld, H.-P., Günther Roßbach, H., and von Maurice, J., editors (2011). *Education as a lifelong process : the German National Educational Panel Study (NEPS)*. VS Verlag für Sozialwissenschaften, Wiesbaden, Germany.
- Bos, W., Lankes, E.-M., Prenzel, M., Schwippert, K., and Valtin, R., editors (2003). *Erste Ergebnisse aus IGLU : Schulleistungen am Ende der vierten Jahrgangsstufe im internationalen Vergleich [The first]*. Waxmann, Münster, Germany.
- Browne, W. J. (2006). MCMC algorithms for constrained variance matrices. *Computational Statistics & Data Analysis*, 50(7):1655–1677.

- Burke, D. L., Ensor, J., and Riley, R. (2016). Meta-analysis using individual participant data: one-stage and two-stage approaches, and why they may differ. *Statistics in Medicine*. To appear.
- Carpenter, J. and Kenward, M. (2013). *Multiple imputation and its application*. Wiley, 1st edition edition.
- Carrig, M. M., Manrique-Vallier, D., Ranby, K. W., Reiter, J., and Hoyle, R. H. (2015). A nonparametric, multiple imputation-based method for the retrospective integration of data sets. *Multivariate Behavioral Research*, 50(4):383–397.
- Curran, P., Hussong, A. M., Cai, L., Huang, W., Chassin, L., Sher, K. J., and Zucker, R. A. (2008). Pooling data from multiple longitudinal studies: The role of item response theory in integrative data analysis. *Developmental psychology*, 44(2):365.
- Curran, P. J. and Hussong, A. M. (2009). Integrative data analysis: the simultaneous analysis of multiple data sets. *Psychological methods*, 14(2):81.
- Debray, T., Moons, K., van Valkenhoef, G., Efthimiou, O., Hummel, N., Groenwold, R., and Reitsma, J. o. (2015a). Get real in individual participant data (IPD) meta-analysis: a review of the methodology. *Research Synthesis Methods*, 6(4):293–309.
- Debray, T., Riley, R., Rovers, M., Reitsma, J., Moons, K., and on behalf of the Cochrane IPD Meta-analysis Methods group (2015b). Individual Participant Data (IPD) Meta- analyses of Diagnostic and Prognostic Modeling Studies: Guidance on Their Use. *PLoS Med.*, 12(10):e1001886.
- DerSimonian, R. and Laird, N. (1986). Meta-analysis in clinical trials. *Controlled clinical trials*, 7(3):177–188.
- Drechsler, J. (2015). Multiple imputation of multilevel missing data - rigor versus simplicity. *Journal of Educational and Behavioral Statistics*, 40(1):69–95.
- Enders, C. (2010). *Applied Missing Data Analysis*. Methodology in the social sciences. Guilford Press.
- Enders, C., Keller, B. T., and Levy, R. (2017). A fully conditional specification approach to multilevel imputation of categorical and continuous variables. *Psychological Methods*.
- Enders, C., Mistler, S., and Keller, B. (2016). Multilevel multiple imputation: A review and evaluation of joint modeling and chained equations imputation. *Psychological Methods*, 21(2):222–240.

- Erler, N., Rizopoulos, D., van Rosmalen, J., Jaddoe, V., Franco, O., and Lesaffre, E. (2016). Dealing with missing covariates in epidemiologic studies: A comparison between multiple imputation and a full Bayesian approach. *Stat Med*.
- Firth, D. (1993). Bias reduction of maximum likelihood estimates. *Biometrika*, 80(1):27–38.
- Gelman, A. (2006). Prior distributions for variance parameters in hierarchical models (comment on article by browne and draper). *Bayesian Analysis*, 1(3):515–534.
- Geman, S. and Geman, D. (1984). Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 6(6):721–741.
- Goldstein, H., Bonnet, G., and Rocher, T. (2007). Multilevel structural equation models for the analysis of comparative data on educational performance. *Journal of Educational and Behavioral Statistics*, 32(3):252–286.
- Goldstein, H., Carpenter, J., Kenward, M. G., and Levin, K. (2009). Multilevel models with multivariate mixed response types. *Statistical Modelling*, 9(3):173–197.
- Graham, J. W. (2012). *Missing Data: Analysis and Design*. Statistics for Social and Behavioral Sciences. Springer-Verlag New York, 1 edition.
- GREAT Network (2013). Managing acute heart failure in the ed - case studies from the acute heart failure academy. <http://www.greatnetwork.org>.
- Grund, S., Lüdtke, O., and Robitzsch, A. (2016). Multiple imputation of missing covariate values in multilevel models with random slopes: a cautionary note. *Behavior Research Methods*, 48(2):640–649.
- Hughes, R. A., White, I. R., Seaman, S., Carpenter, J., Tilling, K., and Sterne, J. (2014). Joint modelling rationale for chained equations. *BMC Medical Research Methodology*, 14(1):28.
- Jackson, D., White, I., and Riley, R. (2013). A matrix-based method of moments for fitting the multivariate random effects model for meta-analysis and meta-regression. *Biometrical Journal*, 55(2):231–245.
- Jolani, S. (2017). Hierarchical imputation of systematically and sporadically missing data: An approximate bayesian approach using chained equations. *Biometrical Journal*.
- Jolani, S., Debray, T., Koffijberg, H., van Buuren, S., and Moons, K. (2015). Imputation of systematically missing predictors in an individual participant data meta-analysis: a generalized approach using MICE. *Statistics in Medicine*, 34(11):1841–1863.

