

CALIBRAÇÃO LINEAR COM MISTURAS DE ESCALA NORMAL ASSIMÉTRICA

Marcos Antonio Alves Pereira

Dissertação submetida como requerimento parcial para obtenção do
grau de Mestre em Estatística pela Universidade Federal de
Pernambuco

Área de Concentração: **Estatística Aplicada**

Orientadora: **Profa. Dra. Betsabé Grimalda Blas Achic**

Recife, fevereiro de 2013

Catálogo na fonte
Bibliotecária Monick Raquel Silvestre da Silva, CRB4-1217

Pereira, Marcos Antonio Alves

**Calibração linear com misturas de escala normal
assimétrica / Marcos Antonio Alves Pereira. - Recife: O
Autor, 2013.**

xi, 73 f.: il., fig., tab.

Orientadora: Betsabé Grimalda Blas Achic.

**Dissertação (mestrado) - Universidade Federal de
Pernambuco. CCEN. Estatística, 2013.**

Inclui bibliografia e apêndice.

**1. Estatística aplicada. 2. Calibração. 3. Modelos lineares
(estatística) I. Achic, Betsabé Grimalda Blas (orientadora). II.
Título.**

310

CDD (23. ed.)

MEI2013 – 028

Universidade Federal de Pernambuco
Pós-Graduação em Estatística

Recife, 19 de fevereiro de 2013.

Nós recomendamos que a dissertação de mestrado de autoria de

Marcos Antonio Alves Pereira

Intitulada

“Calibração Linear com Misturas de Escala Normal Assimétrica”

Seja aceita como cumprimento parcial dos requerimentos para o grau de Mestre em Estatística.

Coordenador da Pós-Graduação em Estatística

Banca Examinadora:

Betsabé Grimalda Blas Achic

Orientadora / UFPE

Claudia Regina Oliveira de Paiva Lima

UFPE

Mônica Carneiro Sandoval

USP

Este documento será anexado à versão final da tese.

Dedicatória

Aos meus familiares e a todos aqueles que de alguma forma contribuíram para a realização deste trabalho.

Agradecimentos

Gostaria de agradecer à minha orientadora Profa. Dra. Betsabé Grimalda Blas Achic pela atenção durante o desenvolvimento desse projeto de formatura para obtenção do grau de Mestre em Estatística.

Agradeço à Universidade de Federal de Pernambuco, em especial, ao Prof. Dr. Francisco Cribari Neto e à Valéria Bittencourt, pelo apoio à realização desse projeto.

Agradeço também ao Prof. Dr. Clécio da Silva Ferreira pelas várias dúvidas sanadas durante a produção deste texto.

As professoras Mônica Carneiro Sandoval e Claudia Regina Oliveira de Paiva Lima, pelas observações, sugestões e correções que muito contribuíram para o melhoramento deste trabalho.

Por fim, agradeço a CAPES pelo suporte financeiro concedido.

Resumo

Neste trabalho apresenta-se um modelo de calibração estatística linear com repetição na variável resposta e assumindo que os erros têm distribuição pertencente a uma classe ou família de distribuições denominada *Misturas de Escala Normal Assimétrica (MENA)*. Essa família de distribuições é uma generalização de várias distribuições que permite a escolha de uma distribuição simétrica ou assimétrica. Na literatura, os modelos de calibração supõem, em grande parte, que os erros têm distribuição *normal*, no entanto, a distribuição *normal* é inadequada para dados com observações destoantes e assimetria. A estimação dos parâmetros do modelo proposto é feita numericamente por meio do *algoritmo EM*, devido a sua facilidade de implementação e eficiência. Realizou-se um estudo de simulação em que o estimador de máxima verossimilhança via *algoritmo EM* mostrou-se consistente. O modelo de calibração linear proposto foi aplicado em dois conjuntos de dados reais relacionados à química analítica.

Palavras chaves: calibração, simetria, assimetria, *MENA*, *algoritmo EM*.

Abstract

This work presents a statistical linear calibration model with replication on the response variable by assuming that the error model follows the family of *Scale Mixtures of the Skew-Normal (SMSN)* distributions. This family of distributions is a generalization of many distributions that allow the choice of a asymmetric or asymmetric distribution. In the literature, most of calibration models assume that the errors are normally distributed, however, the *normal* distribution is inadequate for data atypical observations and asymmetry. The estimation of model parameters is done numerically by the *EM algorithm*, because it is efficient and easy to implement. A simulation study is carried out and it was verified that the maximum likelihood estimators via the *EM algorithm* are consistent. The new approach is applied to two real data set from chemical analysis.

Keywords: calibration, symmetry, asymmetry, *SMSN*, *EM algorithm*.

Lista de Tabelas

4.1	Distribuição t de Student Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$	44
4.2	Distribuição t de Student Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$	45
4.3	Distribuição Normal Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$	46
4.4	Distribuição Normal Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$	46
4.5	Distribuição Slash Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$	47
4.6	Distribuição Slash Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$	48
4.7	Distribuição Exponencial Potência Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$	49
4.8	Distribuição Exponencial Potência Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$	50
5.1	Concentração (mg/l) e absorbâncias das soluções-padrão de zinco.	53
5.2	Estimativas dos parâmetros para as distribuições <i>tA</i> , <i>NA</i> , <i>SLA</i> e <i>EPA</i>	54
5.3	Crerios de informação - Concentração de zinco.	55
5.4	Concentração (mgL ⁻¹) e altura do pico (cm) das soluções-padrão de benzatona.	56

5.5	Estimativas dos parâmetros para as distribuições tA , NA , SLA e EPA	57
5.6	Cr�terios de informa��o - Concentra��o de benzatona.	57

Lista de Figuras

2.1	Densidades de distribuições $MENA(\lambda; \boldsymbol{\tau})$	14
3.1	Diferença entre os erros na calibração <i>clássica</i> (a) e <i>inversa</i> (b)	20

Sumário

1	Introdução	1
2	Família de Distribuições $MENA$	6
2.1	Casos particulares da família $MENA$	12
2.2	Casos particulares da família MEN	15
3	Modelo de Calibração Linear	17
3.1	Modelo de calibração linear simples	17
3.1.1	Estimadores <i>clássico</i> e <i>inverso</i>	18
3.2	Modelo de calibração linear $MENA$	21
3.2.1	<i>Algoritmo EM</i>	23
3.2.2	<i>Algoritmo EM</i> para o modelo proposto	25
3.2.3	Matriz de informação de Fisher observada	32
3.2.4	Matriz hessiana	40

4	Estudo de Simulação	42
5	Aplicações	52
5.1	Concentração de zinco	53
5.2	Concentração de benzatona	55
6	Conclusões	59
6.1	Considerações finais	59
6.2	Perspectivas para pesquisas futuras	60
	Referências Bibliográficas	60

Introdução

Modelos de calibração estatística são frequentemente utilizados em diversas áreas do conhecimento. Por exemplo, em química analítica, calibração é o processo que permite estabelecer a relação entre um conjunto de variáveis, que geralmente é uma série de medidas físicas com o objetivo de determinar a concentração de um composto. Na engenharia, esses modelos têm sido utilizados para calibrar instrumentos para medição de grandezas físicas, assim como na medicina, para calibrar instrumentos para medição de sinais vitais, tais como pressão sanguínea, nível de temperatura e colesterol. Em economia ou administração, modelos de calibração podem ser utilizados para tomar decisões em cenários de instabilidade já vividos por meio de previsões das possíveis causas. Segundo Figueiredo *et al.* (2010), modelos de calibração são úteis quando queremos determinar concentrações de soluções, composições de material, fazer avaliação de instrumentos físicos utilizados para a medição de grandezas físicas, fazer avaliações dos efeitos de doses de medicamentos ou para redução de prejuízo ou maximização de lucro.

Modelos de calibração são, usualmente, constituídos de dois estágios. No primeiro estágio a variável aleatória y_i ($i = 1, \dots, n$) é observada em função dos valores prefixados

x_i ($i = 1, \dots, n$), e no segundo estágio a variável aleatória y_{0i} ($i = 1, \dots, m$), associada a um valor desconhecido x_0 , é observada. Os parâmetros de uma função, previamente estabelecida, que relaciona as variáveis é estimada levando-se em consideração os dois estágios. O interesse principal de um modelo de calibração é estimar o valor do parâmetro x_0 . Os estimadores *clássico* e *inverso* são os mais utilizados em modelos de calibração linear simples. No Capítulo 3 serão levantadas algumas vantagens e desvantagens destes dois estimadores.

Ao longo dos anos, várias propostas sobre calibração foram surgindo a fim de se obter boa concordância entre os dados observados e calculados, inclusive com abordagens computacionais devido ao avanço da tecnologia.

Novas abordagens de calibração têm um grande potencial e aplicabilidade em todas as áreas do conhecimento. O avanço da tecnologia também minimizou erros na leitura e coleta de dados, pois equipamentos de medição mais precisos foram surgindo nas últimas décadas de forma a aumentar a confiabilidade nos resultados das calibrações.

Branco (1997) ilustra a utilização de calibração aplicada à genética, conjecturando a situação de um possível acidente nuclear, onde um certo trabalhador foi exposto a uma dose de radiação desconhecida, x_0 . O número de células efetivamente afetadas pela radiação foi determinado por meio de amostras de sangue do trabalhador e expresso por y_0 . Por questões éticas, o estimador da dose, x_0 , de radiação recebida pelo trabalhador deve ser obtido em um experimento que envolva medições laboratoriais que espelhem as mesmas características do hipotético acidente nuclear com cobaias, por exemplo. Assim, é suficiente expor as cobaias a doses conhecidas de radiação $x_1, \dots, x_n$ e anotar os efeitos citogenéticos produzidos $y_1, \dots, y_n$. Este procedimento é conhecido na literatura por *calibração controlada*.

Alguns pesquisadores deram importantes contribuições ao estudo de calibração estatística, como por exemplo, Galea-Rojas (1995) que apresentou um estudo sobre calibração comparativa estrutural e funcional, Francisconi (1996) que realizou um estudo para comparar instrumentos de medição usando calibração comparativa e Lima (1996) que apresentou um estudo sobre calibração absoluta com erros nas variáveis. Blas *et al.* (2007) estudaram o modelo de calibração com erros de medida do tipo *Berkson* assumindo homoscedasticidade dos erros. Blas e Sandoval (2010) estenderam o estudo de Blas *et al.* (2007) assumindo heteroscedasticidade dos erros de medida.

Recentemente, modelos de calibração que assumem distribuições dos erros com caudas pesadas foram estudadas por Marciano (2012) e Figueiredo *et al.* (2010). Marciano (2012) estudou modelos de calibração linear com erros com distribuições simétricas, tais como a distribuição *normal*, *t de Student*, *exponencial potência* e *logística tipo II*. Figueiredo *et al.* (2010) estudaram modelos de calibração linear considerando distribuições assimétricas para os erros, tais como *normal assimétrica* e *t-normal assimétrica*, sob os enfoques frequentista e bayesiano.

