

UNIVERSIDAD DE CONCEPCIÓN



CENTRO DE INVESTIGACIÓN EN INGENIERÍA MATEMÁTICA (CI²MA)



Multivariate measurement error models based on student-t
distribution under censored responses

CELSO R. B. CABRAL, LUIS M. CASTRO,
VÍCTOR H. LACHOS, LARISSA A. MATOS

PREPRINT 2015-36

SERIE DE PRE-PUBLICACIONES

Multivariate Measurement Error Models Based on Student-t Distribution under Censored Responses

Larissa A. Matos^a, Luis M. Castro^b, Celso R. B. Cabral^{c*} and Víctor H. Lachos^a

^a Department of Statistics, Universidade Estadual de Campinas, Campinas, São Paulo, Brazil

^b Department of Statistics and CI²MA, Universidad de Concepción, Concepción, Chile

^c Department of Statistics, Universidade Federal do Amazonas, Amazonas, Brazil

Abstract

Measurement error models constitute a wide class of models, that include linear and nonlinear regression models. They are very useful to model many real life phenomena, particularly in the medical and biological areas. The great advantage of these models is that, in some sense, they can be represented as mixed effects models, allowing to us the implementation of well-known techniques, like the EM-algorithm for the parameter estimation. In this paper, we consider a class of multivariate measurement error models where the observed response and/or covariate are not fully observed, i.e., the observations are subject to certain threshold values below or above which the measurements are not quantifiable. Consequently, these observations are considered censored. We assume a Student-t distribution for the unobserved true values of the mismeasured covariate and the error term of the model, providing a robust alternative for parameter estimation. Our approach relies on a likelihood-based inference using the EM-algorithm. The proposed method is illustrated through simulation studies and the analysis of a real dataset.

Keywords: Censored responses; EM algorithm; Measurement error models; Student- t distribution.

1. Introduction

Measurement error – hereafter ME – models (also known as error-in-variables models) are defined as regression models where the covariates cannot be measured/observed directly, or are measured with a substantial error. From a practical point of view, such models are very useful because they take into account some notions of randomness inherent to the covariates. For example, in AIDS studies, linear and nonlinear mixed effects models are typically considered to study the relationship between the viral load (HIV-1 RNA) measures and CD4 (T-cells)

* *Address for correspondence:* Celso R. B. Cabral, Departamento de Estatística, ICE, Universidade Federal do Amazonas. Av. Rodrigo Otávio, 6200, Campus Universitário Arthur Virgílio Filho, Coroado I, CEP 69077-000, Manaus, Amazonas, Brazil. E-mail address: celsoromulo@ufam.edu.br

cell count. However, as pointed out by many authors (see for instance Wu, 2010; Bandyopadhyay *et al.*, 2015, and many others), this covariate is measured (in general) with substantial error. This is because, in most HIV clinical trials, cell counts are measured periodically with a substantial amount of variability.

A wide variety of proposals exist in the statistical literature trying to deal with the presence of ME in multivariate data. For example, Carrol *et al.* (1997) proposed a generalized linear mixed ME model and Buonaccorsi *et al.* (2000) (see also Dumitrescu, 2010) studied estimation of the variance components in a linear mixed effect model with ME in a time varying covariate. Zhang *et al.* (2011) introduced a multivariate ME model including the presence of zero inflation. Recently, Abarin *et al.* (2014) proposed a method of moments for the parameter estimation in the linear mixed effect with ME model. Moreover, Cabral *et al.* (2014) studied a multivariate ME model using finite mixtures of skew Student-t distributions. A comprehensive review of ME models can be found in the books of Fuller (1987), Cheng & Van Ness (1999), Carroll *et al.* (2006) and Buonaccorsi (2010).

Although many models for multivariate data consider the existence of mismeasured covariates, many of them do not consider censored observations or detection limits for the response variable. This aspect is relevant, since in many studies the observed response is subject to maximum/minimum detection limits. For that reason, clearly there is a need for a new methodology that takes into account censored responses in multivariate data and mismeasured covariates at the same time. We propose an approach where the random observational errors and the unobserved latent variable are jointly modeled by a Student-t distribution, which has heavier tails than the normal one. Besides this, our estimation approach relies on an exact EM-type algorithm, providing explicit expressions for the E and M steps, obtaining as byproduct the asymptotic covariance of the maximum likelihood estimates. To illustrate the applicability of the method, we analyze a real dataset consisting of measurements of the testicular volume of 42 adolescents using five different techniques.

The paper is organized as follows. Section 2 presents some results about the multivariate Student-t distribution, focusing on its truncated version. Section 3 proposes the ME model for censored multivariate responses under the Student-t distribution. Sections 4 and 5 present the likelihood-based estimation and standard errors of the parameter estimates in the proposed model via an EM-type algorithm, respectively. Section 6 presents the results of simulation studies conducted to examine the performance of the proposed method with respect to the asymptotic properties of the ML estimates, and the consequences of the inappropriateness of the normality assumption. The analysis of a real dataset is presented in Section 7.

2. The multivariate Student-t distribution and truncated related ones

We say that the random vector $\mathbf{Y} : p \times 1$ has a *Student-t distribution* with location vector $\boldsymbol{\mu}$, dispersion matrix $\boldsymbol{\Sigma}$ and ν degrees of freedom, when its probability density function (pdf) is given by

$$t_p(\mathbf{y}|\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu) = \frac{\Gamma(\frac{p+\nu}{2})}{\Gamma(\frac{\nu}{2})\pi^{p/2}} \nu^{-p/2} |\boldsymbol{\Sigma}|^{-1/2} \left(1 + \frac{d_{\boldsymbol{\Sigma}}(\mathbf{y}, \boldsymbol{\mu})}{\nu} \right)^{-(p+\nu)/2}, \quad (1)$$

where $\Gamma(\cdot)$ is the standard gamma function and

$$d_{\boldsymbol{\Sigma}}(\mathbf{y}, \boldsymbol{\mu}) = (\mathbf{y} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \boldsymbol{\mu}),$$

is the Mahalanobis distance. The cumulative distribution function (cdf) of $\mathbf{Y}$ is denoted by $T_p(\cdot | \boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)$. If $\nu > 1$, $\boldsymbol{\mu}$ is the mean of $\mathbf{Y}$, and if $\nu > 2$, $\nu(\nu - 2)^{-1}\boldsymbol{\Sigma}$ is its covariance matrix. We use the notation $\mathbf{Y} \sim t_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)$.

It is possible to show that $\mathbf{Y}$ admits the stochastic representation

$$\mathbf{Y} = \boldsymbol{\mu} + U^{-1/2}\mathbf{Z}, \quad \mathbf{Z} \sim N_p(\mathbf{0}, \boldsymbol{\Sigma}), \quad U \sim \text{Gamma}(\nu/2, \nu/2), \quad (2)$$

where $\mathbf{Z}$ and U are independent, and $\text{Gamma}(a, b)$ denotes the gamma distribution with mean a/b . As ν tends to infinity, U converges to one with probability one and $\mathbf{Y}$ is approximately distributed as a $N_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ distribution. From this representation we can easily deduce that an affine transformation $\mathbf{AY} + \mathbf{b}$ has a $t_q(\mathbf{A}\boldsymbol{\mu} + \mathbf{b}, \mathbf{A}\boldsymbol{\Sigma}\mathbf{A}^\top, \nu)$ distribution, where $\mathbf{A}$ is a $q \times p$ matrix and $\mathbf{b}$ is a q -dimensional vector. For a reference with extensive material regarding the multivariate Student-t distribution, see Kotz & Nadarajah (2004).

The following result shows that the Student-t family of distributions is closed under marginalization and conditioning. The proof can be found in Matos *et al.* (2013, Proposition 1).

Proposition 1. *Let $\mathbf{Y} \sim t_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)$. Consider the partition $\mathbf{Y} = (\mathbf{Y}_1^\top, \mathbf{Y}_2^\top)^\top$, with $\mathbf{Y}_1 : p_1 \times 1$ and $\mathbf{Y}_2 : p_2 \times 1$. Accordingly, consider the partitions $\boldsymbol{\mu} = (\boldsymbol{\mu}_1^\top, \boldsymbol{\mu}_2^\top)^\top$ and $\boldsymbol{\Sigma} = (\boldsymbol{\Sigma}_{ij})$, $i, j = 1, 2$. Then*

$$(i) \quad \mathbf{Y}_1 \sim t_{p_1}(\boldsymbol{\mu}_1, \boldsymbol{\Sigma}_{11}, \nu),$$

$$(ii) \quad \mathbf{Y}_2 | \mathbf{Y}_1 = \mathbf{y}_1 \sim t_{p_2}(\boldsymbol{\mu}_{2.1}, \tilde{\boldsymbol{\Sigma}}_{22.1}, \nu + p_1),$$

where

$$\begin{aligned} \boldsymbol{\mu}_{2.1} &= \boldsymbol{\mu}_2 + \boldsymbol{\Sigma}_{21}\boldsymbol{\Sigma}_{11}^{-1}(\mathbf{y}_1 - \boldsymbol{\mu}_1), \quad \tilde{\boldsymbol{\Sigma}}_{22.1} = \frac{\nu + d_{\boldsymbol{\Sigma}_{11}}(\mathbf{y}_1, \boldsymbol{\mu}_1)}{\nu + p_1}\boldsymbol{\Sigma}_{22.1}, \quad \text{and} \\ \boldsymbol{\Sigma}_{22.1} &= \boldsymbol{\Sigma}_{22} - \boldsymbol{\Sigma}_{21}\boldsymbol{\Sigma}_{11}^{-1}\boldsymbol{\Sigma}_{12}. \end{aligned}$$

Let $\mathbf{Y} \sim t_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)$ and $\mathbb{D}$ be a Borel set in $\mathbb{R}^p$. We say that the random vector $\mathbf{Z}$ has a *truncated Student-t distribution on $\mathbb{D}$* when $\mathbf{Z}$ has the same distribution as $\mathbf{Y} | (\mathbf{Y} \in \mathbb{D})$. In this case, the pdf of $\mathbf{Z}$ is given by

$$Tt_p(\mathbf{z} | \boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu; \mathbb{D}) = \frac{t_p(\mathbf{z} | \boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)}{P(\mathbf{Y} \in \mathbb{D})} \mathbb{I}_{\mathbb{D}}(\mathbf{z}),$$

where $\mathbb{I}_{\mathbb{D}}(\cdot)$ is the indicator function of $\mathbb{D}$, that is, $\mathbb{I}_{\mathbb{D}}(\mathbf{z}) = 1$ if $\mathbf{z} \in \mathbb{D}$ and $\mathbb{I}_{\mathbb{D}}(\mathbf{z}) = 0$ otherwise. We use the notation $\mathbf{Z} \sim Tt_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu; \mathbb{D})$. If $\mathbb{D}$ has the form

$$\mathbb{D} = \{(x_1, \dots, x_p) \in \mathbb{R}^p; \ x_1 \leq d_1, \dots, x_p \leq d_p\}, \quad (3)$$

then we use the notation $(\mathbf{Y} \in \mathbb{D}) = (\mathbf{Y} \leq \mathbf{d})$, where $\mathbf{d} = (d_1, \dots, d_p)^\top$. In this case, $P(\mathbf{Y} \leq \mathbf{d}) = T_p(\mathbf{d} | \boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)$. Notice that we can have $d_i = +\infty$, $i = 1, \dots, p$.

The following propositions are crucial to obtain the expectations in the E step of the EM type algorithm, which will be used to compute maximum likelihood estimates of the parameters in the model proposed in this work. Proofs can be found in Matos *et al.* (2013, Propositions 2 and 3). We will use the notations $\mathbf{Z}^{(0)} = 1$, $\mathbf{Z}^{(1)} = \mathbf{Z}$ and $\mathbf{Z}^{(2)} = \mathbf{Z}\mathbf{Z}^\top$.

Proposition 2. Let $\mathbf{Z} \sim \text{Tt}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu; \mathbb{D})$, where $\mathbb{D}$ is as in (3). Then, for $k = 0, 1, 2$,

$$\mathbb{E} \left[\left(\frac{\nu + p}{\nu + d_{\boldsymbol{\Sigma}}(\mathbf{Z}, \boldsymbol{\mu})} \right)^r \mathbf{Z}^{(k)} \right] = c_p(\nu, r) \frac{T_p(\mathbf{d}|\boldsymbol{\mu}, \boldsymbol{\Sigma}^*, \nu + 2r)}{T_p(\mathbf{d}|\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu)} \mathbb{E}[\mathbf{Y}^{(k)}], \quad (4)$$

where $\nu + 2r > 0$ and

$$\begin{aligned} \mathbf{Y} &\sim \text{Tt}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}^*, \nu + 2r; \mathbb{D}), \\ \boldsymbol{\Sigma}^* &= \frac{\nu}{\nu + 2r} \boldsymbol{\Sigma}, \\ c_p(\nu, r) &= \left(\frac{\nu + p}{\nu} \right)^r \left(\frac{\Gamma((p + \nu)/2) \Gamma((\nu + 2r)/2)}{\Gamma(\nu/2) \Gamma((p + \nu + 2r)/2)} \right). \end{aligned} \quad (5)$$

Observe that the computation of the expectation on the left side of (4) is reduced to the computation of the moments of the truncated Student-t distribution in (5). These moments are available in closed form in Ho *et al.* (2012) and the implementations were done using the R package *TTmoment()*, available on CRAN.