- Kropko, J., Goodrich, B., Gelman, A., and Hill, J. (2014). Multiple imputation for continuous and categorical data: Comparing joint multivariate normal and conditional approaches. *Political Analysis*, 22(4):497–519.
- Kunkel, D. and Kaizar, E. E. (2017). A comparison of existing methods for multiple imputation in individual participant data meta-analysis. *Statistics in Medicine*, 36(22):3507–3532.
- Langan, D., Higgins, J., and Simmonds, M. (2016). Comparative performance of heterogeneity variance estimators in meta-analysis: a review of simulation studies. *Research Synthesis Methods*. To appear.
- Lassus, J., Gayat, E., Mueller, C., Peacock, W., Spinar, J., Harjola, V., van Kimmenade, R., Pathak, A., Mueller, T., and *et al.* (2013). Incremental value of biomarkers to clinical variables for mortality prediction in acutely decompensated heart failure: the Multinational Observational Cohort on Acute Heart Failure (MOCA) study. *International Journal of Cardiology*, 168(3):2186–2194.
- Lee, K. and Carlin, J. (2010). Multiple imputation for missing data: Fully conditional specification versus multivariate normal imputation. *American Journal of Epidemiology*, 171(5):624–632.
- Lee, Y., Nelder, J., and Pawitan, Y. (2006). *Generalized Linear Models with Random Effects: Unified Analysis via H-likelihood*. Chapman & Hall/CRC Monographs on Statistics & Applied Probability. CRC Press, Boca Raton.
- Little, R. (1988). Missing-data adjustments in large surveys. *Journal of Business & Economic Statistics*, 6(3):287–296.
- Little, R. and Rubin, D. (2002). *Statistical Analysis with Missing Data*. Wiley series in probability and statistics, New-York.
- Liu, J., Gelman, A., Hill, J., Su, Y., and Kropko, J. (2014). On the stationary distribution of iterative imputations. *Biometrika*, 101(1):155–173.
- Longford, N. (2008). *Missing Data*, pages 377–399. Springer New York, New York, NY.
- Mathew, T. and Nordstrom, K. (2010). Comparison of one-step and two-step meta-analysis models using individual patient data. *Biometrical Journal*, 52(2):271–287.
- McNeish, D. and Stapleton, L. (2016). Modeling clustered data with very few clusters. *Multivariate Behavioral Research*, 51(4):495–518.
- Mebazaa, A., Gayat, E., Lassus, J., Meas, T., Mueller, C., and *et al.* (2013). Association between elevated blood glucose and outcome in acute heart failure: results from an international observational cohort. *Journal of the American College of Cardiology*, 61(8):820–829.

- Meng, X. (1994). Multiple-imputation inferences with uncongenial sources of input (with discussion). *Statistical Science*, 10:538–573.
- Mullis, I., Martin, M., Gonzalez, E., and Kennedy, A. (2003). Pirls 2001 international report: Iea’s study of reading literacy achievement in primary school in 35 countries.
- Natarajan, R. and Kass, R. (2000). Reference Bayesian methods for generalized linear mixed models. *Journal of the American Statistical Association*, 95(449):227–237.
- Noh, M. and Lee, Y. (2007). REML estimation for binary data in GLMMs. *Journal of Multivariate Analysis*, 98(5):896–915.
- Pinheiro, J. and Bates, D. (2000). *Mixed-effects models in S and S-PLUS*. Springer, New York.
- Pinheiro, J., Bates, D., DebRoy, S., Sarkar, D., and R Core Team (2016). *nlme: Linear and Nonlinear Mixed Effects Models*. R package version 3.1-128.
- Quartagno, M. and Carpenter, J. (2016a). Multiple imputation for IPD meta-analysis: allowing for heterogeneity and studies with missing covariates. *Statistics in Medicine*, 35(17):2938–2954.
- Quartagno, M. and Carpenter, J. (2016b). *jomo: A package for Multilevel Joint Modelling Multiple Imputation*. R package version 2.2-0.
- R Core Team (2016). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. version 3.3.0.
- Raghunathan, T., Lepkowski, J. M., Van Hoewyk, J., and Solenberger, P. (2001). A multivariate technique for multiply imputing missing values using a sequence of regression models. *Survey methodology*, 27(1):85–96.
- Resche-Rigon, M. and White, I. (2016). Multiple imputation by chained equations for systematically and sporadically missing multilevel data. *Statistical Methods in Medical Research*. <http://dx.doi.org/10.1177/0962280216666564>.
- Resche-Rigon, M., White, I., Bartlett, J., Peters, S., Thompson, S., and on behalf of the PROG-IMT Study Group (2013). Multiple imputation for handling systematically missing confounders in meta-analysis of individual participant data. *Statistics in Medicine*, 32(28):4890–4905.
- Riley, R., Lambert, C., Staessen, J., Wang, J., Gueyffier, F., Thijs, L., and Bouitrie, F. (2008). Meta-analysis of continuous outcomes combining individual patient data and aggregate data. *Statistics in Medicine*, 27(11):1870–1893.