O estudo de modelos de regressão é útil para o desenvolvimento de modelos de calibração, pois modelos de calibração, de forma geral, são extensões de modelos de regressão. Segundo Brereton (2003), calibração, na sua forma mais simples, é uma forma de regressão.

Ferreira *et al.* (2011) propuseram um modelo de regressão linear que assume que a distribuição dos erros pertençam a uma classe ou família de distribuições denominada *Misturas de Escala Normal Assimétrica (MENA)* que abrange várias distribuições simétricas e assimétricas, assim vários modelos de regressão linear estudados podem ser considerados como casos particulares. Pesquisadores como Branco e Dey (2001) e Zeller *et al.* (2010) também deram importantes contribuições no desenvolvimento da classe *MENA*.

A maioria dos modelos de calibração presentes na literatura supõem que os erros são normalmente distribuídos, no entanto, a distribuição *normal* é inadequada para dados com observações destoantes ou atípicas e assimetria. O uso da família de distribuições *MENA*, que é uma generalização de várias distribuições simétricas e assimétricas, pode ser uma solução para estes problemas.

Para solucionar esses problemas, neste trabalho é proposto um novo modelo de calibração linear que considera que a distribuição dos erros pertence a classe *MENA* (Ferreira *et al.*, 2011). Este novo modelo de calibração linear acomoda dados provenientes de experimentos com replicatas, podendo-se variar o número destas repetições.

Para a estimação dos parâmetros do modelo proposto é utilizado o *algoritmo EM* por facilitar a maximização do logaritmo da função de verossimilhança.

Este trabalho é composto por seis capítulos e um apêndice, sendo este o primeiro capítulo e os demais estão resumidos a seguir.

O Capítulo 2 apresenta a família de distribuições *MENA* e *MEN* (*Misturas de Escala Normal*), bem como a descrição de alguns casos particulares que são utilizados no modelo de calibração linear proposto.

O Capítulo 3 apresenta um breve resumo sobre os estimadores *clássico* e *inverso* utilizados no modelo de calibração usual. Além disso, apresentação do modelo de calibração linear proposto, os passos para implementação do *algoritmo EM* e as expressões dos elementos da matriz de informação de Fisher observada e da matriz hessiana.

No Capítulo 4 são apresentados os resultados de simulação para as distribuições *t de Student assimétrica*, *normal assimétrica*, *slash assimétrica* e *exponencial potência assimétrica*, pertencentes à família *MENA*, nos quais se verifica a consistência dos estimadores do modelo de calibração linear proposto.

O Capítulo 5 apresenta duas aplicações do modelo de calibração linear proposto utilizando dados reais relacionados à química analítica obtidos em Neto *et al.* (2007) e Oliveira e Aguiar (2009).

Considerações finais e perspectivas de trabalhos futuros são apresentadas no Capítulo 6.

No Apêndice está disponível o código em *R* para geração de amostras do modelo de calibração linear proposto com erros segundo a distribuição *t de Student assimétrica* e para a estimação dos parâmetros via *algoritmo EM*.

Família de Distribuições *MENA*

Distribuições derivadas de misturas de escala da distribuição normal são comuns na literatura e bastante utilizadas em procedimentos estatísticos, como pode ser encontrado em Dempster *et al.* (1980) e Lange *et al.* (1989). A mistura de escala da distribuição normal (Andrews e Mallows, 1974) fornece um grupo de distribuições com caudas pesadas que são frequentemente utilizadas para a inferência robusta de dados simétricos.

Porém, para analisar conjuntos de dados assimétricos essas distribuições não são adequadas, surgindo, então, a necessidade de se obter distribuições assimétricas. Nesse sentido pesquisadores introduziram famílias de distribuições paramétricas flexíveis que acomodem desvios de normalidade e, assim, amenizem a necessidade de transformação de dados.

A família de distribuições *MENA* estudada por Branco e Dey (2001), Zeller *et al.* (2010) e Ferreira *et al.* (2011) é bastante útil para analisar dados com observações destoantes ou atípicas levando-se em consideração sua assimetria, que inclui a distribuição *normal* e *normal assimétrica* como casos particulares, podendo ser útil para a modelagem de dados, análise estatística e estudos de métodos robustos, além da facilidade de geração de dados

de distribuições membros dessa família e da estimação de seus parâmetros pelo método de máxima verossimilhança via *algoritmos EM* genuínos.

A família de distribuições $MENA(\mu, \sigma^2, \lambda; H)$, segundo Ferreira *et al.* (2011), tem a seguinte função densidade de probabilidade

$$f(y) = 2 \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \Phi\left(\lambda \frac{y - \mu}{\sigma}\right) dH(u; \boldsymbol{\tau}), \quad y \in \mathbb{R}, \quad (2.1)$$

onde $\mu \in \mathbb{R}$ é o parâmetro de locação, σ^2 é o parâmetro de dispersão, $\lambda \in \mathbb{R}$ é o parâmetro de assimetria, U é uma variável aleatória positiva com função distribuição acumulada $H(u, \boldsymbol{\tau})$ indexada pelo vetor de parâmetros $\boldsymbol{\tau}$ supostamente conhecido e $\kappa(u)$ é uma função estritamente positiva.

A função geradora de momentos (fgm) de uma variável aleatória $Y \sim MENA(\mu, \sigma^2, \lambda; H)$ é dada pela expressão geral

$$M_y(t) = \mathbb{E}(e^{tY}) = \int_0^\infty 2e^{t\mu + \frac{t^2}{2}\kappa(u)\sigma^2} \Phi\left(t \frac{\sigma \lambda \kappa(u)}{\sqrt{1 + \lambda^2 \kappa(u)}}\right) dH(u), \quad t \in \mathbb{R}. \quad (2.2)$$

A função ϕ , função densidade de probabilidade (fdp) e a função Φ , função distribuição acumulada (fda), da distribuição *normal* são definidas respectivamente por

$$\begin{aligned} \phi(y|\mu, \sigma^2 \kappa(u)) &= \frac{1}{\sqrt{2\pi} \sigma \sqrt{\kappa(u)}} \exp\left(-\frac{(y - \mu)^2}{2\sigma^2 \kappa(u)}\right) \text{ e} \\ \Phi\left(\lambda \frac{y - \mu}{\sigma}\right) &= \int_{-\infty}^{\left(\lambda \frac{y - \mu}{\sigma}\right)} \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{w^2}{2}\right) dw. \end{aligned}$$

Alguns casos particulares e triviais da família de distribuições *MENA*:

- se $\mu = 0$ e $\sigma^2 = 1$, tem-se uma família de distribuições *MENA Padrão*: $MENA(\lambda; H)$,
- se $\lambda = 0$, tem-se uma família de distribuições *Misturas de Escala Normal*:
 $MEN(\mu, \sigma^2; H)$,
- se $\kappa(u) = 1$ e H degenerada, tem-se a distribuição *Normal Assimétrica*: $NA(\mu, \sigma^2, \lambda)$,
- se $\lambda = 0$, $\kappa(u) = 1$ e H degenerada, tem-se a distribuição *Normal*: $N(\mu, \sigma^2)$ e
- se $\mu = \lambda = 0$, $\sigma^2 = \kappa(u) = 1$ e H degenerada, tem-se a distribuição *Normal Padrão*:
 $N(0, 1)$.

Proposição 2.1. *Seja $Y \sim MENA(\mu, \sigma^2, \lambda; H)$ cuja fdp é dada por (2.1), então uma forma alternativa de representar esta fdp é dada por*

$$f(y) = 2 \int_0^\infty \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \phi(t|\lambda(y - \mu), \sigma^2) h(u; \tau) dt du \quad (2.3)$$

Demonstração. Observe que

$$\begin{aligned} \int_{-\infty}^0 \phi(t|\lambda \frac{y - \mu}{\sigma}, 1) dt &= \mathbb{P}(Z \leq 0), \quad Z \sim N(-\lambda \frac{y - \mu}{\sigma}, 1) \\ &= \mathbb{P}(-Z \geq 0) \\ &= \mathbb{P}(T \geq 0), \quad T \sim N(\lambda \frac{y - \mu}{\sigma}, 1) \\ &= \int_0^\infty \phi(t|\lambda \frac{y - \mu}{\sigma}, 1) dt \\ &= \int_0^\infty \phi(t|\lambda(y - \mu), \sigma^2) dt. \end{aligned}$$

Assim,

$$\begin{aligned}
f(y) &= 2 \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \Phi\left(\lambda \frac{y - \mu}{\sigma}\right) dH(u) \\
&= 2 \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) h(u; \tau) du \int_{-\infty}^{\lambda \frac{y - \mu}{\sigma}} \phi(t|0, 1) dt \\
&= 2 \int_0^\infty \int_{-\infty}^0 \phi(y|\mu, \sigma^2 \kappa(u)) \phi(t - \lambda \frac{y - \mu}{\sigma}, 1) h(u; \tau) dt du \\
&= 2 \int_0^\infty \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \phi(t|\lambda(y - \mu), \sigma^2) h(u; \tau) dt du.
\end{aligned}$$

A esperança e a variância da variável aleatória $Y \sim MENA(\mu, \sigma^2, \lambda; H)$ são dadas, respectivamente, por

$$\mathbb{E}(Y) = \mu + \sqrt{\frac{2}{\pi}} \sigma \lambda \mathbb{E}_U \left[\frac{\kappa(U)}{\sqrt{1 + \lambda^2 \kappa(U)}} \right] \quad \text{e} \quad (2.4)$$

$$\text{Var}(Y) = \sigma^2 \left\{ \mathbb{E}_U[\kappa(U)] - \frac{2}{\pi} \lambda^2 \mathbb{E}_U^2 \left[\frac{\kappa(U)}{\sqrt{1 + \lambda^2 \kappa(U)}} \right] \right\}. \quad (2.5)$$

Proposição 2.2. *Seja $Y \sim MENA(\mu, \sigma^2, \lambda; H)$. Então uma representação estocástica para Y é dada por*

$$\begin{aligned}
Y|U = u &\sim NA(\mu, \sigma^2 \kappa(u), \lambda \sqrt{\kappa(u)}), \\
U &\sim H(u; \tau).
\end{aligned} \quad (2.6)$$

Demonstração. De (2.1), a função densidade conjunta de (Y, U) é dada por

$$f(y, u) = 2\phi(y|\mu, \sigma^2 \kappa(u)) \Phi\left(\lambda \frac{y - \mu}{\sigma}\right) h_U(u).$$

Fornecendo $f(y, u) = f(y|u)h_U(u)$, então

$$\begin{aligned} f(y|u) &= 2\phi(y|\mu, \sigma^2\kappa(u))\Phi\left(\lambda\frac{y-\mu}{\sigma}\right) \\ &= 2\phi(y|\mu, \sigma^2\kappa(u))\Phi\left(\lambda\sqrt{\kappa(u)}\frac{y-\mu}{\sigma\sqrt{\kappa(u)}}\right), \text{ logo} \end{aligned}$$

$$Y|U = u \sim NA(\mu, \sigma^2\kappa(u), \lambda\sqrt{\kappa(u)}).$$

Uma forma de apresentar a distribuição *normal assimétrica* é por meio da representação estocástica obtida por Henze (1986) definida a seguir

$$\begin{aligned} Z &= \mu + \sigma\kappa^{1/2}(U)(\delta T + (1 - \delta^2)^{1/2}V) \text{ com } \delta = \frac{\lambda\sqrt{\kappa(u)}}{\sqrt{1 + \lambda^2\kappa(u)}}, \\ Z &\sim NA(\mu, \sigma^2\kappa(u), \lambda\sqrt{\kappa(u)}), \end{aligned} \quad (2.7)$$

com $T \sim NT(0, 1)\mathbb{I}_{(0 < t < \infty)}$ e $V \sim N(0, 1)$, onde $NT(0, 1)$ e $N(0, 1)$ denotam as distribuições *normal truncada padrão* e *normal padrão* respectivamente, sendo $T = |V|$.