Proposition 3. Let $\mathbf{Z} \sim \text{Tt}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu; \mathbb{D})$, where $\mathbb{D}$ is as in (3). Consider the partition $\mathbf{Z} = (\mathbf{Z}_1^\top, \mathbf{Z}_2^\top)^\top$, with $\mathbf{Z}_1 : p_1 \times 1$ and $\mathbf{Z}_2 : p_2 \times 1$. Accordingly, consider the partitions $\boldsymbol{\mu} = (\boldsymbol{\mu}_1^\top, \boldsymbol{\mu}_2^\top)^\top$ and $\boldsymbol{\Sigma} = (\boldsymbol{\Sigma}_{ij})$, $i, j = 1, 2$. Then ,

$$\begin{aligned} \mathbb{E} \left[\left(\frac{\nu + p}{\nu + d_{\boldsymbol{\Sigma}}(\mathbf{Z}, \boldsymbol{\mu})} \right)^r \mathbf{Z}_2^{(k)} | \mathbf{Z}_1 = \mathbf{z}_1 \right] &= \frac{h_p(p_1, \nu, r)}{(\nu + d_{\boldsymbol{\Sigma}_{11}}(\mathbf{z}_1, \boldsymbol{\mu}_1))^r} \\ &\times \frac{T_{p_2}(\mathbf{d}_2 | \boldsymbol{\mu}_{2,1}, \tilde{\boldsymbol{\Sigma}}_{22,1}^*, \nu + p_1 + 2r)}{T_{p_2}(\mathbf{d}_2 | \boldsymbol{\mu}_{2,1}, \tilde{\boldsymbol{\Sigma}}_{22,1}, \nu + p_1)} \mathbb{E}[\mathbf{Y}^{(k)}], \end{aligned}$$

where $\nu + p_1 + 2r > 0$, $\mathbf{d}_2 = (d_{p_1+1}, \dots, d_p)^\top$,

$$\begin{aligned} \mathbf{Y} &\sim \text{Tt}_{p_2}(\boldsymbol{\mu}_{2,1}, \tilde{\boldsymbol{\Sigma}}_{22,1}^*, \nu + p_1 + 2r; \mathbb{D}_2), \\ \mathbb{D}_2 &= \{(x_{p_1+1}, \dots, x_p) \in \mathbb{R}^{p_2}; x_{p_1+1} \leq d_{p_1+1}, \dots, x_p \leq d_p\}, \\ \tilde{\boldsymbol{\Sigma}}_{22,1}^* &= \frac{\nu + d_{\boldsymbol{\Sigma}_{11}}(\mathbf{z}_1, \boldsymbol{\mu}_1)}{\nu + p_1 + 2r} \boldsymbol{\Sigma}_{22,1}, \\ h_p(p_1, \nu, r) &= (\nu + p)^r \left(\frac{\Gamma((p + \nu)/2) \Gamma((p_1 + \nu + 2r)/2)}{\Gamma((p_1 + \nu)/2) \Gamma((p + \nu + 2r)/2)} \right), \end{aligned}$$

$\boldsymbol{\mu}_{2,1}$, $\boldsymbol{\Sigma}_{22,1}$ and $\tilde{\boldsymbol{\Sigma}}_{22,1}$ are given in Proposition 1.

3. Model specification

Let $\mathbf{Y}_i = (Y_{i1}, \dots, Y_{ir})^\top$ be the vector of responses for the i th experimental unit, where Y_{ij} is the j th observed response of unit i (for $i = 1, \dots, n$ and $j = 1, \dots, r$). Let X_i be the i th observed value of the covariate and x_i be the unobserved (true) covariate value for unit i . Following Barnett (1969), the multivariate ME model is formulated as

$$X_i = x_i + \xi_i \quad (6)$$

and

$$\mathbf{Y}_i = \boldsymbol{\alpha} + \boldsymbol{\beta}x_i + \mathbf{e}_i, \quad (7)$$

where $\mathbf{e}_i = (e_{i1}, \dots, e_{ir})^\top$ is a vector of measurement errors, $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_r)^\top$ and $\boldsymbol{\beta} = (\beta_1, \dots, \beta_r)^\top$ are vectors with regression parameters. Let $\boldsymbol{\epsilon}_i = (\xi_i, \mathbf{e}_i^\top)^\top$ and $\mathbf{Z}_i = (X_i, \mathbf{Y}_i^\top)^\top = (Z_{i1}, \dots, Z_{ip})^\top$. Then, equations (6) and (7) imply

$$\mathbf{Z}_i = \mathbf{a} + \mathbf{b}x_i + \boldsymbol{\epsilon}_i = \mathbf{a} + \mathbf{B}\mathbf{r}_i, \quad i = 1, \dots, n, \quad (8)$$

where $\mathbf{a} = (0, \boldsymbol{\alpha}^\top)^\top$ and $\mathbf{b} = (1, \boldsymbol{\beta}^\top)^\top$ are $p \times 1$ vectors, with $p = r + 1$, $\mathbf{B} = [\mathbf{b}; \mathbf{I}_p]$ is a $p \times (p + 1)$ matrix and $\mathbf{r}_i = (x_i, \boldsymbol{\epsilon}_i^\top)^\top$. Thus, from equation (8), the distribution of $\mathbf{Z}_i$ becomes specified once the distribution of $\mathbf{r}_i$ is specified. Usually, a normality assumption is made, such that

$$\mathbf{r}_i \stackrel{\text{iid}}{\sim} N_{1+p} \left(\begin{pmatrix} \mu_x \\ \mathbf{0}_p \end{pmatrix}, \begin{pmatrix} \sigma_x^2 & \mathbf{0}_p^\top \\ \mathbf{0}_p & \boldsymbol{\Omega} \end{pmatrix} \right), \quad i = 1, \dots, n, \quad (9)$$

where $\mathbf{0}_p = (0, \dots, 0)^\top : p \times 1$, $\boldsymbol{\Omega} = \text{diag}(\omega_1^2, \dots, \omega_p^2)$ and $\stackrel{\text{iid}}{\sim}$ denotes independent and identically distributed random vectors. Marginally, we have that $x_i \stackrel{\text{iid}}{\sim} N(\mu_x, \sigma^2)$ and $\boldsymbol{\epsilon}_i \stackrel{\text{iid}}{\sim} N_r(\mathbf{0}, \boldsymbol{\Omega})$ are independent for all $i = 1, \dots, n$. For more details see, for example, Fuller (1987, Sec. 4.1).

In order to obtain robust estimation of the parameters in the model, we propose to replace assumption (9) by

$$\mathbf{r}_i = \begin{bmatrix} x_i \\ \boldsymbol{\epsilon}_i \end{bmatrix} \stackrel{\text{iid}}{\sim} t_{1+p} \left(\begin{pmatrix} \mu_x \\ \mathbf{0}_p \end{pmatrix}, \begin{pmatrix} \sigma_x^2 & \mathbf{0}_p^\top \\ \mathbf{0}_p & \boldsymbol{\Omega} \end{pmatrix}, \nu \right), \quad i = 1, \dots, n. \quad (10)$$

By (2), this formulation implies

$$\begin{aligned} \begin{bmatrix} x_i \\ \boldsymbol{\epsilon}_i \end{bmatrix} \mid U_i = u_i &\sim N_{1+p} \left(\begin{pmatrix} \mu_x \\ \mathbf{0}_p \end{pmatrix}, u_i^{-1} \begin{pmatrix} \sigma_x^2 & \mathbf{0}_p^\top \\ \mathbf{0}_p & \boldsymbol{\Omega} \end{pmatrix} \right), \\ U_i &\sim \text{Gamma} \left(\frac{\nu}{2}, \frac{\nu}{2} \right), \end{aligned}$$

for $i = 1, \dots, n$. Consequently,

$$x_i \mid U_i = u_i \stackrel{\text{ind}}{\sim} N(\mu_x, u_i^{-1} \sigma_x^2) \quad \text{and}, \quad (11)$$

$$\boldsymbol{\epsilon}_i \mid U_i = u_i \stackrel{\text{ind}}{\sim} N_p(\mathbf{0}_p, u_i^{-1} \boldsymbol{\Omega}). \quad (12)$$

Besides this, $\boldsymbol{\epsilon}_i$ and x_i have Student-t marginal distributions, with $\boldsymbol{\epsilon}_i \sim t_p(\mathbf{0}, \boldsymbol{\Omega}, \nu)$ and $x_i \sim t(\mu_x, \sigma_x^2, \nu)$.

Since for each i , $\boldsymbol{\epsilon}_i$ and x_i are indexed by the same scale mixing factor U_i , they are not independent in general. The independence corresponds to the case where $U_i = 1$ (normal case). However, conditional on U_i , $\boldsymbol{\epsilon}_i$ and x_i are independent for each $i = 1, \dots, n$, which implies that $\boldsymbol{\epsilon}_i$ and x_i are not correlated, since $\text{Cov}(\boldsymbol{\epsilon}_i, x_i) = E[\boldsymbol{\epsilon}_i x_i \mid U_i] = \mathbf{0}$. By (8), $\mathbf{Z}_i$ is an affine transformation of $\mathbf{r}_i$. Thus, its distribution is given by

$$\mathbf{Z}_i \sim t_p(\boldsymbol{\mu}_z, \boldsymbol{\Sigma}_z, \nu), \quad i = 1, \dots, n, \quad (13)$$

where

$$\boldsymbol{\mu}_z = \mathbf{a} + \mathbf{b}\mu_x \quad \text{and} \quad \boldsymbol{\Sigma}_z = \sigma_x^2 \mathbf{b}\mathbf{b}^\top + \boldsymbol{\Omega}. \quad (14)$$

As mentioned earlier, our model considers censored observations. Following Matos *et al.* (2013), we consider the case in which the response Z_{ij} is not fully observed for all i, j . What we truly observe, for each $i = 1, \dots, n$, is the random vector $\mathbf{V}_i = (V_{i1}, \dots, V_{ip})^\top$, such that $V_{ij} = \max\{Z_{ij}, \kappa_{ij}\}$, where κ_{ij} is a censoring level, that is,

$$V_{ij} = \begin{cases} Z_{ij} & \text{if } Z_{ij} > \kappa_{ij} \\ \kappa_{ij} & \text{if } Z_{ij} \leq \kappa_{ij}. \end{cases} \quad (15)$$

The model defined by Equations (6), (7) along with (10) and (15) is named *the Student-t Censored Measurement Error Model* – hereafter t-MEC model.

3.1. The likelihood function

In this section we present the likelihood function, which will be used in the model selection computations to compare fitted models.

First, let us partition $\mathbf{Z}_i$ into the observed and censored components, namely, $\mathbf{Z}_i = \text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c)$, where $\mathbf{Z}_i^o : p_o \times 1$ corresponds to the former case, $\mathbf{Z}_i^c : p_c \times 1$ corresponds to the latter and $\text{vec}(\cdot)$ denotes the function which stacks vectors or matrices of the same number of columns. Accordingly, let us consider $\mathbf{V}_i = \text{vec}(\mathbf{V}_i^o, \mathbf{V}_i^c)$ and, recalling that $\mathbf{Z}_i \sim t_p(\boldsymbol{\mu}_z, \boldsymbol{\Sigma}_z, \nu)$, see (13), $\boldsymbol{\mu}_z = \text{vec}(\boldsymbol{\mu}_z^o, \boldsymbol{\mu}_z^c)$ and $\boldsymbol{\Sigma}_z = \begin{pmatrix} \boldsymbol{\Sigma}_z^{oo} & \boldsymbol{\Sigma}_z^{oc} \\ \boldsymbol{\Sigma}_z^{co} & \boldsymbol{\Sigma}_z^{cc} \end{pmatrix}$. $\boldsymbol{\kappa}_i^c$ is the vector with the corresponding censoring levels for $\mathbf{Z}_i^c$. By Proposition 1, we have

$$\mathbf{Z}_i^o \sim t_{p_o}(\boldsymbol{\mu}_z^o, \boldsymbol{\Sigma}_z^{oo}, \nu) \quad \text{and} \quad \mathbf{Z}_i^c | \mathbf{Z}_i^o = \mathbf{z}_i^o \sim t_{p_c}(\boldsymbol{\mu}_z^{co}, \mathbf{S}_z^{co}, \nu + p_o), \quad (16)$$

where

$$\boldsymbol{\mu}_z^{co} = \boldsymbol{\mu}_z^c + \boldsymbol{\Sigma}_z^{co}(\boldsymbol{\Sigma}_z^{oo})^{-1}(\mathbf{z}_i^o - \boldsymbol{\mu}_z^o), \quad (17)$$

$$\mathbf{S}_z^{co} = \left(\frac{\nu + d_{\boldsymbol{\Sigma}_z^{oo}}(\mathbf{z}_i^o, \boldsymbol{\mu}_z^o)}{\nu + p^o} \right) \boldsymbol{\Sigma}_z^{cc.o}, \quad (18)$$

$$\boldsymbol{\Sigma}_z^{cc.o} = \boldsymbol{\Sigma}_z^{cc} - \boldsymbol{\Sigma}_z^{co} \boldsymbol{\Sigma}_z^{oo-1} \boldsymbol{\Sigma}_z^{oc}. \quad (19)$$

The observed sample for the experimental unit i is $\{\mathbf{z}_i^o, \boldsymbol{\kappa}_i^c\}$. The associated likelihood is

$$L_i(\boldsymbol{\theta}) = P(\mathbf{V}_i^c = \boldsymbol{\kappa}_i^c | \mathbf{Z}_i^o = \mathbf{z}_i^o) f(\mathbf{z}_i^o),$$

where $f(\cdot)$ is the marginal density of $\mathbf{Z}_i^o$. But $\mathbf{V}_i^c = \boldsymbol{\kappa}_i^c$ if and only if $\mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c$. By (16), we obtain

$$L_i(\boldsymbol{\theta}) = T_{p_c}(\boldsymbol{\kappa}_i^c | \boldsymbol{\mu}_z^{co}, \mathbf{S}_z^{co}, \nu + p_o) t_{p_o}(\mathbf{z}_i^o | \boldsymbol{\mu}_z^o, \boldsymbol{\Sigma}_z^{oo}, \nu).$$

The log-likelihood associated with the whole sample is

$$\ell(\boldsymbol{\theta}) = \sum_{i=1}^n \log L_i(\boldsymbol{\theta}). \quad (20)$$

4. The ECM algorithm

In this section, we describe how the t-MEC model can be fitted by using the ECM algorithm (Meng & Rubin, 1993). This algorithm considers a simple modification of the traditional EM algorithm initially proposed by Dempster *et al.* (1977) and is an efficient tool to obtain the maximum likelihood estimates under a missing data framework.