- Riley, R. D., Ensor, J., Snell, K. I. E., Debray, T. P. A., Altman, D. G., Moons, K. G. M., and Collins, G. S. (2016). External validation of clinical prediction models using big datasets from e-health records or ipd meta-analysis: opportunities and challenges. *BMJ*, 353.
- Robert, C. (2007). *The Bayesian Choice: From Decision-Theoretic Foundations to Computational Implementation*. Springer, 2nd edition.
- Rubin, D. (1976). Inference and missing data. *Biometrika*, 63:581–592.
- Rubin, D. (1987). *Multiple Imputation for Non-Response in Survey*. Wiley, New-York.
- Schafer, J. (1997). *Analysis of Incomplete Multivariate Data*. Chapman & Hall/CRC, London.
- Schafer, J. and Yucel, R. (2002). Computational strategies for multivariate linear mixed-effects models with missing values. *Journal of Computational and Graphical Statistics*, 11:421–442.
- Simmonds, M., Higgins, J., Stewart, L., Tierney, J., Clarke, M., and Thompson, S. (2005). Meta-analysis of individual patient data from randomized trials: a review of methods used in practice. *Clinical Trials (London, England)*, 2(3):209–217.
- Tanner, M. and Wong, W. (1987). The calculation of posterior distributions by data augmentation. *Journal of the American Statistical Association*, 82:805–811.
- van Buuren, S. (2007). Multiple imputation of discrete and continuous data by fully conditional specification. *Statistical Methods in Medical Research*, 16(3):219–242.
- van Buuren, S. (2010). Multiple imputation of multilevel data. In *The Handbook of Advanced Multilevel Analysis*, chapter forthcoming. Routledge, Milton Park, UK.
- van Buuren, S., Brand, J., Groothuis-Oudshoorn, C., and Rubin, D. (2006). Fully conditional specification in multivariate imputation. *Journal of Statistical Computation and Simulation*, 76:1049–1064.
- van Buuren, S. and Groothuis-Oudshoorn, K. (2011). *mice*: Multivariate imputation by chained equations in R. *Journal of Statistical Software*, 45(3):1–67.
- Vink, G., Lazendic, G., and van Buuren, S. (2015). Partitioned predictive mean matching as a multilevel imputation technique. *Psychological Test and Assessment Modeling*, 57(4):577–594.
- Wagstaff, D. and Harel, O. (2011). A closer examination of three small-sample approximations to the multiple-imputation degrees of freedom. *Stata Journal*, 11(3):403–419.

- Yucel, R. (2011). Random-covariances and mixed-effects models for imputing multivariate multilevel continuous data. *Statistical modelling*, 11(4):351–370.
- Zhao, E. and Yucel, R. (2009). Performance of sequential imputation method in multilevel applications. In *Proceedings of the Survey Research Methods Section*. ASA 2009, Alexandria, VA.
- Zhao, Y. and Long, Q. (2016). Multiple imputation in the presence of high-dimensional data. *Statistical Methods in Medical Research*, 25(5):2021–2035.
- Zhu, J. and Raghunathan, T. (2015). Convergence properties of a sequential regression multiple imputation algorithm. *Journal of the American Statistical Association*, 110(511):1112–1124.

A Posterior distributions

This section reports the prior and posterior distributions of the parameters of MI method as well as technical points to obtain realisations from them.

A.1 JM-jomo

The prior distributions for parameters $\theta = (\beta, \Psi, (\Sigma_k)_{1 \leq k \leq K})$ of model (2) are as follows:

$$\begin{aligned}\beta &\propto 1 \\ \Psi^{-1} &\sim \mathcal{W}(\nu_1, \Lambda_1) \\ \Sigma_k^{-1} | \nu_2, \Lambda_2 &\sim \mathcal{W}(\nu_2, \Lambda_2) \text{ for all } 1 \leq k \leq K \\ \nu_2 &\sim \chi^2(\eta), \quad \Lambda_2^{-1} \sim \mathcal{W}(\nu_3, \Lambda_3)\end{aligned}$$

where $\mathcal{W}(\nu, \Lambda)$ denotes the Wishart distribution with ν degrees of freedom and scale matrix Λ , and η denotes the degrees of freedom of the chi-squared distribution. Hyperparameters are set as $\nu_1 = r$, $\Lambda_1 = \mathbb{I}_r$, $\nu_3 = rK$, $\Lambda_3 = \mathbb{I}_{rK}$ and $\eta = rK$.

If data were complete, conditional posterior distributions for each component of θ would be:

$$\beta | Y, (\Sigma_k)_{1 \leq k \leq K}, b \sim \mathcal{N}(\hat{\beta}, \widehat{\text{var}}(\hat{\beta})) \quad (15)$$

$$\Psi^{-1} | Y, b \sim \mathcal{W}(\nu_1 + K, (\Lambda_1^{-1} + S_1)^{-1}) \quad (16)$$

$$\Sigma_k^{-1} | Y, b, \beta, \nu_2, \Lambda_2 \sim \mathcal{W}(\nu_2 + n, (\Lambda_2^{-1} + S_2)^{-1}) \quad (17)$$

with $\hat{\beta}$ the estimate of the weighted least-squares estimator, $\widehat{\text{var}}(\hat{\beta})$ the estimate of the associated variance, $S_1 = (b_1^V, \dots, b_K^V) (b_1^V, \dots, b_K^V)^\top$, $S_2 = \sum_k \hat{\varepsilon}_k^V (\hat{\varepsilon}_k^V)^\top$ where $\hat{\varepsilon}_k$ denotes the residuals for cluster k .