Esta representação é de grande utilidade na implementação do *algoritmo EM*, como será visto no Capítulo 3.

Observa-se que é possível gerar uma distribuição independente *normal assimétrica* procedendo em dois passos. Gerando primeiro uma distribuição U e depois uma distribuição condicional $Y|U$ usando (2.7).

Proposição 2.3. *Seja a variável aleatória $X \sim N(\mu, \sigma^2)$ truncada nos dois lados, isto é $a \leq X \leq b$, então a fdp da distribuição truncada de X é dada por*

$$f_X(x) = \frac{\frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\}}{\Phi\left(\frac{b-\mu}{\sigma}\right) - \Phi\left(\frac{a-\mu}{\sigma}\right)}, \quad (2.8)$$

Demonstração.

$$\begin{aligned} f(x|a \leq x \leq b) &= \frac{\text{fdp de } X}{\mathbb{P}(a \leq X \leq b)} \\ &= \frac{\frac{1}{\sigma} \phi\left(\frac{x-\mu}{\sigma}\right)}{\mathbb{P}\left(\frac{a-\mu}{\sigma} \leq \frac{X-\mu}{\sigma} \leq \frac{b-\mu}{\sigma}\right)} \\ &= \frac{\frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\}}{\Phi\left(\frac{b-\mu}{\sigma}\right) - \Phi\left(\frac{a-\mu}{\sigma}\right)}, \end{aligned}$$

em que $\frac{X-\mu}{\sigma} \sim N(0,1)$.

Quando $a = 0$ e $b = \infty$, tem-se uma variável com distribuição *normal truncada à esquerda*, $X \sim NT(\mu, \sigma^2)\mathbb{I}_{(0 \leq x \leq \infty)}$, com fdp

$$\begin{aligned} f_X(x) &= \frac{\frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\}}{1 - \Phi\left(-\frac{\mu}{\sigma}\right)} \\ &= \frac{\frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\}}{\Phi\left(\frac{\mu}{\sigma}\right)}. \end{aligned}$$

No caso da distribuição *normal truncada à esquerda*, o primeiro e segundo momentos são dados, respectivamente, por

$$\mathbb{E}(X) = \mu + \sigma \frac{\phi\left(\frac{\mu}{\sigma}\right)}{\Phi\left(\frac{\mu}{\sigma}\right)} \quad \text{e} \quad (2.9)$$

$$\mathbb{E}(X^2) = \mu^2 + \sigma^2 + \mu\sigma \frac{\phi\left(\frac{\mu}{\sigma}\right)}{\Phi\left(\frac{\mu}{\sigma}\right)}. \quad (2.10)$$

Utilizando a representação estocástica (2.7), pode-se reescrever o modelo (2.6) como

$$\begin{aligned}
Y|T=t, U=u &\stackrel{\text{iid}}{\sim} N\left(\mu + \frac{\sigma\lambda\kappa(u)}{\sqrt{1+\lambda^2\kappa(u)}}t, \frac{\sigma^2\kappa(u)}{1+\lambda^2\kappa(u)}\right), \\
U &\stackrel{\text{iid}}{\sim} H(u; \tau), \\
T &\stackrel{\text{iid}}{\sim} NT(0, 1),
\end{aligned} \tag{2.11}$$

2.1 Casos particulares da família $MENA$

A seguir são apresentadas algumas distribuições mais comuns que compõem a família de distribuições $MENA$. Em todos os casos considera-se $d = (y - \mu)^2/\sigma^2$, distância de Mahalanobis ao quadrado.

- **Distribuição normal assimétrica**, denotada por $Y \sim NA(\mu, \sigma^2, \lambda)$, com fdp dada por

$$f(y) = \frac{2}{\sigma\sqrt{2\pi}} e^{-d/2} \Phi\left(\lambda \frac{y - \mu}{\sigma}\right), \quad y \in \mathbb{R}, \tag{2.12}$$

considerando $\kappa(U) = 1$ e H degenerada.

- **Distribuição t de Student generalizada normal assimétrica**, com $\nu > 0$ graus de liberdade, $\gamma > 0$, denotada por $Y \sim GtA(\mu, \sigma^2, \lambda; \nu, \gamma)$, tem fdp dada por

$$f(y) = \frac{2}{\sigma\sqrt{\gamma\pi}} \frac{\Gamma((\nu+1)/2)}{\Gamma(\nu/2)} \left(1 + \frac{d}{\gamma}\right)^{-(\frac{\nu+1}{2})} \Phi\left(\lambda \frac{y - \mu}{\sigma}\right), \quad y \in \mathbb{R}, \tag{2.13}$$

considerando $\kappa(U) = 1/U$ e $U \sim \text{gama}(\nu/2, \gamma/2)$ com densidade

$$h(u; \nu, \gamma) = \frac{(\gamma/2)^{\nu/2}}{\Gamma(\nu/2)} u^{\nu/2-1} e^{-\gamma u/2}. \tag{2.14}$$

Para $\gamma = \nu$, tem-se a distribuição *t de Student assimétrica*. A distribuição *Cauchy normal assimétrica* é um caso particular desta distribuição, obtida quando $\gamma = \nu = 1$. Também quando $\nu \uparrow \infty$, tem-se a distribuição *normal assimétrica* como caso limite.

- **Distribuição slash assimétrica**, com $\nu > 0$ graus de liberdade, denotada por $Y \sim SLA(\mu, \sigma^2, \lambda; \nu)$, tem fdp dada por

$$f(y) = 2\nu\Phi\left(\lambda\frac{y-\mu}{\sigma}\right)\int_0^1 u^{\nu-1}\phi\left(y|\mu, \frac{\sigma^2}{u}\right)du, \quad y \in \mathbb{R}, \quad (2.15)$$

nesse caso $\kappa(U) = 1/U$ e U com densidade

$$h(u; \nu) = \nu u^{\nu-1}\mathbb{I}_{(0,1)}(u). \quad (2.16)$$

A distribuição *slash assimétrica* se reduz à distribuição *normal assimétrica* quando $\nu \uparrow \infty$.

- **Distribuição normal contaminada assimétrica**, denotada por $Y \sim NCA(\mu, \sigma^2, \lambda; \nu, \gamma)$ com $0 \leq \nu, \gamma \leq 1$, tem fdp dada por

$$f(y) = 2\left\{\nu\phi\left(y|\mu, \frac{\sigma^2}{\gamma}\right) + (1-\nu)\phi(y|\mu, \sigma^2)\right\}\Phi\left(\lambda\frac{y-\mu}{\sigma}\right) \quad y \in \mathbb{R}, \quad (2.17)$$

com $\kappa(U) = 1/U$ e U com densidade

$$h(u; \nu, \gamma) = \nu\mathbb{I}_{(u=\gamma)} + (1-\nu)\mathbb{I}_{(u=1)}. \quad (2.18)$$

A distribuição *normal contaminada assimétrica* se reduz à distribuição *normal assimétrica* quando $\gamma = \nu = 1$ ou quando $\gamma = 1$ e $\nu = 1/2$.

- **Distribuição exponencial potência assimétrica**, com parâmetro $0,5 < \nu \leq 1$, denotada por $Y \sim EPA(\mu, \sigma^2, \lambda; \nu)$, tem fdp dada por

$$f(y) = \frac{2\nu}{2^{\frac{1}{2\nu}} \sigma \Gamma(\frac{1}{2\nu})} e^{-d^\nu/2} \Phi\left(\lambda \frac{y - \mu}{\sigma}\right), \quad y \in \mathbb{R}, \quad (2.19)$$

neste caso, $\kappa(u)$ não possui forma explícita e U tem densidade positiva estável $S^p(u|\nu)$ (Branco e Dey, 2001). A distribuição *exponencial potência assimétrica* se reduz à distribuição *normal assimétrica* quando $\nu = 1$.

Na Figura 2.1 estão apresentados os gráficos das funções de densidade de probabilidade das distribuições $NA(3)$, $GtA(3; 2, 2)$, $SLA(3; 0, 5)$, $NCA(3; 0, 9, 1)$ e $EPA(3; 0, 5)$, nas quais foram considerados os parâmetros $\mu = 0$ e $\sigma^2 = 1$.

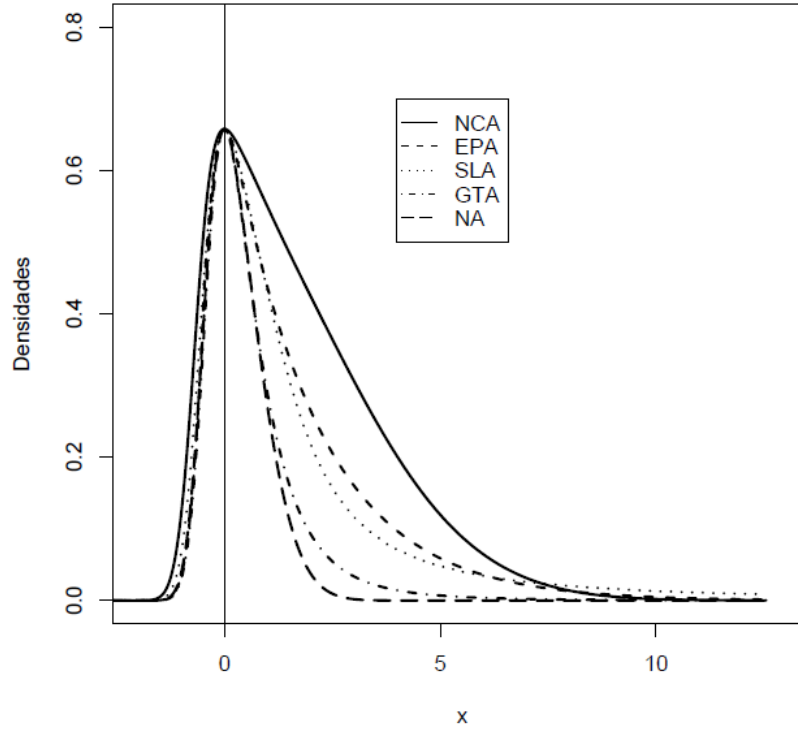


Figura 2.1: Densidades de distribuições $MENA(\lambda; \tau)$.