The t-MEC model can be formulated in a flexible hierarchical representation that is useful for theoretical derivations. It is easily obtained through Equations (8), (11) and (12) and is given by

$$\mathbf{Z}_i \mid x_i, U_i = u_i \stackrel{\text{ind}}{\sim} N_p(\mathbf{a} + \mathbf{b}x_i, u_i^{-1}\mathbf{\Omega}), \quad (21)$$

$$x_i \mid U_i = u_i \stackrel{\text{ind}}{\sim} N(\mu_x, u_i^{-1}\sigma_x^2), \quad (22)$$

$$U_i \stackrel{\text{iid}}{\sim} \text{Gamma}(\nu/2, \nu/2), \quad i = 1, \dots, n. \quad (23)$$

Following the suggestions of Lange *et al.* (1989) and Lucas (1997), who pointed out difficulties in estimating ν due to problems of unbounded and local maxima in the likelihood function, we consider the value of ν to be known.

Now, we enunciate two important results that will be useful in the E step of the EM algorithm.

Proposition 4. *Consider the hierarchical representation of the t-MEC model given in (21)–(23). Then,*

$$x_i \mid U_i = u_i, \mathbf{Z}_i = \mathbf{z}_i \sim N\left(\frac{\mu_x + \sigma_x^2 \mathbf{b}' \mathbf{\Omega}^{-1} (\mathbf{z}_i - \mathbf{a})}{1 + \sigma_x^2 \mathbf{b}' \mathbf{\Omega}^{-1} \mathbf{b}}, \frac{\sigma_x^2}{u_i (1 + \sigma_x^2 \mathbf{b}' \mathbf{\Omega}^{-1} \mathbf{b})}\right).$$

The proof follows directly from the relation $f(x_i \mid u_i, \mathbf{z}_i) \propto f(\mathbf{z}_i \mid x_i, u_i) f(x_i \mid u_i)$, where $f(\cdot)$ denotes a generic pdf.

Proposition 5. *For the t-MEC model,*

$$E(U_i \mid \mathbf{Z}_i = \mathbf{z}_i) = \frac{p + \nu}{d_{\mathbf{\Sigma}_z}(\mathbf{z}_i, \boldsymbol{\mu}_z) + \nu}.$$

To prove this result, recall that $\mathbf{Z}_i \sim t_p(\boldsymbol{\mu}_z, \mathbf{\Sigma}_z, \nu)$, which implies $\mathbf{Z}_i \mid U_i = u_i \sim N_p(\boldsymbol{\mu}_z, u_i^{-1} \mathbf{\Sigma}_z)$ and $U_i \sim \text{Gamma}(\nu/2, \nu/2)$ – see (2). Using the relation $f(u_i \mid \mathbf{z}_i) \propto f(\mathbf{z}_i \mid u_i) f(u_i)$, we can prove that $U_i \mid \mathbf{Z}_i = \mathbf{z}_i \sim \text{Gamma}\left(\frac{p+\nu}{2}, \frac{1}{2}(d_{\mathbf{\Sigma}_z}(\mathbf{z}_i, \boldsymbol{\mu}_z) + \nu)\right)$, and the result follows.

4.1. The E Step

Let $\mathbf{Z} = (\mathbf{Z}_1^\top, \dots, \mathbf{Z}_n^\top)^\top$, $\mathbf{x} = (x_1, \dots, x_n)^\top$ and $\mathbf{u} = (u_1, \dots, u_n)^\top$. Let $\boldsymbol{\theta}$ be the vector with all the parameters in the model. Apart from constants which do not depend on $\boldsymbol{\theta}$, the complete log-likelihood associated with the complete data $\mathbf{Z}_c = \{\mathbf{Z}, \mathbf{x}, \mathbf{u}\}$ is given by

$$\begin{aligned} \ell_c(\boldsymbol{\theta} \mid \mathbf{Z}_c) = & -\frac{n}{2} \sum_{i=1}^p \log \omega_j^2 - \frac{1}{2} \sum_{i=1}^n u_i (\mathbf{Z}_i - \mathbf{a} - \mathbf{b}x_i)^\top \mathbf{\Omega}^{-1} (\mathbf{Z}_i - \mathbf{a} - \mathbf{b}x_i) \\ & - \frac{n}{2} \log \sigma_x^2 - \frac{1}{2\sigma_x^2} \sum_{i=1}^n u_i (x_i - \mu_x)^2. \end{aligned}$$

Suppose that at the k th stage of the algorithm we obtain an estimate $\hat{\boldsymbol{\theta}}^{(k)}$ of $\boldsymbol{\theta}$. The E step consists of the computation of the conditional expectation

$$Q(\boldsymbol{\theta}|\hat{\boldsymbol{\theta}}^{(k)}) = E_{\hat{\boldsymbol{\theta}}^{(k)}} [\ell_c(\boldsymbol{\theta}|\mathbf{Z}_c)|\mathbf{V}],$$

where $E_{\hat{\boldsymbol{\theta}}^{(k)}}$ means that the expectation is being affected using $\hat{\boldsymbol{\theta}}^{(k)}$ as the true parameter value and $\mathbf{V} = (\mathbf{V}_1^\top, \dots, \mathbf{V}_n^\top)^\top$. The M step consists of maximizing $Q(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ in $\boldsymbol{\theta}$. To do so, first observe that the function $Q(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ can be decomposed into

$$Q(\boldsymbol{\theta}|\hat{\boldsymbol{\theta}}^{(k)}) = Q_1(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\omega}|\hat{\boldsymbol{\theta}}^{(k)}) + Q_2(\mu_x, \sigma_x^2|\hat{\boldsymbol{\theta}}^{(k)}), \quad (24)$$

where $\boldsymbol{\omega} = (\omega_1^2, \dots, \omega_p^2)^\top$,

$$Q_1(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\omega}|\hat{\boldsymbol{\theta}}^{(k)}) = E_{\hat{\boldsymbol{\theta}}^{(k)}} \left[-\frac{n}{2} \sum_{j=1}^p \log \omega_j^2 - \frac{1}{2} \sum_{i=1}^n u_i (\mathbf{Z}_i - \mathbf{a} - \mathbf{b}x_i)^\top \boldsymbol{\Omega}^{-1} (\mathbf{Z}_i - \mathbf{a} - \mathbf{b}x_i) | \mathbf{V} \right] \quad \text{and} \quad (25)$$

$$Q_2(\mu_x, \sigma_x^2|\hat{\boldsymbol{\theta}}^{(k)}) = E_{\hat{\boldsymbol{\theta}}^{(k)}} \left[-\frac{n}{2} \log \sigma_x^2 - \frac{1}{2\sigma_x^2} \sum_{i=1}^n u_i (x_i - \mu_x)^2 | \mathbf{V} \right].$$

Given this decomposition, we can reduce the problem to the maximization of two independent functions, searching for critical points of $Q_1(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ and $Q_2(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ separately. Expanding the expressions of $Q_1(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ and $Q_2(\cdot|\hat{\boldsymbol{\theta}}^{(k)})$ and taking expectations, it follows that

$$\begin{aligned} Q_1(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\omega}|\hat{\boldsymbol{\theta}}^{(k)}) &= -\frac{n}{2} \sum_{j=1}^p \log \omega_j^2 - \frac{1}{2} \left\{ \sum_{i=1}^n \left(\text{tr}\{\boldsymbol{\Omega}^{-1} \widehat{u\mathbf{z}_i^2}\} - 2\mathbf{a}^\top \boldsymbol{\Omega}^{-1} \widehat{u\mathbf{z}_i} \right. \right. \\ &\quad \left. \left. - 2\widehat{ux\mathbf{z}_i} \boldsymbol{\Omega}^{-1} \mathbf{b} + \mathbf{a}^\top \boldsymbol{\Omega}^{-1} \mathbf{a} \widehat{u_i} + 2\mathbf{a}^\top \boldsymbol{\Omega}^{-1} \mathbf{b} \widehat{ux_i} + \mathbf{b}^\top \boldsymbol{\Omega}^{-1} \mathbf{b} \widehat{ux_i^2} \right) \right\}, \\ Q_2(\mu_x, \sigma_x^2|\hat{\boldsymbol{\theta}}^{(k)}) &= -\frac{n}{2} \log \sigma_x^2 - \frac{1}{2\sigma_x^2} \left\{ \sum_{i=1}^n \left(\widehat{ux_i^2} - 2\mu_x \widehat{ux_i} + \mu_x^2 \widehat{u_i} \right) \right\}, \end{aligned}$$

where $\text{tr}(\cdot)$ denotes the trace of a matrix,

$$\begin{aligned} \widehat{u\mathbf{z}_i^2} &= E[U_i \mathbf{Z}_i \mathbf{Z}_i^\top | \mathbf{V}_i], \quad \widehat{u\mathbf{z}_i} = E[U_i \mathbf{Z}_i | \mathbf{V}_i], \\ \widehat{u_i} &= E[U_i | \mathbf{V}_i], \quad \widehat{ux\mathbf{z}_i} = E[U_i x_i \mathbf{Z}_i^\top | \mathbf{V}_i], \\ \widehat{ux_i} &= E[U_i x_i | \mathbf{V}_i], \quad \widehat{ux_i^2} = E[U_i x_i^2 | \mathbf{V}_i], \end{aligned}$$

and we have omitted the subscript $\hat{\boldsymbol{\theta}}^{(k)}$ to simplify the notation. To obtain expressions for these expectations, we will use a result from probability theory called *the tower property of conditional expectation*: if $\mathbf{X}$ and $\mathbf{Y}$ are arbitrary random vectors and $f(\cdot)$ is a measurable function, then $E[E(\mathbf{X}|\mathbf{Y})|f(\mathbf{Y})] = E[\mathbf{X}|f(\mathbf{Y})]$. For a proof, see Ash (2000, Theorem 5.5.10). Now, observe that, by (15), $\mathbf{V}_i$ is a function of $\mathbf{Z}_i$. Then, by this property, we can write

$$\widehat{u\mathbf{z}_i^2} = E\{E[U_i \mathbf{Z}_i \mathbf{Z}_i^\top | \mathbf{Z}_i] | \mathbf{V}_i\}, \quad \widehat{u\mathbf{z}_i} = E\{E[U_i \mathbf{Z}_i | \mathbf{Z}_i] | \mathbf{V}_i\}, \quad \text{and} \quad \widehat{u_i} = E\{E[U_i | \mathbf{Z}_i] | \mathbf{V}_i\}. \quad (26)$$

Proposition 5 gives the conditional expectation $E[U_i|\mathbf{Z}_i]$ and, from this result and formulas (26) we obtain the following expressions for $\widehat{u}_i$, $\widehat{u\mathbf{Z}_i}$ and $\widehat{u\mathbf{Z}_i^2}$ (as we will see soon, all expectations involved in the E step are written as functions of these), considering three different cases:

- (i) Individual i does not have censored components. In this case, $\mathbf{V}_i = \mathbf{Z}_i$ – see Equation (15). Thus,

$$\widehat{u}_i = \frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu}, \quad \widehat{u\mathbf{Z}_i} = \frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \mathbf{Z}_i, \quad \text{and} \quad \widehat{u\mathbf{Z}_i^2} = \frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \mathbf{Z}_i \mathbf{Z}_i^\top.$$

- (ii) Individual i has only censored components. By Equation (15), this fact occurs if and only if $\mathbf{Z}_i \leq \boldsymbol{\kappa}_i$, where $\boldsymbol{\kappa}_i$ is the vector with the censoring levels for individual i . Thus,

$$\widehat{u}_i = E\{E[U_i|\mathbf{Z}_i]|\mathbf{Z}_i \leq \boldsymbol{\kappa}_i\} = E\left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \middle| \mathbf{Z}_i \leq \boldsymbol{\kappa}_i\right].$$

By (13) and the definition of a truncated Student-t distribution, we have that $\mathbf{Z}_i | (\mathbf{Z}_i \leq \boldsymbol{\kappa}_i) \sim \text{Tt}_p(\boldsymbol{\mu}_z, \Sigma_z, \nu; \mathbb{D}_i)$, where $\mathbb{D}_i$ is like in (3) with $\mathbf{d} = \boldsymbol{\kappa}_i$. Using $r = 1$ and $k = 0$ in Proposition 2, we get

$$\widehat{u}_i = \frac{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z^*, \nu + 2r)}{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z, \nu)},$$

where $\Sigma_z^* = (\nu/(\nu + 2))\Sigma_z$. Using $r = 1$ and $k = 1$ in Proposition 2, we obtain

$$\widehat{u\mathbf{Z}_i} = \frac{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z^*, \nu + 2r)}{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z, \nu)} E[\mathbf{Y}_i],$$

where $\mathbf{Y}_i \sim \text{Tt}_p(\boldsymbol{\mu}_z, \Sigma_z^*, \nu + 2; \mathbb{D}_i)$. Finally, $r = 1$ and $k = 2$ in Proposition 2 imply

$$\widehat{u\mathbf{Z}_i^2} = \frac{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z^*, \nu + 2r)}{\text{T}_p(\boldsymbol{\kappa}_i | \boldsymbol{\mu}_z, \Sigma_z, \nu)} E[\mathbf{Y}_i \mathbf{Y}_i^\top].$$