To draw the parameters of the imputation model from their posterior distribution with a multivariate missing data pattern, a DA algorithm is used (Tanner and Wong, 1987). Note that the posterior distribution of θ also depends on the random coefficients and on the random parameters (ν_2, Λ_2) , which is why they are updated at each cycle of the algorithm.

To deal with binary variables, a probit link and a latent variables framework have been proposed (Goldstein et al., 2009). Note that drawing $(\Sigma_k)_{1 \leq k \leq K}$ from its posterior distribution is more tricky with binary variables since Σ_k has to respect the constraint of diagonal elements equal to one for latent variables (Browne, 2006).

A.2 FCS-GLM

For continuous incomplete variables, the conditional imputation model is model (1) assuming homoscedastic error terms, *i.e.* $\sigma_k = \sigma$ for all k . From non-informative independent priors, components of the parameter $\theta = (\beta, \Psi, \sigma)$ are drawn from the following posterior distributions:

$$\sigma_j^2 | Y, b \sim \text{Inv-}\Gamma \left(\frac{n-p}{2}, \frac{(n-p)\widehat{\sigma}^2}{2} \right) \quad (18)$$

$$\beta | Y, b, \sigma_j^2 \sim \mathcal{N} \left(\widehat{\beta}, \widehat{\text{var}} \left(\widehat{\beta} \right) \right) \quad (19)$$

$$\Psi^{-1} | Y, b \sim \mathcal{W} \left(K, \widehat{b}\widehat{b}^\top \right) \quad (20)$$

where $\text{Inv-}\Gamma(s, r)$ denotes the inverse Gamma distribution with shape s and rate r , $\widehat{\beta}$ and $\widehat{\sigma}^2$ are the ML estimates of β and σ^2 , $\widehat{\text{var}}(\widehat{\beta})$ is the variance estimate associated with the ML estimator of β , and $\widehat{b} = (\widehat{b}_1, \dots, \widehat{b}_K)$ are the best linear unbiased predictions (BLUP) of the random effects, preferably estimated by using the restricted ML (REML) estimator Jolani et al. (2015).

With systematically missing values on variable y , ML estimates are evaluated from the clusters observed in y and degrees of freedom in (20) and (18) are modified accordingly.

Note that posterior distributions are conditional to the random coefficients that are unknown. Rigorously, simulation of these posterior distributions require the use of a Gibbs sampler. To avoid the use of such iterative algorithm, the marginal distribution of θ is approximated by the one of θ given a fixed value for b (*i.e.* $P(\theta | Y) \approx P(\theta | Y; b = \widehat{b})$). Assuming having mimicked the marginal posterior distribution of θ by making one draw, then we can draw random effects given this one. Furthermore, contrary to the original paper (Jolani et al., 2015) the prior distribution accounts for the dependence between the components σ^2

and β (Jolani, 2017).

Binary variables are imputed in the same way by considering a GLMM model with a logit link.

With sporadically missing continuous variables, steps (18-20) are essentially the same: the parameters of the posterior distribution are tuned by considering the ML estimator evaluated from individuals observed on variable j and degrees of freedom are suitably modified in (18).

A.3 FCS-2step

For a continuous incomplete variable y , FCS-2stage is based on the conditional imputation model:

$$\begin{aligned} y_k &= \mathbf{Z}_k (\boldsymbol{\beta} + b_k) + \varepsilon_k \\ b_k &\sim \mathcal{N}(0, \boldsymbol{\Psi}) \\ \varepsilon_k &\sim \mathcal{N}(0, \sigma_k^2 \mathbb{I}_{n_k}) \end{aligned}$$

The parameter of this model is $\boldsymbol{\theta} = (\boldsymbol{\beta}, \boldsymbol{\Psi}, (\sigma_k)_{1 \leq k \leq K})$.

FCS-2stage draws the parameters of the imputation model by using an asymptotic strategy: a two-stage estimator is evaluated from the observed data and the posterior distribution is then approximated. To fit the two-stage estimator, at stage one, the ML estimator of a linear model is computed on each available cluster:

$$\hat{\boldsymbol{\beta}}_k = (\mathbf{Z}_k^\top \mathbf{Z}_k)^{-1} \mathbf{Z}_k^\top \mathbf{y}_k$$

Then, at stage two, the following random effects model is used:

$$\hat{\boldsymbol{\beta}}_k = \boldsymbol{\beta} + b_k + \varepsilon'_k$$

with $b_k \sim \mathcal{N}(0, \boldsymbol{\Psi})$ and $\varepsilon'_k \sim \mathcal{N}\left(0, \sigma_k^2 (\mathbf{Z}_k \mathbf{Z}_k^\top)^{-1}\right)$. $\boldsymbol{\beta}$ and $\boldsymbol{\Psi}$ may be estimated by REML; alternatively, Resche-Rigon and White (2016) suggest using the method of moments (MM), which is even faster, especially with high dimensional $\boldsymbol{\beta}$ (DerSimonian and Laird, 1986; Jackson et al., 2013).