2.2 Casos particulares da família MEN

Pode-se obter, por meio da família de distribuições $MENA$, versões simétricas de distribuições que compõem a família de distribuições *Misturas de Escala Normal* (MEN). Note que se em (2.1) $\lambda = 0$, tem-se

$$\begin{aligned} f(y) &= 2 \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \Phi\left(0 \frac{y - \mu}{\sigma}\right) dH(u; \tau) \\ &= 2 \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) \frac{1}{2} dH(u; \tau) \\ &= \int_0^\infty \phi(y|\mu, \sigma^2 \kappa(u)) dH(u; \tau). \end{aligned}$$

A fdp da variável aleatória Y não depende de λ , ou seja, $Y \sim MEN(\mu, \sigma^2; H)$.

A seguir algumas distribuições que compõem a família de distribuições MEN . Considerando-se, como na seção anterior, $d = (y - \mu)^2 / \sigma^2$, distância de Mahalanobis ao quadrado.

- **Distribuição normal**, denotada por $Y \sim N(\mu, \sigma^2)$, tem fdp dada por

$$f(y) = \frac{1}{\sigma\sqrt{2\pi}} e^{-d/2}, \quad y \in \mathbb{R}, \quad (2.20)$$

considerando $\kappa(U) = 1$ e H degenerada.

- **Distribuição t de Student generalizada**, com $\nu > 0$ graus de liberdade, $\gamma > 0$, denotada por $Y \sim Gt(\mu, \sigma^2; \nu, \gamma)$, tem fdp dada por

$$f(y) = \frac{1}{\sigma\sqrt{\gamma\pi}} \frac{\Gamma((\nu+1)/2)}{\Gamma(\nu/2)} \left(1 + \frac{d}{\gamma}\right)^{-(\frac{\nu+1}{2})}, \quad y \in \mathbb{R}, \quad (2.21)$$

considerando $\kappa(U) = 1/U$ e $U \sim \text{gama}(\nu/2, \gamma/2)$ com densidade como em (2.14).

Para $\gamma = \nu$, tem-se a distribuição *t de Student*.

- **Distribuição slash**, com $\nu > 0$ graus de liberdade, denotada por $Y \sim SL(\mu, \sigma^2; \nu)$ com $h(u; \nu)$ como em (2.16), $\kappa(U) = 1/U$, tem fdp dada por

$$f(y) = \frac{\nu}{\sqrt{2\pi}} \int_0^1 u^{\nu-1/2} e^{-ud/2} du, \quad y \in \mathbb{R}. \quad (2.22)$$

- **Distribuição normal contaminada**, denotada por $Y \sim NC(\mu, \sigma^2; \nu, \gamma)$ com $0 \leq \nu, \gamma \leq 1$ e $h(u; \nu, \gamma)$ como em (2.18), $\kappa(U) = 1/U$, tem fdp dada por

$$f(y) = \nu \phi\left(y|\mu, \frac{\sigma^2}{\gamma}\right) + (1 - \nu) \phi(y|\mu, \sigma^2) \quad y \in \mathbb{R}. \quad (2.23)$$

- **Distribuição exponencial potência**, com parâmetro $0,5 < \nu \leq 1$, denotada por $Y \sim EP(\mu, \sigma^2; \nu)$, tem fdp dada por

$$f(y) = \frac{\nu}{2^{\frac{1}{2\nu}} \sigma \Gamma(\frac{1}{2\nu})} e^{-d^\nu/2}, \quad y \in \mathbb{R}, \quad (2.24)$$

neste caso, $\kappa(u)$ não possui forma explícita e U tem densidade positiva estável $S^p(u|\nu)$ (Branco e Dey, 2001).

Modelo de Calibração Linear

Inicialmente é apresentado o modelo de calibração linear simples, que é bastante usual, em seguida é apresentado o modelo proposto neste trabalho, o modelo de calibração linear com *Misturas de Escala Normal Assimétrica*. Estas abordagens são úteis em situações em que os dados apresentam uma relação linear.

3.1 Modelo de calibração linear simples

O modelo de calibração linear simples é constituído por dois estágios, a saber

$$Y_i = \alpha + \beta X_i + \epsilon_i, \quad i = 1, \dots, n, \quad (3.1)$$

$$Y_{0i} = \alpha + \beta x_0 + \epsilon_i, \quad i = n + 1, \dots, n + m, \quad (3.2)$$

assumindo que os erros (ϵ_i) são independentes e indenticamente distribuídos (iid) segundo a distribuição $N(0, \sigma^2)$.

Os valores das variáveis aleatórias Y_i e Y_{0i} são observados, os valores X_i são prefixados e os parâmetros α , β , x_0 e σ^2 são desconhecidos, sendo que o interesse principal é estimar o valor de x_0 .

Na literatura é bastante discutido dois tipos de estimadores para o modelo de calibração linear simples (3.1) e (3.2) (veja Shukla (1972), Krutchkoff (1967) e Eisenhart (1939)). Os estimadores *clássico* e *inverso* são os mais utilizados, pois os cálculos são simples e os estimadores dos parâmetros possuem forma fechada.

A seguir serão apresentados estes dois tipos de estimadores para este modelo (Shukla, 1972).

3.1.1 Estimadores *clássico* e *inverso*

Obtém-se por meio do método de máxima verossimilhança, levando-se em consideração os dois estágios do modelo de calibração dado em (3.1) e (3.2) os seguintes estimadores

$$\hat{\alpha} = \bar{Y} - \hat{\beta}\bar{X}, \quad (3.3)$$

$$\hat{\beta} = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sum_{i=1}^n (X_i - \bar{X})^2}, \quad (3.4)$$

$$\widehat{x_{0C}} = \frac{\bar{Y}_0 - \hat{\alpha}}{\hat{\beta}}, \quad \hat{\beta} \neq 0, \quad (3.5)$$

$\widehat{x_{0C}}$ é chamado de estimador *clássico*

sendo

$$\overline{X} = \frac{\sum_{i=1}^n X_i}{n} \quad (3.6)$$

$$\overline{Y} = \frac{\sum_{i=1}^n Y_i}{n} \quad (3.7)$$

$$\overline{Y_0} = \frac{\sum_{i=n+1}^{n+m} Y_{0i}}{m}. \quad (3.8)$$

Pode-se notar que α e β dependem apenas dos dados do primeiro e segundo estgios.

No mtodo *inverso*, como o prprio nome diz, o estimador do parmetro x_0  obtido considerando a regresso de x em Y , apesar de x no ser varivel aleatria. Assim, um problema de calibrao linear  transformado em um problema de regresso linear.

Podemos reescrever o modelo (3.1) como

$$X_i = \gamma + \phi Y_i + \epsilon'_i, \quad i = 1, \dots, n, \quad (3.9)$$

sendo $\epsilon'_i = -\epsilon_i \phi$, $\gamma = -\alpha/\beta$ e $\phi = 1/\beta$.

O estimador de γ e ϕ , obtidos pelo mtodo de mnimos quadrados, e o estimador *inverso* de x_0 so dados, respectivamente, por

$$\widehat{\gamma} = \overline{X} + \widehat{\phi} \overline{Y}, \quad (3.10)$$

$$\widehat{\phi} = \frac{\sum_{i=1}^n (X_i - \overline{X})(Y_i - \overline{Y})}{\sum_{i=1}^n (Y_i - \overline{Y})^2}, \quad (3.11)$$

$$\widehat{x_{0I}} = \widehat{\gamma} + \widehat{\phi} \overline{Y_0}. \quad (3.12)$$

$\widehat{x_{0I}}$  chamado de estimador *inverso*

Observa-se que $\hat{\gamma} = -\hat{\alpha}/\hat{\beta}$ e $\hat{\phi} = 1/\hat{\beta}$ e que $\hat{\alpha}$, $\hat{\beta}$, $\bar{X}$, $\bar{Y}$ e $\bar{Y}_0$ são obtidos por (3.3), (3.4), (3.6), (3.7) e (3.8), respectivamente.

Ambos estimadores apresentam vantagens e desvantagens e são amplamente discutidos na literatura. Segundo Brereton (2003), embora o estimador *clássico* seja o mais utilizado, não é sempre o mais apropriado em abordagem química, por duas razões principais. Em primeiro lugar, o objetivo final é, geralmente, prever a concentração ou uma variável independente (x_0) a partir do espectro ou cromatograma que é a variável resposta, em vez de vice-versa. O segundo diz respeito à distribuição dos erros, como é mostrado na figura a seguir

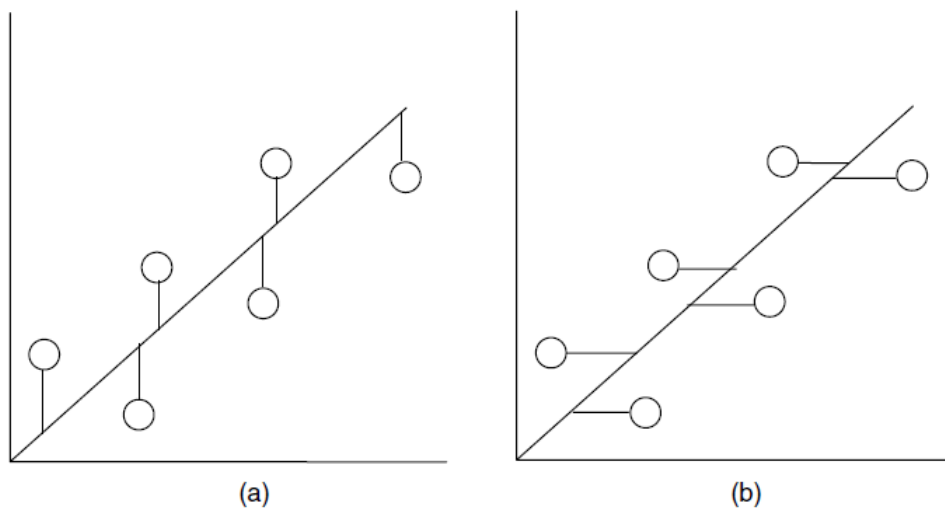


Figura 3.1: Diferença entre os erros na calibração *clássica* (a) e *inversa* (b)

Os erros na variável resposta são obtidos, muitas vezes, devido ao desempenho de instrumentos que, com o avanço da tecnologia, possuem relativa precisão. Por outro lado, a variável independente é, geralmente, determinada por pesagens ou diluições manuais, sendo a maior fonte de erros. No modelo de calibração *clássico* considera-se que todos os erros estão na resposta [Figura 3.1 (a)], ao passo que uma hipótese mais apropriada é que os erros estão, principalmente, na medição da concentração [Figura 3.1 (b)].

Shukla (1972) constatou que o número de observações no experimento de calibração influencia no erro quadrático médio. Para um número de observações pequeno, o estimador *inverso* produz um erro quadrático médio menor que o do estimador *clássico*. Porém, se o número de observações for grande na segunda etapa, não há garantias que o estimador *inverso* seja melhor que o *clássico*. Assim, é aconselhável um estimador consistente para grandes amostras, sendo o estimador *clássico* adequado em tal situação.

O modelo proposto neste trabalho segue a abordagem do estimador *clássico*, como é verificado a seguir.