- (iii) Individual i has censored and uncensored components. As we commented before in Section 3.1, in this case, we decompose the vector $\mathbf{V}_i$ into two subvectors, $\mathbf{Z}_i^o$ and $\mathbf{Z}_i^c$, corresponding to the uncensored observations and the censoring levels, respectively. Accordingly, we partition the vector $\mathbf{Z}_i$ as $\mathbf{Z}_i = \text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c)$. The components are censored if and only if $\mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c$. Thus,

$$\widehat{u}_i = E\{E[U_i|\mathbf{Z}_i]|\mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c\} = E\left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \middle| \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c\right].$$

In this case, we have that $\mathbf{Z}_i | (\mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c) \sim \text{Tt}_p(\boldsymbol{\mu}_z, \Sigma_z, \nu; \mathbb{D}_i^c)$, with

$$\mathbb{D}_i^c = \{(x_1, \dots, x_p) \in \mathbb{R}^p; \ x_i \leq \kappa_i^c, \ i \in \mathcal{C}\}, \quad (27)$$

where $\mathcal{C}$ is the set of indices for the censored components – consequently, we make $d_i = +\infty$ for $i \notin \mathcal{C}$ in (3). Thus, $\widehat{u}_i$ can be calculated using Proposition 3, with $\mathbf{Z}_i^o$ and $\mathbf{Z}_i^c$ playing the role of $\mathbf{Z}_1$ and $\mathbf{Z}_2$, respectively, taking $r = 1$ and $k = 0$. Then, we get

$$\widehat{u}_i = \frac{p^o + \nu}{\nu + d_{\Sigma_z^{oo}}(\mathbf{z}_i^o, \boldsymbol{\mu}_z^o)} \frac{\text{T}_{p^c}(\boldsymbol{\kappa}_i^c | \boldsymbol{\mu}_z^{co}, \widetilde{\mathbf{S}}_z^{co}, \nu + p_o + 2)}{\text{T}_{p^c}(\boldsymbol{\kappa}_i^c | \boldsymbol{\mu}_z^{co}, \mathbf{S}_z^{co}, \nu + p_o)},$$

where p^o and p^c are the dimensions of the vectors $\mathbf{Z}_i^o$ e $\mathbf{Z}_i^c$, respectively, $\nu + p^o + 2 > 0$,

$$\tilde{\mathbf{S}}_z^{co} = \frac{\nu + d_{\Sigma_z^{oo}}(\mathbf{z}_i^o, \boldsymbol{\mu}_z^o)}{\nu + p^o + 2} \Sigma_z^{cc.o},$$

$\boldsymbol{\mu}_z^o$, $\mathbf{S}_z^{co}$ and $\Sigma_z^{cc.o}$ are given in (17), (18) and (19), respectively. Regarding $\widehat{u\mathbf{z}}_i$, we have that

$$\begin{aligned} \widehat{u\mathbf{z}}_i &= \mathbb{E} \left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c) | \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c \right] \\ &= \text{vec} \left(\mathbb{E} \left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \mathbf{z}_i^o | \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c \right], \mathbb{E} \left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \mathbf{Z}_i^c | \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c \right] \right) \\ &= \text{vec}(\widehat{u_i \mathbf{z}_i^o}, \mathbb{E}[\mathbf{Y}_i]), \end{aligned}$$

where

$$\mathbf{Y}_i \sim \text{Tt}_{p^c}(\boldsymbol{\mu}_z^{co}, \tilde{\mathbf{S}}_z^{co}, \nu + p^o + 2; \mathbb{D}_i^c). \quad (28)$$

Finally, to compute $\widehat{u\mathbf{z}}_i^2$, observe that

$$\begin{aligned} \widehat{u\mathbf{z}}_i^2 &= \mathbb{E} \left[\frac{p + \nu}{d_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu} \begin{pmatrix} \mathbf{Z}_i^o \mathbf{Z}_i^{o\top} & \mathbf{Z}_i^o \mathbf{Z}_i^{c\top} \\ \mathbf{Z}_i^c \mathbf{Z}_i^{o\top} & \mathbf{Z}_i^c \mathbf{Z}_i^{c\top} \end{pmatrix} | \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c \right] \\ &= \begin{pmatrix} \widehat{u_i \mathbf{z}_i^o} \mathbf{z}_i^{o\top} & \widehat{u_i \mathbf{z}_i^o} \mathbb{E}[\mathbf{Y}_i]^\top \\ \widehat{u_i} \mathbb{E}[\mathbf{Y}_i] \mathbf{z}_i^{o\top} & \widehat{u_i} \mathbb{E}[\mathbf{Y}_i \mathbf{Y}_i^\top] \end{pmatrix}, \end{aligned}$$

where $\mathbf{Y}_i$ is as in (28).

Regarding the remaining expectations, we have

$$\begin{aligned} \mathbb{E}[x_i U_i | \mathbf{V}_i = \mathbf{v}_i] &= \iint x_i u_i \pi(x_i, u_i | \mathbf{v}_i) dx_i du_i \\ &= \int x_i \pi(x_i | u_i, \mathbf{v}_i) dx_i \int u_i \pi(u_i | \mathbf{v}_i) du_i \\ &= \mathbb{E}[x_i | U_i = u_i, \mathbf{V}_i = \mathbf{v}_i] \mathbb{E}[U_i | \mathbf{V}_i = \mathbf{v}_i]. \end{aligned} \quad (29)$$

By the tower property, we have

$$\mathbb{E}[x_i | U_i, \mathbf{V}_i] = \mathbb{E}[\mathbb{E}(x_i | U_i, \mathbf{Z}_i) | U_i, \mathbf{V}_i].$$

Consequently,

$$\begin{aligned} \widehat{u\mathbf{x}}_i &= \mathbb{E}[x_i U_i | \mathbf{V}_i] = \mathbb{E} \left[\frac{\mu_x + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} (\mathbf{Z}_i - \mathbf{a})}{1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b}} | U_i, \mathbf{V}_i \right] \mathbb{E}[U_i | \mathbf{V}_i] \\ &= \frac{\mu_x \mathbb{E}[U_i | \mathbf{V}_i] + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbb{E}[\mathbf{Z}_i | U_i, \mathbf{V}_i] \mathbb{E}[U_i | \mathbf{V}_i] - \mathbf{a} \mathbb{E}[U_i | \mathbf{V}_i]}{1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b}} \\ &= \frac{\mu_x \widehat{u_i} + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \widehat{u\mathbf{z}}_i - \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{a} \widehat{u_i}}{1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b}} \\ &= \mu_x \widehat{u_i} + \boldsymbol{\varphi}(\widehat{u\mathbf{z}}_i - \boldsymbol{\mu}_z \widehat{u_i}), \end{aligned} \quad (30)$$

where

$$\boldsymbol{\varphi} = \frac{\sigma_x^2 \mathbf{b}^\top \boldsymbol{\Omega}^{-1}}{1 + \sigma_x^2 \mathbf{b}^\top \boldsymbol{\Omega}^{-1} \mathbf{b}} \quad (31)$$

and the equality in (30) is obtained by proving that $E[\mathbf{Z}_i|U_i, \mathbf{V}_i]E[U_i|\mathbf{V}_i] = E[U_i\mathbf{Z}_i|\mathbf{V}_i] \equiv \widehat{u\mathbf{Z}_i}$, which can be done following the same paths that led to (29), replacing x_i with $\mathbf{Z}_i$.

In a similar fashion, we get

$$\begin{aligned} \widehat{ux_i^2} &= \Lambda + \mu_x^2 \widehat{u_i} + 2\boldsymbol{\varphi} [\widehat{u\mathbf{Z}_i} - \boldsymbol{\mu}_z \widehat{u_i}] + \boldsymbol{\varphi} \left[\widehat{u\mathbf{Z}_i^2} - \widehat{u\mathbf{Z}_i} \boldsymbol{\mu}_z^\top - \boldsymbol{\mu}_z \widehat{u\mathbf{Z}_i}^\top + \boldsymbol{\mu}_z \boldsymbol{\mu}_z^\top \widehat{u_i} \right] \boldsymbol{\varphi}^\top, \text{ and} \\ \widehat{ux\mathbf{Z}_i} &= \mu_x \widehat{u\mathbf{Z}_i} + \boldsymbol{\varphi} \left[\widehat{u\mathbf{Z}_i^2} - \boldsymbol{\mu}_z \widehat{u\mathbf{Z}_i} \right], \end{aligned}$$

with

$$\Lambda = \frac{\sigma_x^2}{1 + \sigma_x^2 \mathbf{b}^\top \boldsymbol{\Omega}^{-1} \mathbf{b}}. \quad (32)$$

4.2. The CM Step

Given the current estimate $\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}^{(k)}$ at the k th stage, the CM-step of the ECM algorithm (Meng & Rubin, 1993) consists of the *conditional maximization* of the Q function given in (24). More precisely, ECM replaces each M-step of the EM algorithm of Dempster *et al.* (1977) by a sequence of S conditional maximization steps, called CM-steps, each of which maximizes the Q function over $\boldsymbol{\theta}$ but with some vector function of $\boldsymbol{\theta}$, $(g_1(\boldsymbol{\theta}), \dots, g_S(\boldsymbol{\theta}))$ say, fixed at its previous value. In our case, for example, we first maximize conditionally the function $Q_1(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\omega} | \widehat{\boldsymbol{\theta}}^{(k)})$ in (25) over $\boldsymbol{\alpha}$ fixing the values $\boldsymbol{\beta} = \widehat{\boldsymbol{\beta}}^{(k)}$ and $\boldsymbol{\omega} = \widehat{\boldsymbol{\omega}}^{(k)}$. Then we maximize $Q_1(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\omega} | \widehat{\boldsymbol{\theta}}^{(k)})$ over $\boldsymbol{\beta}$ fixing the values $\boldsymbol{\alpha} = \widehat{\boldsymbol{\alpha}}^{(k+1)}$ and $\boldsymbol{\omega} = \widehat{\boldsymbol{\omega}}^{(k)}$ and so on. We get the following closed expressions:

$$\begin{aligned} \widehat{\boldsymbol{\alpha}}^{(k+1)} &= \bar{\mathbf{z}}_u^{(k)} - \bar{x}_u^{(k)} \widehat{\boldsymbol{\beta}}^{(k)}, \\ \widehat{\boldsymbol{\beta}}^{(k+1)} &= \frac{n \widehat{u}^{(k)} \sum_{i=1}^n \widehat{ux\mathbf{Z}_i}^{*(k)} - \sum_{i=1}^n \widehat{u\mathbf{Z}_i}^{*(k)} \sum_{i=1}^n \widehat{ux_i}^{(k)}}{n \widehat{u}^{(k)} \sum_{i=1}^n \widehat{ux_i^2}^{(k)} - \left(\sum_{i=1}^n \widehat{ux_i}^{(k)} \right)^2}, \\ \widehat{\omega_1^2}^{(k+1)} &= \frac{1}{n} \sum_{i=1}^n (\widehat{u\mathbf{Z}_i^2}^{(k)} - 2\widehat{ux\mathbf{Z}_i}^{(k)} + \widehat{ux_i^2}^{(k)}), \\ \widehat{\omega_{j+1}^2}^{(k+1)} &= \frac{1}{n} \sum_{i=1}^n \left(\widehat{u\mathbf{Z}_i^2}^{(k)}_{(j+1)(j+1)} + \widehat{u_i}^{(k)} \widehat{\alpha_j^2}^{(k+1)} + \widehat{ux_i^2}^{(k)} \widehat{\beta_j^2}^{(k+1)} + 2\widehat{ux_i}^{(k)} \widehat{\alpha_j}^{(k+1)} \widehat{\beta_j}^{(k+1)} \right. \\ &\quad \left. - 2\widehat{ux\mathbf{Z}_i}^{(k)}_{(j+1)j} \widehat{\beta_j}^{(k+1)} - 2\widehat{u\mathbf{Z}_i}^{(k)}_{(j+1)j} \widehat{\alpha_j}^{(k+1)} \right), \quad j = 1, \dots, r, \\ \widehat{\mu_x}^{(k+1)} &= \bar{x}_u^{(k)}, \\ \widehat{\sigma_x^2}^{(k+1)} &= \frac{1}{n} \sum_{i=1}^n (\widehat{ux_i^2}^{(k)} - 2\widehat{ux_i}^{(k)} \widehat{\mu_x}^{(k+1)} + \widehat{ux_i}^{(k)} \widehat{\mu_x^2}^{(k+1)}), \end{aligned}$$

where $\bar{\mathbf{z}}_u^{(k)} = \frac{\sum_{i=1}^n \widehat{u\mathbf{z}}_i^{\star(k)}}{\sum_{i=1}^n \widehat{u}_i^{(k)}}$, $\bar{x}_u^{(k)} = \frac{\sum_{i=1}^n \widehat{ux}_i^{(k)}}{\sum_{i=1}^n \widehat{u}_i^{(k)}}$ and $\bar{u}^{(k)} = \frac{1}{n} \sum_{i=1}^n \widehat{u}_i^{(k)}$, with $\widehat{u\mathbf{z}}_i^{\star(k)} = (\widehat{u\mathbf{z}}_{i2}, \dots, \widehat{u\mathbf{z}}_{ip})^\top$ and $\widehat{ux\mathbf{z}}_i^{\star(k)} = (\widehat{ux\mathbf{z}}_{i2}, \dots, \widehat{ux\mathbf{z}}_{ip})^\top$.