The two-stage estimator is typically used by assuming $(\sigma_k)_{1 \leq k \leq K}$ to be known (Burke et al., 2016). In our framework, $(\sigma_k)_{1 \leq k \leq K}$ is unknown and cannot be identified at stage one with systematically missing values. Thus, instead of fixing σ_k ($1 \leq k \leq K$), the assumption is made that $(\sigma_k)_{1 \leq k \leq K}$ are the independent realisations of a unique random variable so that

$$\log \sigma_k \sim \mathcal{N}(\log \sigma, \Phi). \quad (21)$$

As with estimation of β and Ψ in (8-9), $\log \sigma$ and Φ are estimated by a two-stage method. At stage one, the ML estimate of σ_k , denoted $\hat{\sigma}_k$, is computed and log-transformed on each cluster, and its variance is derived (using the Delta method). At stage two, these estimates are used in the random effects model

$$\log \hat{\sigma}_k = \log \sigma + s_k + \varepsilon_k'' \quad (22)$$

with $s_k \sim \mathcal{N}(0, \Phi)$, $\varepsilon_k'' \sim \mathcal{N}(0, \text{var}(\log \hat{\sigma}_k))$.

To draw missing values of y from their predictive distribution using this estimator, first, $\log \sigma$ and Φ are estimated by fitting the two-stage estimator to the data and their posterior distribution is then approximated by

$$\log \sigma | Y^{\text{obs}} \sim \mathcal{N}\left(\widehat{\log \sigma}, \widehat{\text{var}}\left(\widehat{\log \sigma}\right)\right) \quad (23)$$

$$\Phi | Y^{\text{obs}} \sim \mathcal{N}\left(\hat{\Phi}, \widehat{\text{var}}\left(\hat{\Phi}\right)\right). \quad (24)$$

Next, from model (22), s_k can be straightforwardly drawn conditionally on $\log \hat{\sigma}_k$ since $\log \hat{\sigma}_k$ and s_k are Gaussian, providing values for $(\sigma_k)_{1 \leq k \leq K}$. Then, β and Ψ are drawn from their posterior distribution in the same way:

$$\beta | Y^{\text{obs}}, \sigma_k \sim \mathcal{N}\left(\hat{\beta}, \sigma_k^2 \left(\mathbf{Z}_k \mathbf{Z}_k^\top\right)^{-1}\right) \quad (25)$$

$$\Psi^{\text{chol}} | Y^{\text{obs}}, \sigma_k \sim \mathcal{N}\left(\hat{\Psi}^{\text{chol}}, \widehat{\text{var}}\left(\hat{\Psi}^{\text{chol}}\right)\right) \quad (26)$$

where Ψ^{chol} is the Cholesky decomposition of Ψ .

To handle binary variables with both sporadically and systematically missing values by applying a logit link in the imputation model (7). In this case the parameters σ_k are not used, thus the distribution assumption (21) does not hold and the two-stage estimator described in step (22) is not needed.

B Description of GREAT data

The GREAT Network performed an IPD meta-analysis to explore risk factors associated with short-term mortality in acute heart failure (AHF) ([GREAT Network, 2013](#)). Their dataset consists of 28 studies: 8 were carried out in Western Europe (2 in Italy, 2 in Spain, and 1 in each of France, Finland, Switzerland, Netherlands), 13 in Central Europe (12 in Czech Republic and 1 in Austria), 3 in America (2 in the United States and 1 in Argentina), 3 in Asia (China, Japan, Korea), and 1 in Africa (Tunisia) ([Mebazaa et al., 2013](#)). The principal investigators of each study provided the original data collected for each patient, including a list of patient characteristics and potential risk factors ([Lassus et al., 2013](#)).

One biomarker of interest was brain natriuretic peptide (BNP), which is known to be elevated in acute heart failure. Since measuring the left ventricular ejection fraction (LVEF) requires an ultrasound examination, one objective consists in explaining LVEF from biomarkers that are easier to measure,

such as BNP or electrocardiographic characteristics such as the atrial fibrillation (AFIB). The generalized linear mixed effect model (GLMM) (Pinheiro and Bates, 2000; Lee et al., 2006) is a suitable statistical model to achieve this goal.

The dataset contains 2 binary variables (AFIB and Gender) and 7 continuous variables (BMI, Age, Systolic blood pressure (SBP), Diastolic blood pressure (DBP), Heart rate (HR), LVEF and BNP). Variables are described in Table 8 in Appendix B. The total number of individuals is 11685 and study sizes range from 18 to 1834.

Each study is incomplete, leading to sporadically missing values on all variables except gender and LVEF. However, BNP measurement is a recent technique, so this variable has been collected on 10 studies only, leading to systematically missing values. Four other variables are systematically missing on some studies (Table 7), notably the binary variable AFIB.

Table 7: GREAT data: percentages of missing values by variable and study.