3.2 Modelo de calibração linear *MENA*

Os modelos de calibração presentes na literatura, em grande parte, supõem que os erros são normalmente distribuídos, no entanto, a distribuição *normal* é inadequada para dados com observações destoantes ou atípicas e assimetria. O uso da família de distribuições *MENA*, que é uma generalização de várias distribuições simétricas e assimétricas, inclusive das distribuições *normal* e *normal assimétrica*, pode ser uma solução para estes problemas. A seguir é apresentado um novo modelo de calibração linear no qual é utilizada a classe ou família de distribuições denominada *MENA* ou *Misturas de Escala Normal Assimétrica* (Ferreira *et al.*, 2011) para a distribuição dos erros.

O modelo de calibração linear proposto é constituído por dois estágios, a saber

$$y_{ij} = \alpha + \beta x_i + \epsilon_{ij}, \quad i = 1, 2, \dots, n \text{ e } j = 1, 2, \dots, r_i, \quad (3.13)$$

$$y_{0i} = \alpha + \beta x_0 + \epsilon_i, \quad i = n+1, n+2, \dots, n+m, \quad (3.14)$$

assumindo que os erros são iid segundo uma distribuição pertencente à família de distribuições $MENA(\sigma^2, \lambda; H)$, isto é, $\epsilon_{ij}, \epsilon_i \stackrel{\text{iid}}{\sim} MENA(\sigma^2, \lambda; H)$.

Assim como no modelo de calibração linear simples, os valores das variáveis aleatórias y_{ij} e y_{0i} são observados, os valores x_i são prefixados e os parâmetros $\alpha, \beta, \sigma^2, \lambda$ e x_0 são desconhecidos, sendo que o interesse principal é estimar o valor de x_0 .

Para $\epsilon_{ij}, \epsilon_i \stackrel{\text{iid}}{\sim} MENA(\sigma^2, \lambda; H)$, tem-se que $\mathbb{E}(\epsilon_{ij}) = \sqrt{\frac{2}{\pi}} \lambda \mathbb{E}_U \left[\frac{\kappa(U)}{\sqrt{1+\lambda^2 \kappa(U)}} \right] \neq 0$ e $\mathbb{E}(\epsilon_i) = \sqrt{\frac{2}{\pi}} \lambda \mathbb{E}_{U_0} \left[\frac{\kappa(U_0)}{\sqrt{1+\lambda^2 \kappa(U_0)}} \right] \neq 0$, para $\lambda \neq 0$.

Seja $(R \times 1)$ a dimensão dos vetores $\mathbf{y}$ e $\mathbf{x}$, o modelo proposto considera situações em que há repetição na variável resposta, sendo r_i o número de repetições. Se os dados não possuem repetição na variável resposta, então $r_i = 1$ e $R = n$. Para o caso em que o número de repetições é constante ($r_i = r$), tem-se que $R = nr$. Porém, não é necessário que o número de repetições seja constante, neste caso tem-se $R = \sum_{i=1}^n r_i$, sendo esta uma definição geral para R .

A função de log-verossimilhança de dados observados do vetor de parâmetros $\boldsymbol{\theta} = (\alpha, \beta, \sigma^2, \lambda, x_0)$ do modelo de calibração proposto, de acordo com (2.3), é dada por

$$\begin{aligned}
\ell(\boldsymbol{\theta}) = & \sum_{i=1}^n \sum_{j=1}^{r_i} \log \int_0^\infty \int_0^\infty 2\phi(y_{ij}|\alpha + \beta x_i, \sigma^2 \kappa(u_{ij})) \phi(t_{ij}|\lambda(y_{ij} - \alpha - \beta x_i), \sigma^2) \times \\
& h(u_{ij}; \tau) dt_{ij} du_{ij} + \\
& \sum_{i=n+1}^{n+m} \log \int_0^\infty \int_0^\infty 2\phi(y_{0i}|\alpha + \beta x_0, \sigma^2 \kappa(u_{0i})) \phi(t_{0i}|\lambda(y_{0i} - \alpha - \beta x_0), \sigma^2) \times \\
& h(u_{0i}; \tau) dt_{0i} du_{0i}.
\end{aligned} \tag{3.15}$$

Observa-se que não é uma tarefa fácil obter as estimativas de máxima verossimilhança do vetor de parâmetros $\boldsymbol{\theta}$ maximizando diretamente a função de log-verossimilhança. Neste caso é preferível uma solução numérica via *algoritmo EM*.

3.2.1 Algoritmo EM

O *algoritmo EM* tem sua origem no trabalho realizado por Hartley (1958), porém após o trabalho de Dempster *et al.* (1977), que detalhou sua estrutura básica e ilustrou seu uso em ampla variedade de aplicações, o *algoritmo EM* entrou efetivamente em proeminência estatística.

Trata-se de uma ferramenta computacional utilizada para o cálculo do estimador de máxima verossimilhança (EMV) de forma iterativa, sendo principalmente utilizado em problemas envolvendo dados incompletos. Para isso é necessário obter o conjunto dos dados completos, que é o conjunto dos dados observados aumentado com o conjunto dos dados faltantes e assim obter a função de log-verossimilhança associada aos dados completos.

A utilização mais comum do *algoritmo EM* é computacionalmente simples é quando a maximização da função de verossimilhança é analiticamente problemática, mas a função de probabilidade pode ser simplificada, admitindo a existência de valores adicionais, ou seja, a utilização de *modelos hierárquicos* que é uma forma de modelar processos complicados como uma sequência de modelos relativamente simples. Segundo Casella e Berger (2011), o *EM* é um algoritmo que certamente converge para o EMV e tem como base a ideia de substituir uma difícil maximização da função de verossimilhança por uma sequência de maximizações mais fáceis, cujo limite é a resposta para o problema original.

Cada iteração do *algoritmo EM* envolve dois passos, um passo E (Esperança) e um passo M (Maximização), definidos como

- **Passo E:** Calcule $Q(\boldsymbol{\theta}|\hat{\boldsymbol{\theta}}) = \mathbb{E} \left[\ell_c(\boldsymbol{\theta}|\mathbf{y}_c) | \mathbf{y}, \mathbf{y}_0, \hat{\boldsymbol{\theta}}^{(p)} \right]$, onde a esperança é tomada com respeito a distribuições condicionais.
- **Passo M:** Encontre $\hat{\boldsymbol{\theta}}^{(p+1)}$, que maximiza $Q(\boldsymbol{\theta}|\hat{\boldsymbol{\theta}})$.

Estes passos devem ser repetidos até se atingir uma convergência. Pode ser adotado como critério de parada, por exemplo, $|\boldsymbol{\theta}^{(p+1)} - \boldsymbol{\theta}^{(p)}| < \epsilon^*$ onde ϵ^* é um valor determinado e maior que zero.

3.2.2 Algoritmo EM para o modelo proposto

De acordo com (2.11), o modelo (3.13) e (3.14) pode ser reescrito hierarquicamente como

$$\begin{aligned} Y_{ij}|T_{ij} = t_{ij}, U_{ij} = u_{ij} &\stackrel{\text{iid}}{\sim} N\left(\alpha + \beta x_i + t_{ij} \frac{\sigma \lambda \kappa(u_{ij})}{\sqrt{1 + \lambda^2 \kappa(u_{ij})}}, \frac{\sigma^2 \kappa(u_{ij})}{1 + \lambda^2 \kappa(u_{ij})}\right), \\ U_{ij} &\stackrel{\text{iid}}{\sim} H(u_{ij}; \boldsymbol{\tau}), \\ T_{ij} &\stackrel{\text{iid}}{\sim} NT(0, 1), \quad i = 1, \dots, n \text{ e } j = 1, \dots, r \end{aligned} \quad (3.16)$$

e

$$\begin{aligned} Y_{0i}|T_{0i} = t_{0i}, U_{0i} = u_{0i} &\stackrel{\text{iid}}{\sim} N\left(\alpha + \beta x_0 + t_{0i} \frac{\sigma \lambda \kappa(u_{0i})}{\sqrt{1 + \lambda^2 \kappa(u_{0i})}}, \frac{\sigma^2 \kappa(u_{0i})}{1 + \lambda^2 \kappa(u_{0i})}\right), \\ U_{0i} &\stackrel{\text{iid}}{\sim} H(u_{0i}; \boldsymbol{\tau}), \\ T_{0i} &\stackrel{\text{iid}}{\sim} NT(0, 1), \quad i = n + 1, \dots, n + m, \end{aligned} \quad (3.17)$$

onde o $NT(0, 1)$ denota a distribuição *normal truncada padrão* e assumindo-se $\boldsymbol{\tau}$ conhecido.

Sejam $\mathbf{y}$, $\mathbf{u}$ e $\mathbf{t}$ vetores de dimensão $(R \times 1)$, em que $R = \sum_{i=1}^n r_i$ e $\mathbf{y}_0$, $\mathbf{u}_0$ e $\mathbf{t}_0$ vetores de dimensão $(m \times 1)$. Tratando $\mathbf{u}$, $\mathbf{t}$, $\mathbf{u}_0$ e $\mathbf{t}_0$ como faltantes, segue que a função de log-verossimilhança completa associada a $\mathbf{y}_c = (\mathbf{y}^\top, \mathbf{y}_0^\top, \mathbf{u}^\top, \mathbf{u}_0^\top, \mathbf{t}^\top, \mathbf{t}_0^\top)^\top$ é dada por

$$\begin{aligned} \ell_c(\boldsymbol{\theta}|\mathbf{y}_c) &\propto -R \log \sigma^2 - m \log \sigma^2 \\ &\quad - \frac{1}{2\sigma^2} \sum_{i=1}^n \sum_{j=1}^{r_i} \frac{(y_{ij} - \alpha - \beta x_i)^2}{\kappa(u_{ij})} - \frac{1}{2\sigma^2} \sum_{i=1}^n \sum_{j=1}^{r_i} (t_{ij} - \lambda(y_{ij} - \alpha - \beta x_i))^2 \\ &\quad - \frac{1}{2\sigma^2} \sum_{i=n+1}^{n+m} \frac{(y_{0i} - \alpha - \beta x_0)^2}{\kappa(u_{0i})} - \frac{1}{2\sigma^2} \sum_{i=n+1}^{n+m} (t_{0i} - \lambda(y_{0i} - \alpha - \beta x_0))^2. \end{aligned} \quad (3.18)$$

Passo E

Dado que, no primeiro estágio tem-se $T_{ij}|y_{ij} \sim NT(\lambda(y_{ij} - \alpha - \beta x_i), \sigma^2)\mathbb{I}_{(0,\infty)}(t_{ij})$ e sendo $\widehat{t_{ij}} = \mathbb{E}(T_{ij}|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{ij})$ e $\widehat{t_{ij}^2} = \mathbb{E}(T_{ij}^2|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{ij})$, obtém-se, usando os momentos da distribuição *normal truncada* como em (2.9) e (2.10),

$$\widehat{t_{ij}} = \widehat{\lambda}\widehat{\eta_{ij}} + \widehat{\sigma}W_{\Phi}\left(\frac{\widehat{\lambda}\widehat{\eta_{ij}}}{\widehat{\sigma}}\right) \quad \text{e} \quad (3.19)$$

$$\widehat{t_{ij}^2} = \widehat{\lambda}^2\widehat{\eta_{ij}}^2 + \widehat{\sigma}^2 + \widehat{\lambda}\widehat{\sigma}\widehat{\eta_{ij}}W_{\Phi}\left(\frac{\widehat{\lambda}\widehat{\eta_{ij}}}{\widehat{\sigma}}\right), \quad (3.20)$$

sendo $\widehat{\eta_{ij}} = y_{ij} - \widehat{\alpha} - \widehat{\beta}x_i$ e $W_{\Phi}(a) = \phi(a)/\Phi(a)$.