4.3. Imputation of censored components

Let $\mathbf{Z}_i^{(c)}$ be the true (partially or completely unobserved) response vector for the censored components of the i th unit. We define a predictor for $\mathbf{Z}_i^{(c)}$ as

$$\tilde{\mathbf{Z}}_i^{(c)} = \mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i = \mathbf{v}_i].$$

We have the following particular cases:

1. If unit i has only censored components then, if we make $r = 0$ and $k = 1$ in Proposition 2, we get

$$\tilde{\mathbf{Z}}_i^{(c)} = \mathbb{E}[\mathbf{Y}_i], \quad \text{with } \mathbf{Y}_i \sim \text{Tt}_p(\hat{\boldsymbol{\mu}}_z, \hat{\boldsymbol{\Sigma}}_z, \nu; \mathbb{D}_i),$$

$\hat{\boldsymbol{\mu}}_z$ and $\hat{\boldsymbol{\Sigma}}_z$ are the EM estimates of $\boldsymbol{\mu}_z$ and $\boldsymbol{\Sigma}_z$, respectively, and $\mathbb{D}_i$ is as in (3) with $\mathbf{d} = \boldsymbol{\kappa}_i$, where $\boldsymbol{\kappa}_i$ is the vector with censoring levels for unit i .

2. Unit i has uncensored and censored components. In this case, we partition the vector $\mathbf{Z}_i$ as $\mathbf{Z}_i = \text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c)$. Components are censored if and only if $\mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c$, such that

$$\tilde{\mathbf{Z}}_i^{(c)} = \mathbb{E}[\text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c) | \mathbf{Z}_i^o = \mathbf{z}_i^o, \mathbf{Z}_i^c \leq \boldsymbol{\kappa}_i^c] = \text{vec}(\mathbf{z}_i^o, \hat{\mathbf{y}}_i^c),$$

where, by Proposition 3 with $r = 0$ and $k = 1$,

$$\hat{\mathbf{y}}_i^c = \mathbb{E}[\mathbf{Y}_i], \quad \text{with } \mathbf{Y}_i \sim \text{Tt}_{p^c}(\boldsymbol{\mu}_z^{co}, \mathbf{S}_z^{co}, \nu + p^o; \mathbb{D}_i^c), \quad (33)$$

where $\boldsymbol{\mu}_z^{co}$ and $\mathbf{S}_z^{co}$ are given in (17) and (18), respectively, and $\mathbb{D}_i^c$ is given in (27).

4.4. Estimation of x_i

Following Lin & Lee (2006), Ho *et al.* (2012) and recently Castro *et al.* (2015), we consider the conditional mean to estimate the unobserved latent covariate. Using the tower property and Proposition 4, we have that an estimator for x_i can be obtained through

$$\begin{aligned} \hat{x}_i &= \mathbb{E}[x_i | \mathbf{V}_i] = \mathbb{E}[\mathbb{E}(x_i | U_i, \mathbf{Z}_i) | \mathbf{V}_i] \\ &= \mathbb{E} \left[\frac{\mu_x + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} (\mathbf{Z}_i - \mathbf{a})}{1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b}} \middle| \mathbf{V}_i \right] \\ &= \mu_x + \boldsymbol{\varphi}(\hat{\mathbf{Z}}_i - \mathbf{a} - \mathbf{b}\mu_x), \end{aligned} \quad (34)$$

where $\boldsymbol{\varphi}$ is given in (31) and $\hat{\mathbf{Z}}_i = \mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i]$. Observe that, if individual i does not have censored components, then $\hat{\mathbf{Z}}_i$ is the first moment of a $\text{t}_p(\boldsymbol{\mu}_z, \boldsymbol{\Sigma}_z, \nu)$ distribution. If all its components are censored, then $\mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i] = \mathbb{E}[\mathbf{Z}_i | \mathbf{Z}_i \leq \boldsymbol{\kappa}_i]$, which can be computed using Proposition 2 with $r = 0$ and $k = 1$. Finally, if it has censored and uncensored components, then $\mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i] = \text{vec}(\mathbf{Z}_i^o, \hat{\mathbf{y}}_i^c)$, see (33). The parameter values in (34) must be replaced with the respective EM estimates.

Moreover, the conditional covariance matrix of x_i given $\mathbf{V}_i$ is

$$\text{Var}[x_i | \mathbf{V}_i] = \mathbb{E}[x_i^2 | \mathbf{V}_i] - (\mathbb{E}[x_i | \mathbf{V}_i])^2.$$

By the tower property and Proposition 4, we have

$$\begin{aligned} \mathbb{E}[x_i^2 | \mathbf{V}_i] &= \mathbb{E}\{\mathbb{E}[\mathbb{E}(x_i^2 | U_i, \mathbf{Z}_i) | \mathbf{Z}_i] | \mathbf{V}_i\} \\ &= \mathbb{E}\left\{\mathbb{E}\left[\frac{\sigma_x^2}{U_i(1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b})} | \mathbf{Z}_i\right] | \mathbf{V}_i\right\} + \mathbb{E}\left[\left(\frac{\mu_x + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} (\mathbf{Z}_i - \mathbf{a})}{1 + \sigma_x^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b}}\right)^2 | \mathbf{V}_i\right]. \end{aligned}$$

It is easy to show that $\mathbb{E}[U_i^{-1} | \mathbf{Z}_i] = (\mathbf{d}_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu) / (p + \nu - 2)$ – recall that $U_i | \mathbf{Z}_i = \mathbf{z}_i \sim \text{Gamma}\left(\frac{p+\nu}{2}, \frac{1}{2}(\mathbf{d}_{\Sigma_z}(\mathbf{z}_i, \boldsymbol{\mu}_z) + \nu)\right)$, see the result after Proposition 5. After some lengthy algebra, we can prove that

$$\text{Var}[x_i | \mathbf{V}_i] = \Lambda \mathbb{E}\left[\frac{\mathbf{d}_{\Sigma_z}(\mathbf{Z}_i, \boldsymbol{\mu}_z) + \nu}{p + \nu - 2} | \mathbf{V}_i\right] + \Lambda^2 \mathbf{b}' \boldsymbol{\Omega}^{-1} \text{Var}[\mathbf{Z}_i | \mathbf{V}_i] \boldsymbol{\Omega}^{-1} \mathbf{b}, \quad (35)$$

where Λ is given in (32) and

$$\text{Var}[\mathbf{Z}_i | \mathbf{V}_i] = \mathbb{E}[\mathbf{Z}_i \mathbf{Z}_i^\top | \mathbf{V}_i] - \mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i] \mathbb{E}[\mathbf{Z}_i | \mathbf{V}_i]^\top.$$

If individual i has only uncensored components, then expression (35) can be computed using the moments of the $t_p(\boldsymbol{\mu}_z, \Sigma_z, \nu)$ distribution using Proposition 2: it is enough to make $d_1 = \dots = d_p = +\infty$, $r = 1$ and $k = 0$ to obtain the first expectation in (35) and $r = 0$, $k = 1$ ($k = 2$) to obtain the other one. If the components are all censored, we again use Proposition 2, but now considering the moments of a $\text{Tt}_p(\boldsymbol{\mu}, \Sigma, \nu; \mathbb{D}_i)$ distribution. Finally, if there are censored and uncensored components, then the expectations can be computed through Proposition 3, using the partition $\mathbf{Z}_i = \text{vec}(\mathbf{Z}_i^o, \mathbf{Z}_i^c)$. Besides this, the parameter values in (35) must be replaced with the respective EM estimates.

5. The observed information matrix

Under some general regularity conditions, we follow Lin (2010) to provide an information-based method to obtain the asymptotic covariance of ML estimates of the t-MEC model's parameters. As defined by Meilijson (1989), the empirical information matrix can be computed as

$$\mathbf{I}_e(\boldsymbol{\theta} | \mathbf{Z}) = \sum_{i=1}^n \mathbf{s}(\mathbf{Z}_i | \boldsymbol{\theta}) \mathbf{s}^\top(\mathbf{Z}_i | \boldsymbol{\theta}) - \frac{1}{n} \mathbf{S}(\mathbf{Z}_i | \boldsymbol{\theta}) \mathbf{S}^\top(\mathbf{Z}_i | \boldsymbol{\theta}),$$

where $\mathbf{S}(\mathbf{Z}_i | \boldsymbol{\theta}) = \sum_{i=1}^n \mathbf{s}(\mathbf{Z}_i | \boldsymbol{\theta})$ and $\mathbf{s}(\mathbf{Z}_i | \boldsymbol{\theta})$ is the empirical score function for the i th unit. According to Louis (1982) it is possible to relate the score function of the incomplete data log-likelihood with the conditional expectation of the complete data log-likelihood function. Therefore, the individual score can be determined as

$$\mathbf{s}(\mathbf{Z}_i | \boldsymbol{\theta}) = \frac{\partial \log f(\mathbf{Z}_i | \boldsymbol{\theta})}{\partial \boldsymbol{\theta}} = \mathbb{E}\left[\frac{\partial \ell_{ic}(\boldsymbol{\theta} | \mathbf{Z}_i^c)}{\partial \boldsymbol{\theta}} | \mathbf{V}_i, \mathbf{C}_i, \boldsymbol{\theta}\right], \quad (36)$$

where $\ell_{ic}(\boldsymbol{\theta} | \mathbf{Z}_i^c)$ is the complete data log-likelihood formed from the single observation $\mathbf{Z}_i$, $i = 1, \dots, n$. Using the EM estimates $\hat{\boldsymbol{\theta}}$, $\mathbf{S}(\mathbf{Z}_i | \hat{\boldsymbol{\theta}}) = 0$, and then (36) is given by

$$\mathbf{I}_e(\hat{\boldsymbol{\theta}} | \mathbf{Z}) = \sum_{i=1}^n \hat{\mathbf{s}}_i \hat{\mathbf{s}}_i^\top, \quad (37)$$

where $\widehat{\mathbf{s}}_i = (\widehat{\mathbf{s}}_{i,\boldsymbol{\alpha}}, \widehat{\mathbf{s}}_{i,\boldsymbol{\beta}}, \widehat{\mathbf{s}}_{i,\boldsymbol{\omega}}, \widehat{\mathbf{s}}_{i,\mu_x}, \widehat{\mathbf{s}}_{i,\sigma_x^2})^\top$ is a $3p$ -dimensional vector, with components given by

$$\begin{aligned}\widehat{\mathbf{s}}_{i,\boldsymbol{\alpha}} &= (\widehat{\mathbf{s}}_{i,\alpha_1}, \dots, \widehat{\mathbf{s}}_{i,\alpha_r})^\top = \mathbb{I}_{(p)} \widehat{\boldsymbol{\Omega}}^{-1} (\widehat{u\mathbf{z}}_i - \widehat{u}_i \widehat{\mathbf{a}} - \widehat{ux}_i \widehat{\mathbf{b}}), \\ \widehat{\mathbf{s}}_{i,\boldsymbol{\beta}} &= (\widehat{\mathbf{s}}_{i,\beta_1}, \dots, \widehat{\mathbf{s}}_{i,\beta_r})^\top = \mathbb{I}_{(p)} \widehat{\boldsymbol{\Omega}}^{-1} (\widehat{ux\mathbf{z}}_i - \widehat{ux}_i \widehat{\mathbf{a}} - \widehat{ux}_i^2 \widehat{\mathbf{b}}), \\ \widehat{\mathbf{s}}_{i,\boldsymbol{\omega}} &= (\widehat{\mathbf{s}}_{i,\omega_1^2}, \dots, \widehat{\mathbf{s}}_{i,\omega_p^2})^\top = -\frac{1}{2} \widehat{\boldsymbol{\Omega}}^{-1} \mathbf{1}_p + \frac{1}{2} \widehat{\boldsymbol{\Omega}}^{-2} \text{diag}(\widehat{a}_i), \\ \widehat{\mathbf{s}}_{i,\mu_x} &= \frac{1}{\widehat{\sigma_x^2}} (\widehat{ux}_i - \widehat{u}_i \widehat{\mu}_x), \\ \widehat{\mathbf{s}}_{i,\sigma_x^2} &= -\frac{1}{2\widehat{\sigma_x^2}} + \frac{1}{2\widehat{\sigma_x^4}} (\widehat{ux}_i^2 - 2\widehat{ux}_i \widehat{\mu}_x + \widehat{u}_i \widehat{\mu}_x^2),\end{aligned}$$

with $\mathbb{I}_{(p)} = [\mathbf{0}, \mathbf{I}_{p-1}]_{(p-1) \times p}$, $\mathbf{1}_p = (1, \dots, 1)_{p \times 1}^\top$ and $\widehat{a}_i = \widehat{u\mathbf{z}}_i^2 - 2\widehat{u\mathbf{z}}_i \widehat{\mathbf{a}}^\top - 2\widehat{ux\mathbf{z}}_i \widehat{\mathbf{b}}^\top + 2\widehat{ux}_i \widehat{\mathbf{a}} \widehat{\mathbf{b}}^\top + \widehat{u}_i \widehat{\mathbf{a}} \widehat{\mathbf{a}}^\top + \widehat{ux}_i^2 \widehat{\mathbf{b}} \widehat{\mathbf{b}}^\top$.

6. Simulation studies

In order to study the performance of our proposed method, we present three simulation studies. The first one shows the asymptotic behavior of the EM estimates for the proposed model. The second one investigates the consequences on parameter inference when the normality assumption is inappropriate. Finally, the third one is designed to investigate the effect of including the censoring component in the model.