k	n_k	Gender	BMI	Age	SBP	DBP	HR	BNP	AFIB	LVEF
1	410	0	36	< 1	1	2	3	57	< 1	0
2	567	0	19	0	2	3	1	10	0	0
3	210	0	43	0	1	2	1	0	100	0
4	375	0	2	0	1	1	2	4	42	0
5	107	0	1	0	0	0	0	100	0	0
6	267	0	100	0	100	100	< 1	100	0	0
7	203	0	< 1	0	1	2	1	< 1	0	0
8	354	0	44	1	16	16	19	12	22	0
9	137	0	1	0	0	0	0	0	0	0
10	48	0	100	0	0	0	4	100	0	0
11	208	0	24	0	0	< 1	0	100	0	0
12	622	0	27	0	< 1	< 1	1	100	0	0
13	78	0	60	0	0	0	0	100	100	0
14	670	0	77	< 1	1	1	2	100	< 1	0
15	1000	0	13	0	2	2	2	82	< 1	0
16	1093	0	0	0	0	0	0	100	0	0
17	18	0	6	0	0	0	0	22	0	0
18	1834	0	19	0	1	1	< 1	92	< 1	0
19	358	0	7	0	0	0	0	99	0	0
20	54	0	6	0	2	2	2	100	2	0
21	588	0	10	0	< 1	< 1	0	97	< 1	0
22	651	0	24	0	2	2	2	73	2	0
23	455	0	2	0	0	0	< 1	86	< 1	0
24	294	0	4	0	< 1	< 1	< 1	81	0	0
25	397	0	1	0	0	0	0	100	0	0
26	295	0	11	0	0	0	0	66	0	0
27	303	0	11	0	< 1	< 1	0	79	0	0
28	89	0	0	0	0	0	0	38	0	0

Table 8: GREAT data: description of variables. Binary variables are presented by counts and percentages, while continuous variables by their median and quartiles.

variable	value	size	summary
gender	0	6865	58.65%
	1	4820	42.35%
AFIB	0	7704	69.18%
	1	3431	30.81%
BMI		9259	26.58 [23.66143862 ; 30.12]
Age		11678	72.7 [62.7 ; 80]
SBP		11278	130 [111 ; 153]
DBP		11262	80 [68 ; 90]
HR		11518	87 [72 ; 105]
LVEF		11685	0.38 [0.27 ; 0.5]
BNP		2776	2.99 [2.66 ; 3.29]

C Simulation

C.1 Simulation design

C.1.1 Investigated configurations

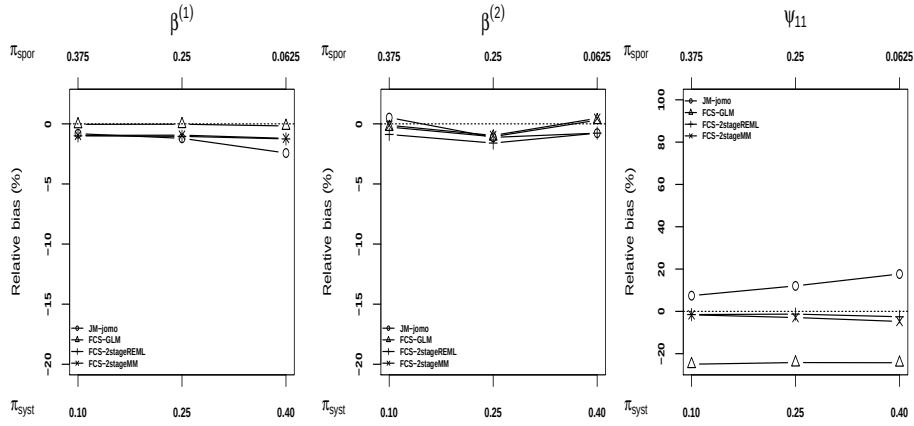
Table 9: Tuning parameters for the several configurations: K the number of clusters, n the number of individuals, λ a scalar multiplying the variance of random effects, ρ_v the correlation between random intercepts generating variables x_1, x_2, x_3 , ρ_b the correlation between random coefficients of the analysis model, the type of the outcome, the type of the covariates (x_1 or x_2), the proportion of systematically missing values on the covariates (x_1 or x_2), the proportion of sporadically missing values on the covariates (x_1 or x_2), the nature of the missing data mechanism R , the intra cluster correlation for continuous variables (ICC), ρ_{x_1, x_3} the correlation between x_1 and x_3 in each cluster, the presence of a random effect on the covariate x_2 , the link function used to generate the binary covariate. Parameters varying from the base-case configuration are in boldface.

case	y					x_1				x_2				link function			
	K	n	λ	ρ_v	ρ_b	type	π_{sys}	π_{spor}	R	ICG	$p(x_1, x_3)$	type	π_{sys}		π_{spor}	R	random effect on x_2
1 (base-case)	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
2	28	11685	1	.07	-.87	cont	.1	.375	MCAR	.25	.3	bin	.1	.375	MCAR	no	logistic
3	28	11685	1	.07	-.87	cont	.4	.0625	MCAR	.25	.3	bin	.4	.0625	MCAR	no	logistic
4	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
5	28	11685	1	.07	-.87	bin	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
6	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
7	28	11685	1	.07	-.87	cont	0	0	MAR	.25	.3	bin	.25	.25	MCAR	no	logistic
8	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MAR	no	logistic
9	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	yes	logistic
10	14	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
11	28	5845	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
12	28	2923	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
13	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.5	.3	bin	.25	.25	MCAR	no	logistic
14	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.1	.3	bin	.25	.25	MCAR	no	logistic
15	28	11685	1	.3	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
16	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.5	bin	.25	.25	MCAR	no	logistic
17	28	11685	2	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
18	28	11685	1	.07	-.3	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic
19	28	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	probit
20	7	11685	1	.07	-.87	cont	.25	.25	MCAR	.25	.3	bin	.25	.25	MCAR	no	logistic