Do mesmo modo, para o segundo estágio, tem-se $T_{0i}|y_{0i} \sim NT(\lambda(y_{0i} - \alpha - \beta x_0), \sigma^2)\mathbb{I}_{(0,\infty)}(t_{0i})$ e sendo $\widehat{t_{0i}} = \mathbb{E}(T_{0i}|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{0i})$ e $\widehat{t_{0i}^2} = \mathbb{E}(T_{0i}^2|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{0i})$, obtém-se, usando também os momentos da distribuição *normal truncada* como em (2.9) e (2.10),

$$\widehat{t_{0i}} = \widehat{\lambda}\widehat{\eta_{0i}} + \widehat{\sigma}W_{\Phi}\left(\frac{\widehat{\lambda}\widehat{\eta_{0i}}}{\widehat{\sigma}}\right) \quad \text{e} \quad (3.21)$$

$$\widehat{t_{0i}^2} = \widehat{\lambda}^2\widehat{\eta_{0i}}^2 + \widehat{\sigma}^2 + \widehat{\lambda}\widehat{\sigma}\widehat{\eta_{0i}}W_{\Phi}\left(\frac{\widehat{\lambda}\widehat{\eta_{0i}}}{\widehat{\sigma}}\right), \quad (3.22)$$

sendo $\widehat{\eta_{0i}} = y_{0i} - \widehat{\alpha} - \widehat{\beta}x_0$ e $W_{\Phi}(a) = \phi(a)/\Phi(a)$.

Denota-se por $\boldsymbol{\theta}^{(p)} = (\alpha^{(p)}, \beta^{(p)}, \sigma^{2(p)}, \lambda^{(p)}, x_0^{(p)})$ a estimativa de $\boldsymbol{\theta}$ na p -ésima iteração do algoritmo EM.

Fazendo as devidas substituições em (3.18), sendo

$$\widehat{\kappa_{ij}} = \widehat{\kappa(u_{ij})} = \mathbb{E}(\kappa^{-1}(U_{ij})|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{ij}) \text{ e} \quad (3.23)$$

$$\widehat{\kappa_{0i}} = \widehat{\kappa(u_{0i})} = \mathbb{E}(\kappa^{-1}(U_{0i})|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{0i}), \quad (3.24)$$

segue que a esperança com respeito à $\mathbf{u}$, $\mathbf{t}$, $\mathbf{u}_0$ e $\mathbf{t}_0$, condicionada em $\mathbf{y}$ e em $\mathbf{y}_0$, da função de log-verossimilhança completa tem a forma

$$\begin{aligned} Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}}) &= \mathbb{E} \left[\ell_c(\boldsymbol{\theta}|\mathbf{y}_c)|\mathbf{y}, \mathbf{y}_0, \widehat{\boldsymbol{\theta}}^{(p)} \right] = -R \log \sigma^{2(p)} - m \log \sigma^{2(p)} \\ &\quad - \frac{1}{2\sigma^{2(p)}} \sum_{i=1}^n \sum_{j=1}^{r_i} (y_{ij} - \alpha^{(p)} - \beta^{(p)} x_i)^2 (\widehat{\kappa_{ij}}^{(p)} + \lambda^{(p)^2}) \\ &\quad - \frac{1}{2\sigma^{2(p)}} \sum_{i=1}^n \sum_{j=1}^{r_i} \widehat{t_{ij}^2}^{(p)} + \frac{\lambda^{(p)}}{\sigma^{2(p)}} \sum_{i=1}^n \sum_{j=1}^{r_i} (y_{ij} - \alpha^{(p)} - \beta^{(p)} x_i) \widehat{t_{ij}}^{(p)} \\ &\quad - \frac{1}{2\sigma^{2(p)}} \sum_{i=n+1}^{n+m} (y_{0i} - \alpha^{(p)} - \beta^{(p)} x_0^{(p)})^2 (\widehat{\kappa_{0i}}^{(p)} + \lambda^{(p)^2}) \\ &\quad - \frac{1}{2\sigma^{2(p)}} \sum_{i=n+1}^{n+m} \widehat{t_{0i}^2}^{(p)} + \frac{\lambda^{(p)}}{\sigma^{2(p)}} \sum_{i=n+1}^{n+m} (y_{0i} - \alpha^{(p)} - \beta^{(p)} x_0^{(p)}) \widehat{t_{0i}}^{(p)}. \quad (3.25) \end{aligned}$$

Passo M

Maximizando a função $Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})$ em relação a $\boldsymbol{\theta}^{(p)}$, obtemos as expressões a seguir. Considere $\boldsymbol{\eta}^{(p)} = \mathbf{y} - \alpha^{(p)} \mathbf{1}_R - \beta^{(p)} \mathbf{x}$, $\boldsymbol{\eta}_0^{(p)} = \mathbf{y}_0 - \left(\alpha^{(p)} + \beta^{(p)} x_0^{(p)} \right) \mathbf{1}_m$, $\mathbf{D}(\mathbf{A}) = \text{Diag}(a_1, a_2, \dots)$, $\mathbf{1}_n$ é um vetor de uns de dimensão $(n \times 1)$ e $\mathbf{I}_R$ é a matriz identidade de ordem R .

$$\begin{aligned}
\widehat{\alpha}^{(p+1)} &= \left[\widehat{\kappa}^{(p)\top} \mathbf{1}_R + \widehat{\kappa}_0^{(p)\top} \mathbf{1}_m + (R+m)\lambda^{(p)2} \right]^{-1} \\
&\quad \left[\left(\mathbf{y}^\top - \beta^{(p)} \mathbf{x}^\top \right) \widehat{\kappa}^{(p)} + \lambda^{(p)} \left(\lambda^{(p)} \mathbf{y}^\top - \widehat{\mathbf{t}}^{(p)\top} - \lambda^{(p)} \beta^{(p)} \mathbf{x}^\top \right) \mathbf{1}_R + \mathbf{y}_0^\top \widehat{\kappa}_0^{(p)} \right. \\
&\quad \left. + \left(\lambda^{(p)2} \mathbf{y}_0^\top - \lambda^{(p)} \widehat{\mathbf{t}}_0^{(p)\top} - \beta^{(p)} x_0^{(p)} \widehat{\kappa}_0^{(p)\top} \right) \mathbf{1}_m - m \beta^{(p)} x_0^{(p)} \lambda^{(p)2} \right] \quad (3.26)
\end{aligned}$$

$$\begin{aligned}
\widehat{\beta}^{(p+1)} &= \left[\mathbf{x}^\top \left(\mathbf{D}(\widehat{\kappa}^{(p)}) + \lambda^{(p)2} \mathbf{I}_R \right) \mathbf{x} + x_0^{(p)2} \left(\widehat{\kappa}_0^{(p)\top} \mathbf{1}_m + \lambda^{(p)2} m \right) \right]^{-1} \\
&\quad \left\{ \mathbf{x}^\top \left[\mathbf{D}(\widehat{\kappa}^{(p)}) \mathbf{y} - \alpha^{(p)} \widehat{\kappa}^{(p)} + \lambda^{(p)2} \left(\mathbf{y} - \alpha^{(p)} \mathbf{1}_R \right) - \lambda^{(p)} \widehat{\mathbf{t}}^{(p)} \right] + x_0^{(p)} \times \right. \\
&\quad \left. \left[\mathbf{y}_0^\top \widehat{\kappa}_0^{(p)} + \left(\lambda^{(p)2} \mathbf{y}_0^\top - \alpha^{(p)} \widehat{\kappa}_0^{(p)\top} - \lambda^{(p)} \widehat{\mathbf{t}}_0^{(p)\top} \right) \mathbf{1}_m - m \lambda^{(p)2} \alpha^{(p)} \right] \right\} \quad (3.27)
\end{aligned}$$

$$\begin{aligned}
\widehat{\sigma^2}^{(p+1)} &= [2(R+m)]^{-1} \left[\left(\boldsymbol{\eta}^{(p)\top} \mathbf{D}(\widehat{\kappa}^{(p)}) + \lambda^{(p)2} \boldsymbol{\eta}^{(p)\top} - 2\lambda^{(p)} \widehat{\mathbf{t}}^{(p)\top} \right) \boldsymbol{\eta}^{(p)} + \widehat{\mathbf{t}}_0^{(p)\top} \mathbf{1}_R \right. \\
&\quad \left. + \left(\boldsymbol{\eta}_0^{(p)\top} \mathbf{D}(\widehat{\kappa}_0^{(p)}) + \lambda^{(p)2} \boldsymbol{\eta}_0^{(p)\top} - 2\lambda^{(p)} \widehat{\mathbf{t}}_0^{(p)\top} \right) \boldsymbol{\eta}_0^{(p)} + \widehat{\mathbf{t}}_0^{(p)\top} \mathbf{1}_m \right] \quad (3.28)
\end{aligned}$$

$$\widehat{\lambda}^{(p+1)} = \left[\boldsymbol{\eta}^{(p)\top} \boldsymbol{\eta}^{(p)} + \boldsymbol{\eta}_0^{(p)\top} \boldsymbol{\eta}_0^{(p)} \right]^{-1} \left[\widehat{\mathbf{t}}^{(p)\top} \boldsymbol{\eta}^{(p)} + \widehat{\mathbf{t}}_0^{(p)\top} \boldsymbol{\eta}_0^{(p)} \right] \quad (3.29)$$

$$\begin{aligned}
\widehat{x}_0^{(p+1)} &= \left[\beta^{(p)} \left(\widehat{\kappa}_0^{(p)\top} \mathbf{1}_m + m \lambda^{(p)2} \right) \right]^{-1} \\
&\quad \left[\mathbf{y}_0^\top \widehat{\kappa}_0^{(p)} - \left(\alpha^{(p)} \widehat{\kappa}_0^{(p)\top} + \lambda^{(p)} \widehat{\mathbf{t}}_0^{(p)\top} \right) \mathbf{1}_m + \lambda^{(p)2} \left(\mathbf{y}_0^\top \mathbf{1}_m - m \alpha^{(p)} \right) \right] \quad (3.30)
\end{aligned}$$

Observa-se que os estimadores de α e β dependem dos dados do primeiro e do segundo estgios.

Distribuições condicionais para o *algoritmo EM*

A seguir serão apresentadas as esperanças condicionais $\mathbb{E}(U|y)$ para as distribuições apresentadas no Capítulo 2 que serão utilizadas no *algoritmo EM*.