6.1. Asymptotic properties

In this simulation study, we analyze the absolute bias (BIAS) and mean square error (MSE) of the regression coefficient estimates obtained from the t-MEC model for six different sample sizes n , namely 50, 100, 200, 300, 400 and 600. These measures are defined by

$$\text{BIAS}_k = \frac{1}{M} \sum_{j=1}^M |\widehat{\boldsymbol{\theta}}_k^{(j)} - \boldsymbol{\theta}_k| \quad \text{and} \quad \text{MSE}_k = \frac{1}{M} \sum_{j=1}^M \left(\widehat{\boldsymbol{\theta}}_k^{(j)} - \boldsymbol{\theta}_k \right)^2, \quad (38)$$

where $\widehat{\boldsymbol{\theta}}_k^{(j)}$ is the EM estimate of the parameter $\boldsymbol{\theta}_k$, $k = 1, \dots, 3p$, for the j -th sample. The key idea of this simulation is to provide empirical evidence about consistency of the EM estimators under the proposed t-MEC model. For each sample size, we generate $M = 100$ datasets with 10% censoring proportion. Using the ECM algorithm, the absolute bias and mean squared error for each parameter over the 100 datasets were computed. The parameter setup (see Section 3), is

$$\boldsymbol{\alpha} = (3, 2, 1, 2)^\top, \quad \boldsymbol{\beta} = (1.5, 1, 1.5, 1)^\top, \quad \mu_x = 4, \quad \sigma_x^2 = 2 \quad \text{and} \quad \boldsymbol{\Omega} = \text{diag}(0.5, 0.5, 0.5, 0.5, 0.5). \quad (39)$$

The degrees of freedom were fixed at the value $\nu = 5$.

The results are presented in Figure 1. From this figure we can see that the MSE tends to zero as the sample size increases. Similar results were obtained after the analysis of the absolute bias (BIAS) as can be seen from Figure 4 in the Appendix. As expected, the proposed ECM algorithm provides ML estimates with good asymptotic properties for the t-MEC model.

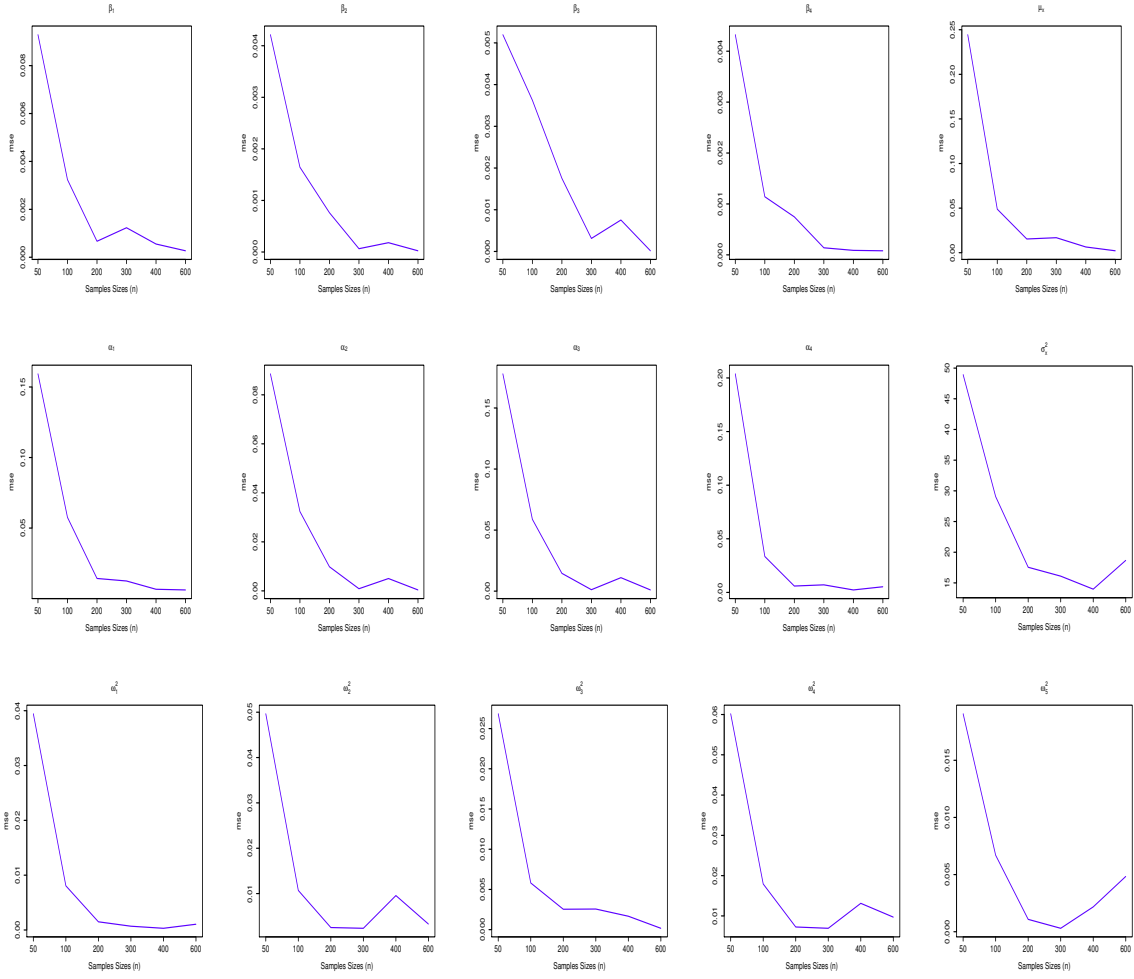


Figure 1: Simulation study 6.1. MSE of parameter estimates under the t -MEC model considering 10% censoring.

6.2. Parameter inference

In this study we investigate the consequences on parameter inference when the normality assumption is inappropriate, as well the ability of some model choice criteria (AIC and BIC) to select the correct model. In addition, we study the effect of different censoring proportions on the EM estimates. For this purpose, we consider a heavy-tail distribution for the random errors. In this context, we generate $M = 100$ datasets coming from a slash distribution with parameter $\nu = 1.5$ and censoring proportions 0%, 10%, 20% and 30%. The slash distribution arises when we change the distribution of U in (2) to $U \sim \text{Beta}(\nu, 1)$, with pdf $f(u|\nu) = \nu u^{\nu-1}$, $u \in (0, 1)$, and $\nu > 0$. See Wang & Genton (2006) for details. The parameter values are set as in the previous experimental study.

For each simulated dataset we fitted the t -MEC (with $\nu = 5$ degrees of freedom) and the N-MEC models. The model selection criteria AIC and BIC as well as the estimates of the model parameters were recorded at each simulation. Summary statistics such as the Monte Carlo mean estimate (MC mean), coverage probability (MC CP) and the approximate standard error obtained through the information-based method (IM SE), discussed in Section 5, for the

Table 1: Simulation study 6.2. Summary statistics based on 100 simulated samples from the slash distribution for different levels of censoring (0%, 10%, 20%, 30%).

Censoring	Fit		Simulated data									
			α_1	α_2	α_3	α_4	β_1	β_2	β_3	β_4	μ_x	σ_x^2
0%	Normal	MC Mean	3.096	2.457	1.274	2.238	1.482	0.933	1.466	0.988	3.463	212.602
		IM SE	0.332	0.243	0.343	0.302	0.038	0.029	0.045	0.038	4.783	0.002
		MC CP	100%	24%	97%	100%	100%	24%	100%	100%	100%	
	Student-t	MC Mean	2.712	1.853	0.964	1.894	1.542	1.022	1.521	1.022	4.586	9.299
		IM SE	0.375	0.245	0.331	0.258	0.066	0.044	0.060	0.046	0.404	0.022
		MC CP	100%	100%	100%	100%	100%	100%	100%	100%	79%	
10%	Normal	MC Mean	2.440	1.570	0.887	1.701	1.608	1.097	1.535	1.090	4.797	31.343
		IM SE	0.345	0.328	0.456	0.405	0.044	0.051	0.072	0.062	0.847	0.010
		MC CP	43%	61%	100%	61%	71%	68%	100%	100%	100%	
	Student-t	MC Mean	2.612	1.626	0.880	1.849	1.565	1.069	1.540	1.031	4.539	7.927
		IM SE	0.415	0.312	0.397	0.294	0.073	0.058	0.072	0.053	0.351	0.026
		MC CP	100%	100%	100%	100%	100%	100%	100%	100%	75%	
20%	Normal	MC Mean	2.435	1.539	0.905	1.600	1.610	1.103	1.533	1.103	4.754	32.254
		IM SE	0.415	0.417	0.532	0.492	0.051	0.063	0.081	0.076	0.832	0.010
		MC CP	55%	61%	100%	61%	83%	85%	100%	98%	100%	
	Student-t	MC Mean	2.506	1.605	0.816	1.783	1.580	1.073	1.549	1.041	4.576	7.599
		IM SE	0.503	0.399	0.494	0.357	0.084	0.070	0.084	0.061	0.350	0.027
		MC CP	99%	100%	100%	100%	100%	100%	100%	100%	64%	
30%	Normal	MC Mean	2.511	1.490	0.999	1.532	1.608	1.112	1.528	1.114	4.703	32.890
		IM SE	0.521	0.543	0.662	0.606	0.061	0.077	0.096	0.093	0.841	0.011
		MC CP	61%	61%	100%	61%	87%	89%	100%	88%	100%	
	Student-t	MC Mean	2.522	1.578	0.864	1.840	1.580	1.080	1.545	1.034	4.554	7.211
		IM SE	0.622	0.506	0.641	0.433	0.096	0.082	0.099	0.069	0.352	0.029
		MC CP	100%	100%	100%	100%	100%	99%	100%	100%	84%	

parameter estimates are presented in Table 1.

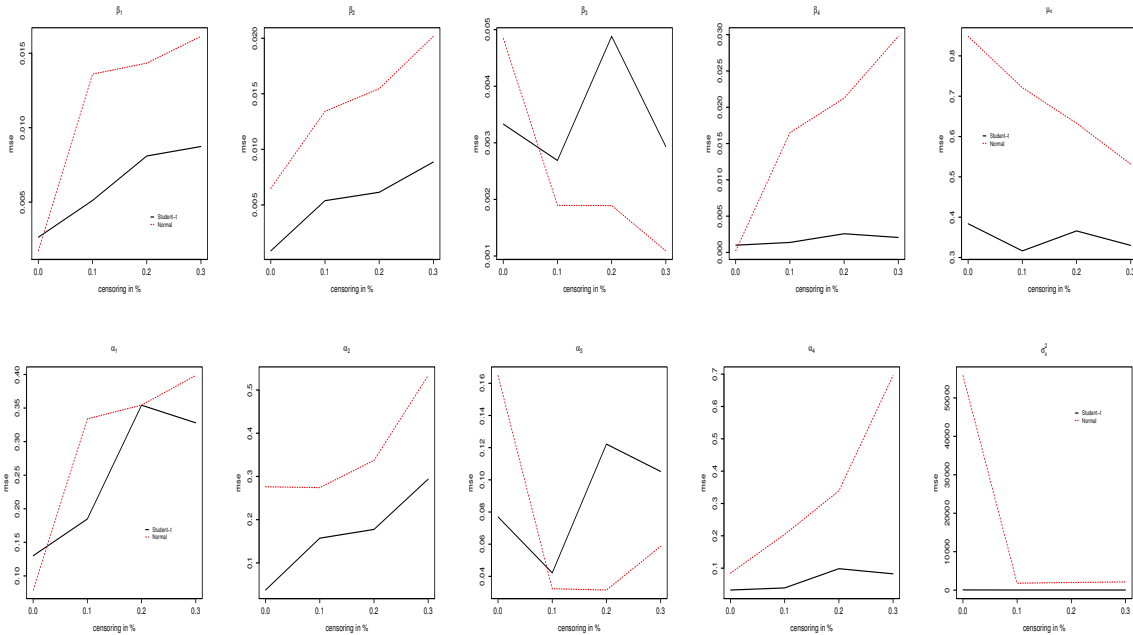


Figure 2: Simulation study 6.2. MSE of $\beta, \alpha, \mu_x, \sigma_x^2$ estimates under normal and Student-t models for different levels of censoring (0%, 10%, 20%, 30%).

From these results we can observe that for all considered levels of censoring, the t-MEC model is chosen as the correct model. Under the t-MEC model, the MC CP for α and β are stable, but the MC CP of μ_x is lower than the nominal level (95%). In general, the MC CP

values are higher than those obtained under the normal model. Figure 2 shows the MSE for some parameter estimates (the biases are presented in Figure 6 in the Appendix). Note that, the MSE under the t -MEC model is lower than the obtained under the normal, for different levels of censoring.

Regarding the model choice, the t -MEC model was chosen as the best by the two criteria for all samples.

6.3. Censored model

In this section, the main goal is to study the effect of taking into account censored data on the parameter estimates. We generated $M = 100$ samples from the t -MEC model with $\nu = 5$, setting the censoring level at 20%. The other parameter values are set as in (39). For each dataset, we fitted two models: in *case 1* we use a naive model, where censored responses are not taken into account. In *case 2* we fit a t -MEC model.