C.2 Complementary results

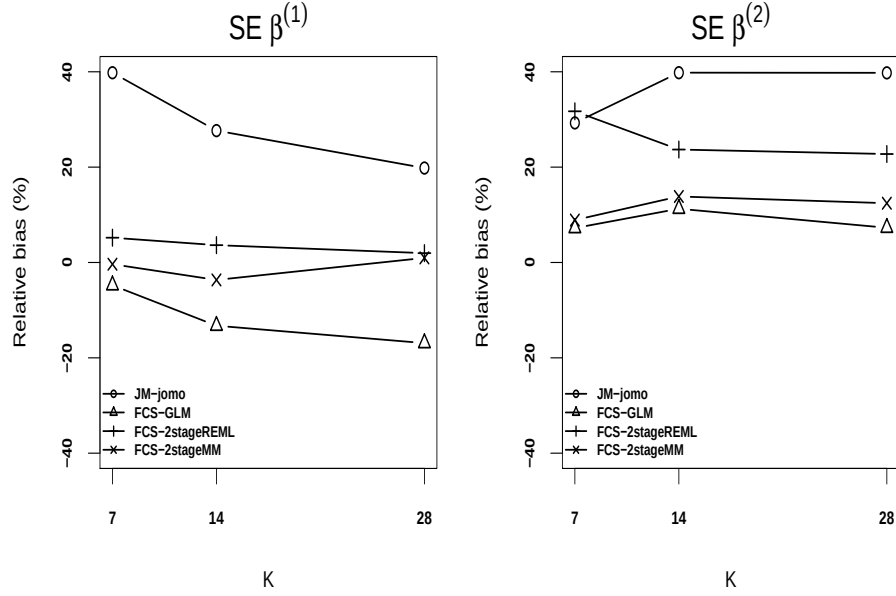
C.2.1 Robustness to the proportion of systematically missing values

Figure 4: Robustness to the proportion of systematically missing values: relative bias for the estimate of $\beta^{(1)}$ (left), $\beta^{(2)}$ (middle) and ψ_{11} (right) according to π_{syst} for each MI method. The proportion of sporadically missing values is modified to keep a constant proportion of missing values (in expectation).



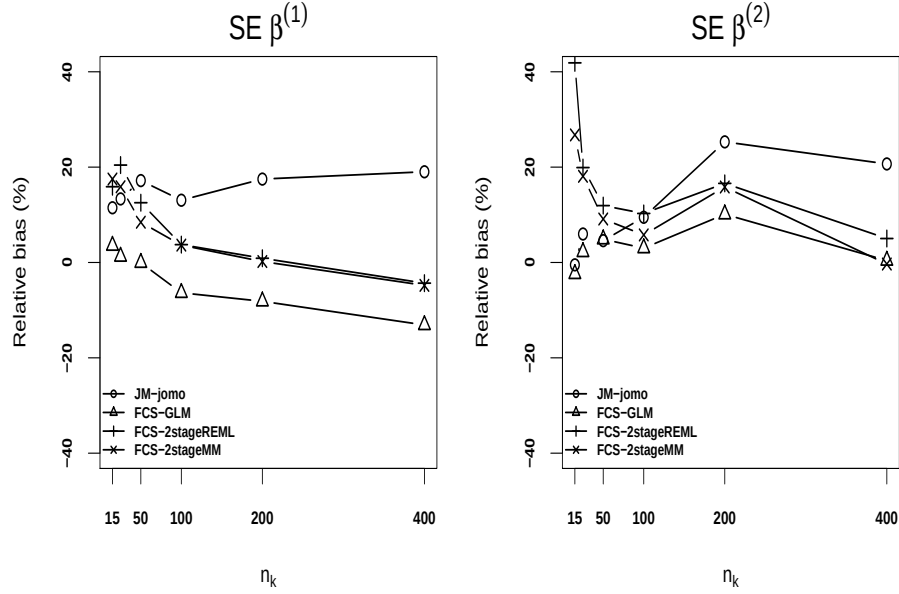
C.2.2 Robustness to the number of clusters

Figure 5: Robustness to the number of clusters: estimate of the relative bias for the SE estimate for $\widehat{\beta}^{(1)}$ (left), $\widehat{\beta}^{(2)}$ (right) according to K for each MI method. The estimated relative bias is calculated by the difference between the model SE and the empirical SE, divided by the empirical SE. Criteria are based on 500 incomplete datasets.