Proposição 3.1. *Se $Y \sim MENA(\mu, \sigma^2, \lambda; H)$, então a distribuição condicional $U|Y = y$ não depende de λ .*

Demonstração.

$$\begin{aligned} f(u|y, \boldsymbol{\theta}) &= \frac{f(u, y|\boldsymbol{\theta})}{f(y|\boldsymbol{\theta})} \\ &= \frac{2\phi(y|\mu, \sigma^2\kappa(u))\Phi(\lambda(y-\mu)/\sigma)h(u;\boldsymbol{\tau})}{2\int_0^\infty \phi(y|\mu, \sigma^2\kappa(u))\Phi(\lambda(y-\mu)/\sigma)h(u;\boldsymbol{\tau})du} \\ &= \frac{\phi(y|\mu, \sigma^2\kappa(u))h(u;\boldsymbol{\tau})}{\int_0^\infty \phi(y|\mu, \sigma^2\kappa(u))h(u;\boldsymbol{\tau})du} \end{aligned}$$

Segue, da Proposição 3.1, que a distribuição $U|Y$ do primeiro estágio se reduz à distribuição *MEN*. Essa peculiaridade simplifica bastante a implementação do *algoritmo EM*. Analogamente, para o segundo estágio, $U_0|Y_0$ também possui distribuição *MEN*.

No passo E do *algoritmo EM* é necessário o cálculo de $\widehat{\kappa}_{ij} = \mathbb{E}(\kappa^{-1}(U_{ij})|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{ij})$ e $\widehat{\kappa}_{0i} = \mathbb{E}(\kappa^{-1}(U_{0i})|\boldsymbol{\theta} = \widehat{\boldsymbol{\theta}}, y_{0i})$. Uma forma geral, segundo Salgado (2006), para o cálculo destas esperanças condicionais é dada por

$$\mathbb{E}(\kappa^{-1}(u)|\mathbf{Y}) = \frac{\int_0^\infty \kappa^{-3/2}(u)\exp\{-\kappa^{-1}(u)d/2\}dH(u)}{\int_0^\infty \kappa^{-1/2}(u)\exp\{-\kappa^{-1}(u)d/2\}dH(u)} \quad \text{e} \quad (3.31)$$

$$\mathbb{E}(\kappa^{-1}(u_0)|\mathbf{Y}_0) = \frac{\int_0^\infty \kappa^{-3/2}(u_0)\exp\{-\kappa_0^{-1}(u_0)d_0/2\}dH(u_0)}{\int_0^\infty \kappa^{-1/2}(u_0)\exp\{-\kappa_0^{-1}(u_0)d_0/2\}dH(u_0)}. \quad (3.32)$$

Para as distribuições discutidas a seguir, dado $\boldsymbol{\theta} = (\alpha, \beta, \sigma^2, \lambda, x_0)$, $d = (y - \alpha - \beta x)^2 / \sigma^2$ e $d_0 = (y_0 - \alpha - \beta x_0)^2 / \sigma^2$, as expressões das esperanças condicionais (3.31) e (3.32) possuem forma fechada, como segue

- **Distribuição normal assimétrica.** Se $Y \sim NA(\mu, \sigma^2, \lambda)$, então

$$\widehat{\kappa} = 1 \quad (3.33)$$

$$\widehat{\kappa}_0 = 1 \quad (3.34)$$

- **Distribuição *t* de Studente generalizada normal assimétrica.** Se $Y \sim GtA(\mu, \sigma^2, \lambda; \nu, \gamma)$, então

$$\begin{aligned} f(u|y, \boldsymbol{\theta}) &= \frac{(\gamma/2)^{\nu/2}}{\Gamma(\nu/2)} u^{\nu/2-1} e^{-\gamma u/2} \frac{u^{1/2}}{\sqrt{2\pi}\sigma} e^{-ud/2} \\ U|Y=y &\sim Gama\left(\frac{\nu+1}{2}, \frac{\gamma+d}{2}\right). \end{aligned}$$

$$\widehat{\kappa} = \frac{\nu+1}{\gamma+d} \quad (3.35)$$

$$\widehat{\kappa}_0 = \frac{\nu+1}{\gamma+d_0} \quad (3.36)$$

- **Distribuição slash assimétrica.** Se $Y \sim SLA(\mu, \sigma^2, \lambda; \nu)$, obtem-se

$$\begin{aligned} f(u|y, \boldsymbol{\theta}) &= \nu u^{\nu-1} \mathbb{I}_{(0,1)}(u) \frac{u^{1/2}}{\sqrt{2\pi}\sigma} e^{-ud/2} \\ U|Y=y &\sim Gama(\nu+1/2, d/2) \mathbb{I}_{(0,1)}(u). \end{aligned}$$

$$\widehat{\kappa} = \frac{(2\nu + 1)}{d} \frac{\mathbb{P}_1(\nu + 3/2, d/2)}{\mathbb{P}_1(\nu + 1/2, d/2)} \quad (3.37)$$

$$\widehat{\kappa}_0 = \frac{(2\nu + 1)}{d_0} \frac{\mathbb{P}_1(\nu + 3/2, d_0/2)}{\mathbb{P}_1(\nu + 1/2, d_0/2)} \quad (3.38)$$

onde $\mathbb{P}_1(a, b)$ é a fda de uma variável aleatória $Gama(a, b)$ avaliada em 1.

- **Distribuição normal contaminada assimétrica.** Se $Y \sim NCA(\mu, \sigma^2, \lambda; \nu, \gamma)$, tem-se que

$$\begin{aligned} f(u|y, \boldsymbol{\theta}) &= \nu p \mathbb{I}_{(u=\gamma)} + (1 - \nu) p \mathbb{I}_{(u=1)}, \text{ com} \\ p &= \frac{u^{1/2} \exp(-\frac{du}{2})}{\nu \gamma^{1/2} \exp(-\frac{d\gamma}{2}) + (1 - \nu) \exp(-\frac{d}{2})}. \end{aligned}$$

$$\widehat{\kappa} = \frac{1 - \nu + \nu \gamma^{3/2} \exp[(1 - \gamma)d/2]}{1 - \nu + \nu \gamma^{1/2} \exp[(1 - \gamma)d/2]} \quad (3.39)$$

$$\widehat{\kappa}_0 = \frac{1 - \nu + \nu \gamma^{3/2} \exp[(1 - \gamma)d_0/2]}{1 - \nu + \nu \gamma^{1/2} \exp[(1 - \gamma)d_0/2]} \quad (3.40)$$

- **Distribuição exponencial potência assimétrica.** Se $Y \sim EPA(\mu, \sigma^2, \lambda; \nu)$, tem-se

$$f(u|y, \boldsymbol{\theta}) \propto \exp \left[\frac{1}{2}(-du + d^\nu) + c(u, \nu) \right].$$

$$\widehat{\kappa} = \nu d^{\nu-1} \quad (3.41)$$

$$\widehat{\kappa}_0 = \nu d_0^{\nu-1} \quad (3.42)$$

Implementação do *algoritmo EM*

A seguir tem-se um resumo dos passos do processo iterativo do *algoritmo EM*

1. Atribua valores iniciais $\alpha^{(0)}$, $\beta^{(0)}$, $\sigma^{2(0)}$, $\lambda^{(0)}$ e $x_0^{(0)}$, para o vetor de parâmetros $\boldsymbol{\theta}$. Sugere-se os valores obtidos por meio do estimador *clássico*,
2. Calcule as esperanças condicionais $\widehat{t}_{ij}^{(p)}$, $\widehat{t}_{ij}^{2(p)}$, $\widehat{t}_{0i}^{(p)}$, $\widehat{t}_{0i}^{2(p)}$ usando (3.19) a (3.22), e calcule $\widehat{\kappa}_{ij}^{(p)}$ e $\widehat{\kappa}_{0i}^{(p)}$ usando (3.33) a (3.42),
3. Obtenha novos valores para os estimadores $\widehat{\alpha}$, $\widehat{\beta}$, $\widehat{\sigma^2}$, $\widehat{\lambda}$ e $\widehat{x_0}$ usando (3.26) a (3.30),
4. Repita os passos 2 e 3 até a convergência.

Os valores obtidos após a convergência para os estimadores $\widehat{\alpha}$, $\widehat{\beta}$, $\widehat{\sigma^2}$, $\widehat{\lambda}$ e $\widehat{x_0}$ são as estimativas de máxima verossimilhança do modelo de calibração linear proposto.

3.2.3 Matriz de informação de Fisher observada

Para se obter a variância assintótica do EMV do vetor de parâmetros $\boldsymbol{\theta}$ será utilizada a matriz de informação de Fisher observada ($I_F(\boldsymbol{\theta})$) por satisfazer, sob condições de regularidade, a desigualdade $\text{var}(\widehat{\boldsymbol{\theta}}) \geq I_F^{-1}(\boldsymbol{\theta})$, ou seja, o EMV de $\boldsymbol{\theta}$ atinge assintoticamente o Limite Inferior de Cramer-Rao, como é discutido amplamente na literatura.

Considerando $Y_{11}, \dots, Y_{1r_1}, \dots, Y_{n1}, \dots, Y_R$ onde $Y_{ij} \sim MENA(\mu_{ij}, \sigma^2, \lambda; H)$, $\mu_{ij} = \alpha + \beta x_i$ e $Y_{01}, \dots, Y_{0m}$ onde $Y_{0i} \sim MENA(\mu_{0i}, \sigma^2, \lambda; H)$, $\mu_{0i} = \alpha + \beta x_0$, então a função de

log-verossimilhança $\ell(\boldsymbol{\theta}) = \sum_{i=1}^{n+m} \sum_{j=1}^{r_i} \ell_{ij}(\boldsymbol{\theta})$, com $\boldsymbol{\theta} = (\alpha, \beta, \sigma^2, \lambda, x_0)$, é da forma

$$\ell_{ij}(\boldsymbol{\theta}) = 2 \log 2 + \ell_{1_{ij}}(\boldsymbol{\theta}) + \log[\Phi(\ell_{2_{ij}}(\boldsymbol{\theta}))] + \ell_{3_i}(\boldsymbol{\theta}) + \log[\Phi(\ell_{4_i}(\boldsymbol{\theta}))], \quad (3.43)$$

onde $\ell_{1_{ij}}(\boldsymbol{\theta})$ e $\ell_{3_i}(\boldsymbol{\theta})$ são funções de log-verossimilhança da distribuição $MEN(\alpha + \beta x_i, \sigma^2; H)$ e $MEN(\alpha + \beta x_0, \sigma^2; H)$, respectivamente, $\ell_{2_{ij}}(\boldsymbol{\theta}) = \lambda \left(\frac{y_{ij} - \alpha - \beta x_i}{\sigma} \right)$ e $\ell_{4_i}(\boldsymbol{\theta}) = \lambda \left(\frac{y_{0i} - \alpha - \beta x_0}{\sigma} \right)$. A primeira derivada de $\ell_{ij}(\boldsymbol{\theta})$ em relação a um dos componentes de $\boldsymbol{\theta}$ é dada por