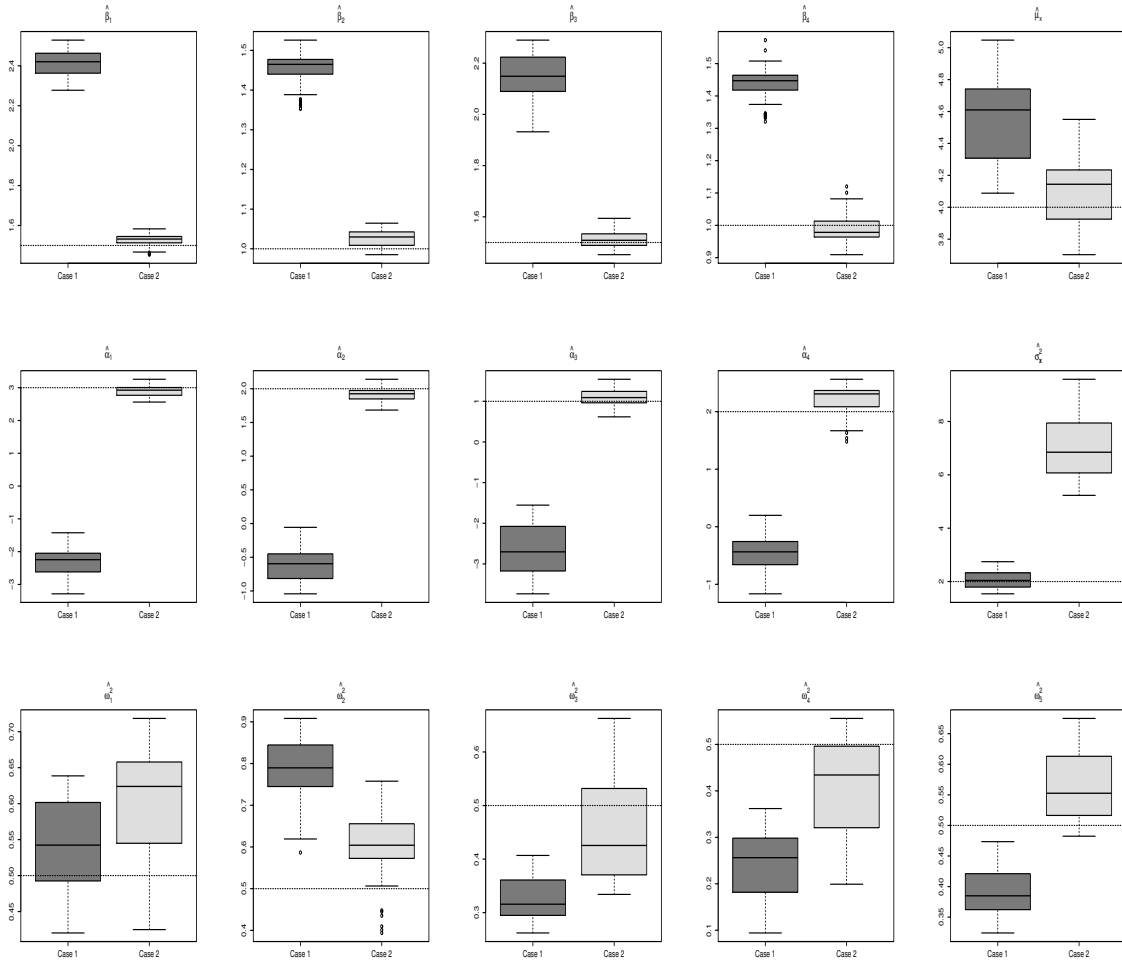


Figure 3: Simulation study 6.3. Boxplots of the parameter estimates. Dotted lines indicate the true parameter value.

Figure 3 shows the box plots corresponding to each parameter estimate considering the

$M = 100$ datasets. Note that, the estimates in *Case 2* are, in general, more precise than those obtained in *Case 1*. It is also possible to note that, in *case 2*, the variability observed in the estimations is smaller than in *case 1*, except for some dispersion parameters. We point out that it is important to consider the effect of censoring in data modeling, avoiding ad-hoc methods.

7. Analysis of case studies

We illustrate the proposed method with a dataset from Chipkevitch *et al.* (1996). The data consist of measurements of the testicular volume of 42 adolescents by using five different techniques: ultrasound (US), graphical method proposed by the authors (I), dimensional measurement (II), Prader orchidometer (III), and ring orchidometer (IV). The ultrasound approach was assumed to be the reference measurement device. Galea-Rojas *et al.* (2002) analyzed the same dataset by fitting the usual normal ME model and recommended considering a data transformation in order to obtain normality. Lachos *et al.* (2010) also analyzed this dataset with the aim of providing a better fit, attempting to avoid possibly unnecessary data transformation. In fact, they considered a joint model of the latent variable and observational errors by using the scale mixtures of skew-normal (SMSN) class of distributions. They also showed evidence of the heavy-tailed behavior of the data (see also Cabral *et al.*, 2014).

Table 2: Chipkevitch data. Testicular volume data (in ml).

Methods						Methods					
i	US	I	II	III	IV	i	US	I	II	III	IV
1	5.0	7.5	5.9	8.0	9.0	22	16.5	10.0	15.3	15.0	15.0
2	5.7	5.0	4.8	6.0	10.0	23	4.5	4.4 (3.5)	4.4 (3.9)	6.0	7.0
3	7.4	5.0	6.8	9.0	12.0	24	5.6	5.0	4.5	4.5	6.0
4	4.4 (2.6)	4.4 (3.5)	4.4 (3.1)	4.4 (4.0)	4.4 (4.0)	25	11.0	7.5	9.7	9.0	11.0
5	5.7	5.0	5.0	6.0	7.0	26	9.2	10.0	11.3	12.0	13.5
6	6.1	5.0	4.4 (4.4)	7.0	8.0	27	8.5	7.5	8.8	12.0	12.0
7	6.2	5.0	6.0	8.0	9.0	28	5.4	5.0	6.1	8.0	8.0
8	10.4	10.0	8.8	10.0	10.0	29	6.7	7.5	7.2	10.0	8.0
9	9.1	7.5	7.9	10.0	11.0	30	5.3	5.0	5.9	8.0	10.0
10	14.8	10.0	13.0	12.0	15.0	31	20.0	20.0	16.3	25.0	22.5
11	16.4	12.5	10.3	17.5	17.5	32	18.8	15.0	16.3	20.0	25.0
12	9.6	7.5	8.2	10.0	11.0	33	13.9	12.5	12.2	15.0	17.5
13	15.7	15.0	19.8	20.0	20.0	34	9.4	10.0	10.3	12.0	13.5
14	4.4 (3.0)	4.4 (2.0)	4.4 (2.0)	4.4 (3.0)	4.4 (4.0)	35	9.1	7.5	10.8	12.0	12.0
15	16.4	15.0	17.3	20.0	20.0	36	14.1	15.0	13.0	13.5	15.0
16	17.6	15.0	17.3	20.0	22.5	37	9.3	10.0	8.4	10.0	10.0
17	10.0	7.5	7.9	12.0	12.0	38	20.9	20.0	22.1	25.0	25.0
18	4.4 (4.1)	4.4 (3.5)	4.4 (4.4)	4.4 (4.0)	6.0	39	11.5	10.0	10.6	15.0	13.5
19	12.7	10.0	11.4	12.0	12.0	40	9.7	10.0	9.7	11.0	12.0
20	4.4 (2.7)	4.4 (3.5)	4.4 (4.1)	4.4 (2.5)	6.0	41	13.7	12.5	11.6	17.5	15.0
21	10.2	10.0	11.1	12.0	13.5	42	8.9	10.0	8.1	12.0	12.0

To apply our method to this dataset, we censored (randomly) 10% (21 observations) of the data. As a consequence, the detection limit κ_{ij} was fixed at 4.4 for all i and j . Table 2 shows the testicular volume data with the true value in parentheses for the censored observations. We fitted the t-MEC (with $\nu = 6$) and N-MEC models. The EM estimates for the parameters of the two model, as well as their corresponding standard errors (SE) obtained via the empirical information matrix are reported in Table 3. This table shows that the estimates of β , α , ω for the t-MEC and N-MEC models are close. However, the standard errors (SE) of the t-MEC are smaller than those of the N-MEC model, indicating that the our robust model seem to produce more precise estimates.

Table 4 compares the fit of the two models using the model selection criteria (AIC and BIC) discussed in Subsection 6.2. Note that, as expected, the t-MEC model outperform the normal one.

Table 3: Chipkevitch data. ML and SE for parameter estimates.

	t -MEC		N-MEC	
	Estimate	SE	Estimate	SE
α_1	-0.0510	1.1501	-0.0584	1.0995
α_2	-0.6674	0.9077	-0.4205	1.2257
α_3	0.2815	0.9361	0.1172	0.9931
α_4	1.9037	0.9853	1.8075	1.0288
β_1	0.9067	0.1166	0.8959	0.0997
β_2	1.0214	0.0809	0.9792	0.0848
β_3	1.1400	0.1017	1.1371	0.0951
β_4	1.0645	0.1038	1.0619	0.0954
μ_x	9.1089	0.8979	9.9222	1.0681
σ_x^2	18.4174	0.0168	25.0263	0.0124
ω_1	1.1068	0.6503	1.4442	0.8227
ω_2	1.1179	0.4447	1.4313	0.5369
ω_3	1.1339	0.3945	1.9156	0.6718
ω_4	0.9437	0.4708	1.1390	0.5460
ω_5	1.1536	0.4261	1.5493	0.5580

Table 4: Chipkevitch data. Model comparison criteria.

	t -MEC	N-MEC
Log-likelihood	-398.4389	-401.4635
AIC	826.8777	832.9269
BIC	877.0843	883.1336

8. Conclusions

In this paper, we introduce the multivariate ME model with censored responses based on the Student-t distribution, the so-called t -MEC model. This model considers the possibility of censoring in the surrogate covariate and the response. Moreover, we assume that the latent unobserved covariate and random observational errors follow a multivariate Student-t distribution, which provides a robust alternative to the usual Gaussian model. For the parameter estimation, an ECM algorithm based on some statistical properties of the multivariate truncated Student-t distribution is developed to obtain ML estimates. Some simulation studies revealed that our proposed method generates less biased estimates of model parameters than the case when the censoring scheme is not taken into account. Moreover, we showed that the use of the Student-t distribution generates better results than the normal one, in the context of the censored ME models.

Of course, further extensions of the current work are possible. For example, the proposed method can be naturally extended by considering the family of scale mixtures of normal (SMN) and skew-normal (SMSN) distributions. An efficient estimation procedure to obtain ML estimates of model parameters can be implemented by using a stochastic approximation of the traditional EM (SAEM) algorithm. Other extensions include, a Bayesian treatment via Markov chain Monte Carlo (MCMC) sampling methods in the context of SMN-MEC and SMSN-MEC models (Lachos *et al.*, 2010).

Acknowledgements

L. A. Matos acknowledges support from FAPESP-Brazil (Grants 11/22063-9 and 15/05385-3). L.M. Castro was partially supported by CONICYT-Chile through BASAL project CMM, Universidad de Chile and Grant FONDECYT 1130233 from the Chilean government. The research of V. H. Lachos was partially supported by CNPq-Brazil (Grant 305054/2011-2) and FAPESP-Brazil (Grant 2014/02938-9). The research of Celso R. B. Cabral was supported by CNPq (via BPPesq and Universal Project 2014), and FAPEAM (via Universal Amazonas Project).

References

- Abarin, T., Li, H., Wang, L. & Briollais, L. (2014). On method of moments estimation in linear mixed effects models with measurement error on covariates and response with application to a longitudinal study gene-environment interaction. *Statistics in Biosciences*, **6**, 1–18.
- Ash, R. B. (2000). *Probability and Measure Theory*. Academic Press, San Diego.
- Bandyopadhyay, D., Castro, L. M., Lachos, V. H. & Pinheiro, H. (2015). Robust joint nonlinear mixed-effects models and diagnostics for censored HIV viral loads with CD4 measurement error. *Journal of Agricultural, Biological and Environmental Statistics*, **20**, 121–139.
- Barnett, V. D. (1969). Simultaneous pairwise linear structural relationships. *Biometrics*, **25**, 129–142.
- Buonaccorsi, J. (2010). *Measurement Error: Models, Methods and Applications*. Chapman & Hall/CRC, Boca Raton.
- Buonaccorsi, J., Demidenko, E. & Tosteson, T. (2000). Estimation in longitudinal random effects models with measurement error. *Statistica Sinica*, **10**, 885–903.
- Cabral, C. R. B., Lachos, V. H. & Borelli, C. (2014). Multivariate measurement error models using finite mixtures of skew-Student t distributions. *Journal of Multivariate Analysis*, **124**, 179–198.
- Carroll, R. J., Lin, X. & Wang, N. (1997). Generalized linear mixed measurement error models. In *Modelling Longitudinal and Spatially Correlated Data*, volume 122 of *Lecture Notes in Statistics*, pages 321–330. Springer-Verlag.
- Carroll, R. J., Ruppert, D., Stefanski, L. A. & Crainiceanu, C. M. (2006). *Measurement Error in Nonlinear Models*. Chapman & Hall/CRC, Boca Raton, second edition.
- Castro, L., Costa, D. R., Prates, M. O. & Lachos, V. H. (2015). Likelihood-based inference for Tobit confirmatory factor analysis using the multivariate distribution. *Statistics and Computing*. DOI 10.1007/s11222-014-9502-0.
- Cheng, C. L. & Van Ness, J. W. (1999). *Statistical Regression with Measurement Error*. Arnold, London.