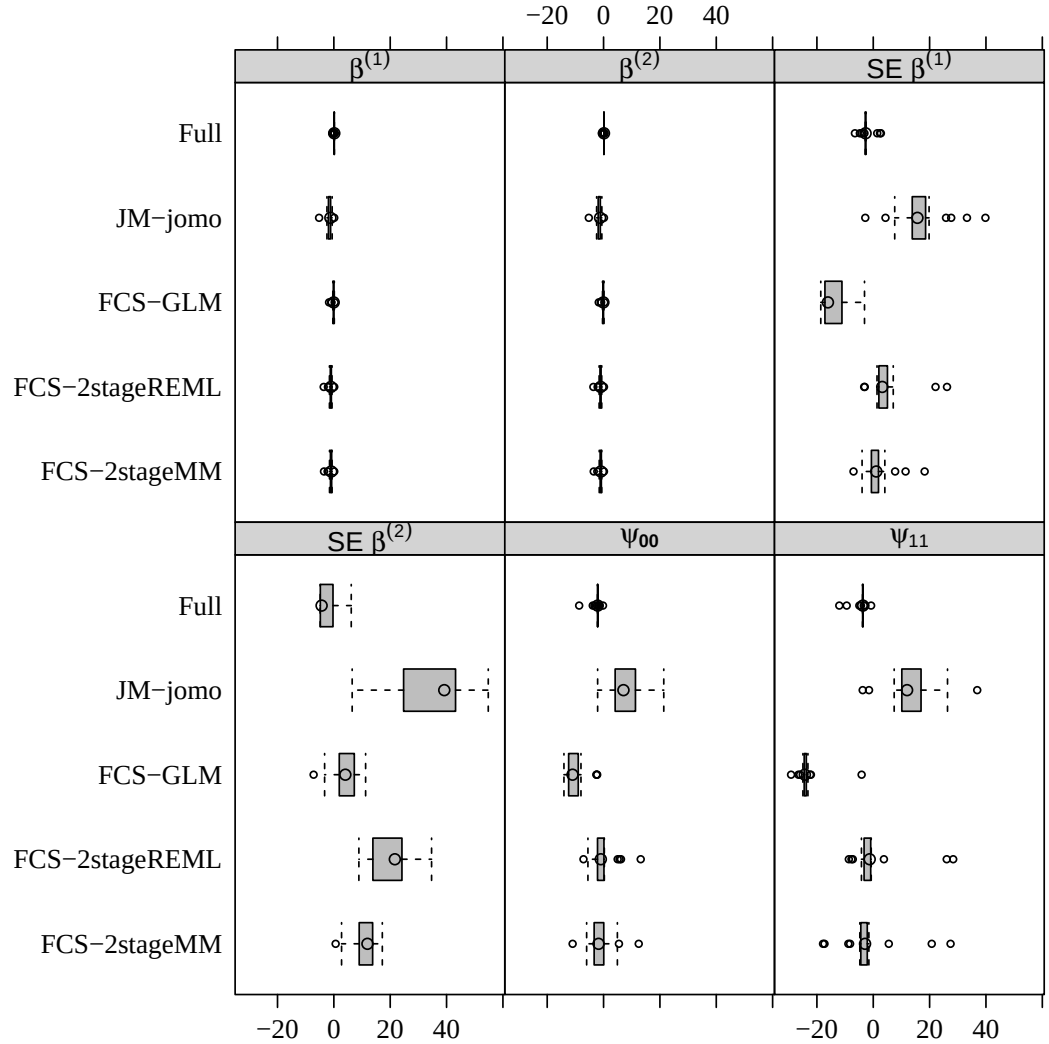
C.2.3 Robustness to the cluster size

Figure 6: Robustness to the cluster size: estimate of the relative bias for the SE estimate for $\widehat{\beta}^{(1)}$ (left), $\widehat{\beta}^{(2)}$ (right) according to n_k for each MI method. The estimated relative bias is calculated by the difference between the model SE and the empirical SE, divided by the empirical SE. Criteria are based on 500 incomplete datasets.



C.2.4 Other configurations

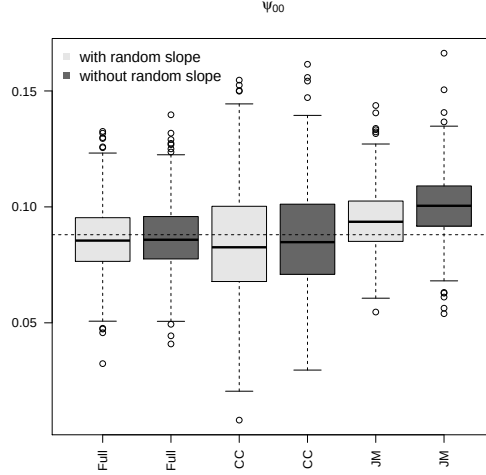
Figure 7: Distribution of the relative bias over the 19 configurations for several methods (Full, CC, JM-jomo, FCS-GLM, FCS-2stageREML, FCS-2stageMM) and several parameters of interest ($\beta^{(1)}, \beta^{(2)}, \psi_{00}, \psi_{11}, \text{SE } \beta^{(1)}, \text{SE } \beta^{(2)}$). One point represents the relative bias observed for one configuration.



C.2.5 Influence of the random slope

The multivariate version of model (1) used in JM-jomo requires that all missing variables are in the left hand side of the model (2). However, if the analysis model includes a random slope in the right hand side (like in the simulation study 3.1), then the imputation model is misspecified. To assess the influence of this misspecification on the random effect variance, we compare the estimates of ψ_{00} when the outcome of the model is generated according to an analysis model including a random slope (base-case configuration), with the estimates of ψ_{00} when the outcome of the model is generated according to an analysis model with a random intercept only. Estimates are reported in Figure 8.

Figure 8: Distribution of the estimate of $\sqrt{\psi_{00}}$ over the 500 generated datasets for Full, CC and JM-jomo methods. Both configurations are considered: a one with a random slope (in grey, corresponding to the base-case configuration) and one without random slope (in blue). The red dashed line represents the true value of $\sqrt{\psi_{00}}$.



A bias is observed even if the outcome is generated from a model with no random slope, indicating that misspecification of the imputation model is not the main reason for the observed bias in the base-case configuration. As shown in Section 3.2.6 it is more likely that the use of wrongly informative prior distributions biases the inference in presence of very small values for the level-2 variances.