- Chipkevitch, E., Nishimura, R. T., Tu, D. G. S. & Galea-Rojas, M. (1996). Clinical measurement of testicular volume in adolescents: Comparison of the reliability of 5 methods. *The Journal of Urology*, **156**, 2050–2053.
- Dempster, A., Laird, N. & Rubin, D. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society, Series B*, **39**, 1–38.
- Dumitrescu, L. (2010). Estimation for a longitudinal linear model with measurement errors. *Electronic Journal of Statistics*, **4**, 486–524.
- Fuller, W. A. (1987). *Measurement Error Models*. John Wiley and Sons, New York.
- Galea-Rojas, M., Bolfarine, H. & de Castro, M. (2002). Local influence in comparative calibration models. *Biometrical journal*, **44**, 59–81.
- Ho, H. J., Lin, T. I., Chen, H. Y. & Wang, W. L. (2012). Some results on the truncated multivariate t distribution. *Journal of Statistical Planning and Inference*, **142**, 25–40.
- Kotz, S. & Nadarajah, S. (2004). *Multivariate t Distributions and Their Applications*. Cambridge University Press, Cambridge.
- Lachos, V. H., Labra, F. V., Bolfarine, H. & Ghosh, P. (2010). Multivariate measurement error models based on scale mixtures of the skew-normal distribution. *Statistics*, **44**, 541–556.
- Lange, K. L., Little, R. J. A. & Taylor, J. M. G. (1989). Robust statistical modeling using t distribution. *Journal of the American Statistical Association*, **84**, 881–896.
- Lin, T. I. (2010). Robust mixture modeling using multivariate skew t distributions. *Statistics and Computing*, **20**, 343–356.
- Lin, T. I. & Lee, J. C. (2006). A robust approach to t linear mixed models applied to multiple sclerosis data. *Statistics in Medicine*, **25**, 1397–1412.
- Louis, T. A. (1982). Finding the observed information matrix when using the EM algorithm. *Journal of the Royal Statistical Society, Series B*, **44**, 226–233.
- Lucas, A. (1997). Robustness of the Student-t based M-estimator. *Communications in Statistics-Theory and Methods*, **26**, 1165–1182.
- Matos, L. A., Prates, M. O., Chen, M. H. & Lachos, V. H. (2013). Likelihood-based inference for mixed-effects models with censored response using the multivariate-t distribution. *Statistica Sinica*, **23**, 1323–1342.
- Meilijson, I. (1989). A fast improvement to the EM algorithm on its own terms. *Journal of the Royal Statistical Society, Series B*, pages 127–138.
- Meng, X. & Rubin, D. (1993). Maximum likelihood estimation via the ECM algorithm: a general framework. *Biometrika*, **80**, 267–278.
- Wang, J. & Genton, M. G. (2006). The multivariate skew-slash distribution. *Journal of Statistical Planning and Inference*, **136**, 209–220.

- Wu, L. (2010). *Mixed Effects Models for Complex Data*. Chapman & Hall/CRC, Boca Raton, FL.
- Zhang, S., Midthune, D., Guenther, P. M., Krebs-Smith, S. M., Kipnis, V., Dodd, K. W., Buckman, D. W., Tooze, J. A., Freedman, L. & Carroll, R. J. (2011). A new multivariate measurement error model with zero-inflated dietary data, and its application to dietary assessment. *The Annals of Applied Statistics*, **5**, 1456–1487.

Appendix

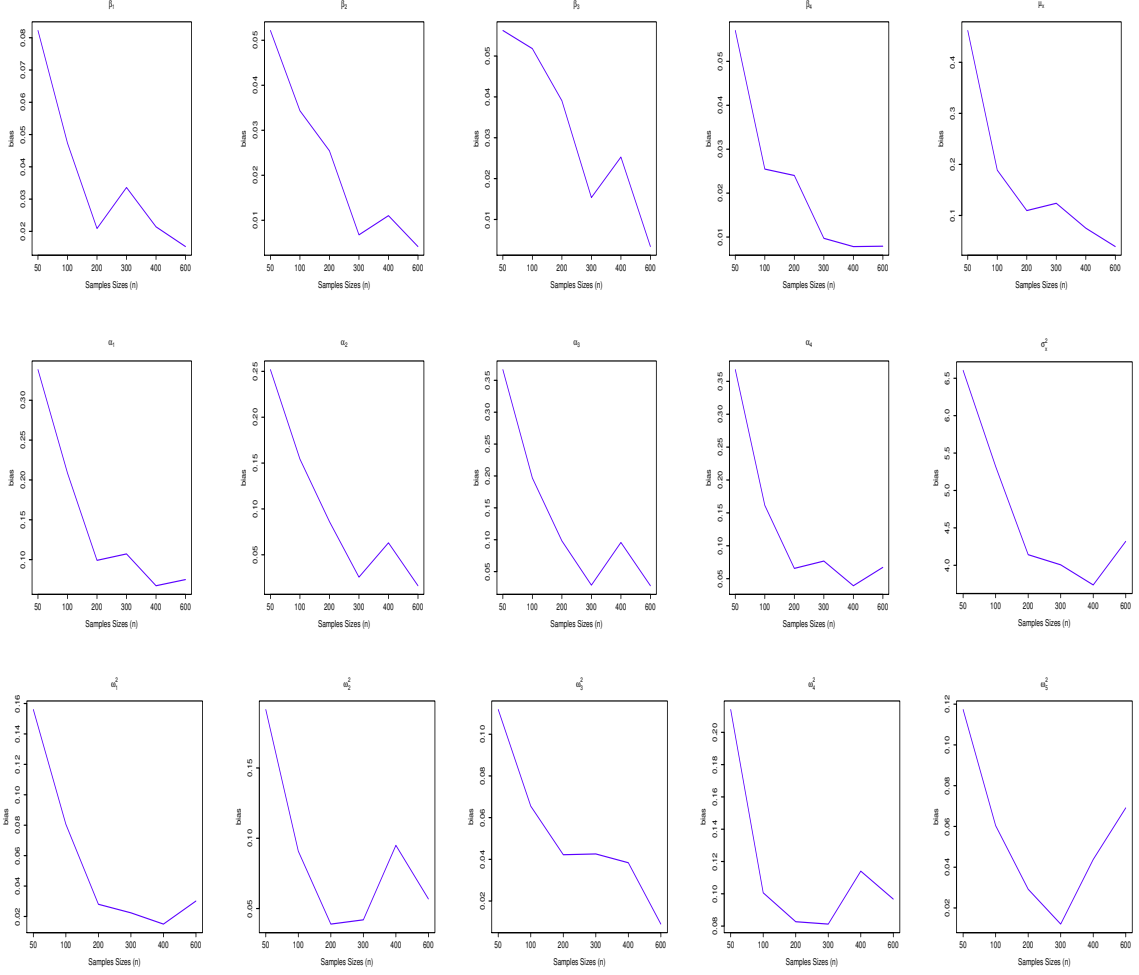


Figure 4: Simulation 6.1. Bias of parameter estimates under the t -MEC model considering 10% of censoring.

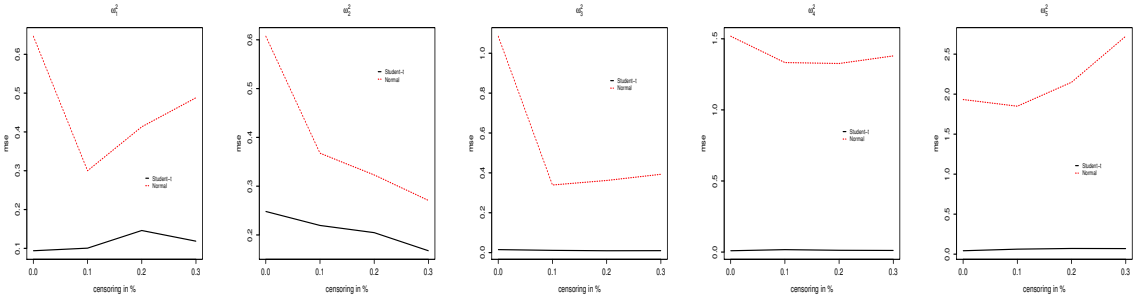


Figure 5: Simulation 6.2. MSE of parameter estimates under the t -MEC model considering 10% of censoring.

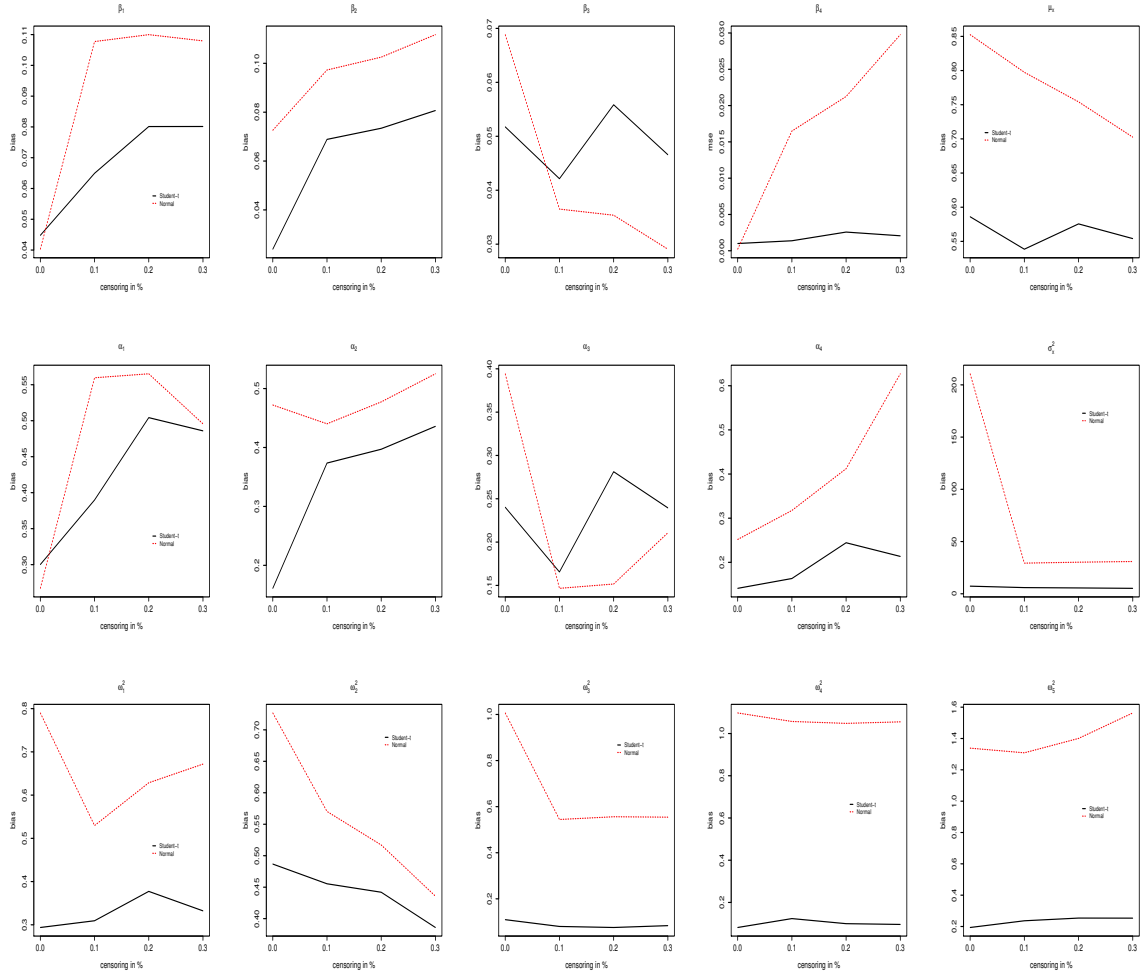


Figure 6: Simulation 6.2 Bias of parameter estimates under the t -MEC model considering different levels of censoring .

Centro de Investigación en Ingeniería Matemática (CI²MA)

PRE-PUBLICACIONES 2015

- 2015-25 JULIO ARACENA, ADRIEN RICHARD, LILIAN SALINAS: *Fixed points in conjunctive networks and maximal independent sets in graph contractions*
- 2015-26 MARGARETH ALVES, JAIME MUÑOZ-RIVERA, MAURICIO SEPÚLVEDA: *Exponential stability to Timoshenko system with shear boundary dissipation*
- 2015-27 RICARDO OYARZÚA, PAULO ZUÑIGA: *Analysis of a conforming finite element method for the Boussinesq problem with temperature-dependent parameters*
- 2015-28 FERNANDO BETANCOURT, FERNANDO CONCHA, LINA URIBE: *Settling velocities of particulate systems part 17. Settling velocities of individual spherical particles in power-law non-Newtonian fluids*
- 2015-29 RAIMUND BÜRGER, PEP MULET, LIHKI RUBIO: *Polynomial viscosity methods for multispecies kinematic flow models*
- 2015-30 ELIGIO COLMENARES, GABRIEL N. GATICA, RICARDO OYARZÚA: *An augmented fully-mixed finite element method for the stationary Boussinesq problem*
- 2015-31 JESSIKA CAMAÑO, CRISTIAN MUÑOZ, RICARDO OYARZÚA: *Analysis of a mixed finite element method for the Poisson problem with data in L^p , $2n/(n+2) < p < 2$, $n = 2, 3$*
- 2015-32 GABRIEL N. GATICA, FILANDER A. SEQUEIRA: *A priori and a posteriori error analyses of an augmented HDG method for a class of quasi-Newtonian Stokes flows*
- 2015-33 MARIO ÁLVAREZ, GABRIEL N. GATICA, RICARDO RUIZ-BAIER: *A posteriori error analysis for a viscous flow-transport problem*
- 2015-34 ANDREA BARTH, RAIMUND BÜRGER, ILJA KRÖKER, CHRISTIAN ROHDE: *A hybrid stochastic Galerkin method for uncertainty quantification applied to a conservation law modelling a clarifier-thickener unit with several random sources*
- 2015-35 RAIMUND BÜRGER, JULIO CAREAGA, STEFAN DIEHL, CAMILO MEJÍAS, INGMAR NOPENS, PETER VANROLLEGHEM: *A reduced model and simulations of reactive settling of activated sludge*
- 2015-36 CELSO R. B. CABRAL, LUIS M. CASTRO, VÍCTOR H. LACHOS, LARISSA A. MATOS: *Multivariate measurement error models based on student-t distribution under censored responses*

Para obtener copias de las Pre-Publicaciones, escribir o llamar a: DIRECTOR, CENTRO DE INVESTIGACIÓN EN INGENIERÍA MATEMÁTICA, UNIVERSIDAD DE CONCEPCIÓN, CASILLA 160-C, CONCEPCIÓN, CHILE, TEL.: 41-2661324, o bien, visitar la página web del centro: <http://www.ci2ma.udec.cl>



**CENTRO DE INVESTIGACIÓN EN
INGENIERÍA MATEMÁTICA (CI²MA)
Universidad de Concepción**



Casilla 160-C, Concepción, Chile
Tel.: 56-41-2661324/2661554/2661316
<http://www.ci2ma.udec.cl>