$$\frac{\partial \ell_{ij}(\boldsymbol{\theta})}{\partial \psi} = \frac{\partial \ell_{1_{ij}}(\boldsymbol{\theta})}{\partial \psi} + W_{\Phi}(\ell_{2_{ij}}(\boldsymbol{\theta})) \frac{\partial \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \psi} + \frac{\partial \ell_{3_i}(\boldsymbol{\theta})}{\partial \psi} + W_{\Phi}(\ell_{4_i}(\boldsymbol{\theta})) \frac{\partial \ell_{4_i}(\boldsymbol{\theta})}{\partial \psi},$$

onde $W_{\Phi}(a) = \phi(a)/\Phi(a)$. A segunda derivada de $\ell_{ij}(\boldsymbol{\theta})$ é dada por

$$\begin{aligned} \frac{\partial^2 \ell_{ij}(\boldsymbol{\theta})}{\partial \psi \partial \omega} &= \frac{\partial^2 \ell_{1_{ij}}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}(\ell_{2_{ij}}(\boldsymbol{\theta})) \frac{\partial^2 \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}^{(1)}(\ell_{2_{ij}}(\boldsymbol{\theta})) \frac{\partial \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \omega} \frac{\partial \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \psi} + \\ &\quad \frac{\partial^2 \ell_{3_i}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}(\ell_{4_i}(\boldsymbol{\theta})) \frac{\partial^2 \ell_{4_i}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}^{(1)}(\ell_{4_i}(\boldsymbol{\theta})) \frac{\partial \ell_{4_i}(\boldsymbol{\theta})}{\partial \omega} \frac{\partial \ell_{4_i}(\boldsymbol{\theta})}{\partial \psi}, \end{aligned}$$

onde $W_{\Phi}^{(1)}(a) = -W_{\Phi}(a)(a + W_{\Phi}(a))$ é a derivada de $W_{\Phi}(a)$ e $\omega, \psi = \alpha, \beta, \sigma^2, \lambda, x_0$.

As expressões dos elementos da matriz de informação de Fisher são dadas por

$$\begin{aligned} I_{\psi\omega}^1 &= - \sum_{i=1}^n \sum_{j=1}^{r_i} \frac{\partial^2 \ell_{1_{ij}}(\boldsymbol{\theta})}{\partial \psi \partial \omega}, \\ I_{\psi\omega}^2 &= - \sum_{i=1}^n \sum_{j=1}^{r_i} \left[W_{\Phi}(\ell_{2_{ij}}(\boldsymbol{\theta})) \frac{\partial^2 \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}^{(1)}(\ell_{2_{ij}}(\boldsymbol{\theta})) \frac{\partial \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \omega} \frac{\partial \ell_{2_{ij}}(\boldsymbol{\theta})}{\partial \psi} \right], \\ I_{\psi\omega}^3 &= - \sum_{i=n+1}^{n+m} \frac{\partial^2 \ell_{3_i}(\boldsymbol{\theta})}{\partial \psi \partial \omega} \text{ e} \\ I_{\psi\omega}^4 &= - \sum_{i=n+1}^{n+m} \left[W_{\Phi}(\ell_{4_i}(\boldsymbol{\theta})) \frac{\partial^2 \ell_{4_i}(\boldsymbol{\theta})}{\partial \psi \partial \omega} + W_{\Phi}^{(1)}(\ell_{4_i}(\boldsymbol{\theta})) \frac{\partial \ell_{4_i}(\boldsymbol{\theta})}{\partial \omega} \frac{\partial \ell_{4_i}(\boldsymbol{\theta})}{\partial \psi} \right]. \end{aligned}$$

Portanto, a matriz de informação de Fisher observada para $\boldsymbol{\theta}$ pode ser escrita como

$$I_F(\boldsymbol{\theta}) = I_1(\boldsymbol{\theta}) + I_2(\boldsymbol{\theta}) + I_3(\boldsymbol{\theta}) + I_4(\boldsymbol{\theta}),$$

$$\text{com } I_k(\boldsymbol{\theta}) = \begin{pmatrix} I_{\alpha\alpha}^k & I_{\beta\alpha}^k & I_{\sigma^2\alpha}^k & I_{\lambda\alpha}^k & I_{x_0\alpha}^k \\ & I_{\beta\beta}^k & I_{\sigma^2\beta}^k & I_{\lambda\beta}^k & I_{x_0\beta}^k \\ & & I_{\sigma^2\sigma^2}^k & I_{\lambda\sigma^2}^k & I_{x_0\sigma^2}^k \\ & & & I_{\lambda\lambda}^k & I_{x_0\lambda}^k \\ & & & & I_{x_0x_0}^k \end{pmatrix},$$

para $k = 1, 2, 3, 4$.

A seguir as expressões dos elementos da matriz de informação de Fisher observada para as distribuições *Normal Assimétrica*, *t de Student Assimétrica*, *Exponencial Potência Assimétrica* e *Slash Assimétrica* são apresentadas. Foram utilizadas as notações $\boldsymbol{\eta} = \mathbf{y} - \alpha \mathbf{1}_R - \beta \mathbf{x}$, $\boldsymbol{\eta}_0 = \mathbf{y}_0 - (\alpha + \beta x_0) \mathbf{1}_m$, $\mathbf{V} = (1 + \frac{d}{\gamma})^{-1}$, $\mathbf{V}_0 = (1 + \frac{d_0}{\gamma})^{-1}$, $\mathbf{D}(\mathbf{A}) = \text{Diag}(a_1, a_2, \dots)$ e $\mathbf{1}_n$ é um vetor de uns de dimensão $(n \times 1)$.

As matrizes $I_2(\boldsymbol{\theta})$ e $I_4(\boldsymbol{\theta})$ são comuns às distribuições consideradas, nas quais $I_{x_0\boldsymbol{\theta}}^2 = 0$, com $\boldsymbol{\theta} = (\alpha, \beta, \sigma^2, \lambda, x_0)$, como segue

$$\begin{aligned}
I_{\alpha\alpha}^2 &= -\frac{\lambda^2}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right)^\top \mathbf{1}_R, \\
I_{\alpha\beta}^2 &= -\frac{\lambda^2}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right)^\top \mathbf{x}, \\
I_{\alpha\sigma^2}^2 &= -\frac{\lambda}{2\sigma^3} W_\Phi \left(\frac{\lambda\eta}{\sigma} \right)^\top \mathbf{1}_R - \frac{\lambda^2}{2\sigma^4} W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right)^\top \eta, \\
I_{\alpha\lambda}^2 &= \frac{1}{\sigma} W_\Phi \left(\frac{\lambda\eta}{\sigma} \right)^\top \mathbf{1}_R + \frac{\lambda}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right)^\top \eta, \\
I_{\beta\beta}^2 &= -\frac{\lambda^2}{\sigma^2} \mathbf{x}^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \mathbf{x}, \\
I_{\beta\sigma^2}^2 &= -\frac{\lambda}{2\sigma^3} \mathbf{x}^\top W_\Phi \left(\frac{\lambda\eta}{\sigma} \right) - \frac{\lambda^2}{2\sigma^4} \mathbf{x}^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \eta, \\
I_{\beta\lambda}^2 &= \frac{1}{\sigma} \mathbf{x}^\top W_\Phi \left(\frac{\lambda\eta}{\sigma} \right) + \frac{\lambda}{\sigma^2} \mathbf{x}^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \eta, \\
I_{\sigma^2\sigma^2}^2 &= -\frac{3\lambda}{4\sigma^5} \eta^\top W_\Phi \left(\frac{\lambda\eta}{\sigma} \right) - \frac{\lambda^2}{4\sigma^6} \eta^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \eta, \\
I_{\sigma^2\lambda}^2 &= \frac{1}{2\sigma^3} \eta^\top W_\Phi \left(\frac{\lambda\eta}{\sigma} \right) + \frac{\lambda}{2\sigma^4} \eta^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \eta, \\
I_{\lambda\lambda}^2 &= -\frac{1}{\sigma^2} \eta^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta}{\sigma} \right) \right] \eta,
\end{aligned}$$

$$\begin{aligned}
I_{\alpha\alpha}^4 &= -\frac{\lambda^2}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m, \\
I_{\alpha\beta}^4 &= -\frac{\lambda^2 x_0}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m, \\
I_{\alpha\sigma^2}^4 &= -\frac{\lambda}{2\sigma^3} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m - \frac{\lambda^2}{2\sigma^4} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{\alpha\lambda}^4 &= \frac{1}{\sigma} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m + \frac{\lambda}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{\alpha x_0}^4 &= -\frac{\lambda^2 \beta}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m, \\
I_{\beta\beta}^4 &= -\frac{\lambda^2 x_0^2}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m, \\
I_{\beta\sigma^2}^4 &= -\frac{\lambda x_0}{2\sigma^3} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m - \frac{\lambda^2 x_0}{2\sigma^4} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{\beta\lambda}^4 &= \frac{x_0}{\sigma} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m + \frac{\lambda x_0}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{\beta x_0}^4 &= \frac{\lambda}{\sigma} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m - \frac{\lambda^2 \beta x_0}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m, \\
I_{\sigma^2\sigma^2}^4 &= -\frac{3\lambda}{4\sigma^5} \eta_0^\top W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right) - \frac{\lambda^2}{4\sigma^6} \eta_0^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right) \right] \eta_0, \\
I_{\sigma^2\lambda}^4 &= \frac{1}{2\sigma^3} \eta_0^\top W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right) + \frac{\lambda}{2\sigma^4} \eta_0^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right) \right] \eta_0, \\
I_{\sigma^2 x_0}^4 &= -\frac{\lambda \beta}{2\sigma^3} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m - \frac{\lambda^2 \beta}{2\sigma^4} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{\lambda\lambda}^4 &= -\frac{1}{\sigma^2} \eta_0^\top \mathbf{D} \left[W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right) \right] \eta_0, \\
I_{\lambda x_0}^4 &= \frac{\beta}{\sigma} W_\Phi \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m + \frac{\lambda \beta}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \eta_0, \\
I_{x_0 x_0}^4 &= -\frac{\lambda^2 \beta^2}{\sigma^2} W_\Phi^{(1)} \left(\frac{\lambda\eta_0}{\sigma} \right)^\top \mathbf{1}_m.
\end{aligned}$$

As matrizes $I_1(\boldsymbol{\theta})$ e $I_3(\boldsymbol{\theta})$ podem ser calculadas para cada distribuição MEN considerada no Capítulo 2, nas quais $I_{\omega\boldsymbol{\theta}}^1 = 0$, $\omega = \lambda$ e x_0 e $I_{\lambda\boldsymbol{\theta}}^3 = 0$, com $\boldsymbol{\theta} = (\alpha, \beta, \sigma^2, \lambda, x_0)$, como segue

• *Normal*, $Y \sim N(\mu, \sigma^2)$

$$\begin{aligned} I_{\alpha\alpha}^1 &= \frac{R}{\sigma^2}, \\ I_{\alpha\beta}^1 &= \frac{1}{\sigma^2} \mathbf{x}^\top \mathbf{1}_R, \\ I_{\alpha\sigma^2}^1 &= \frac{1}{\sigma^4} \boldsymbol{\eta}^\top \mathbf{1}_R, \\ I_{\beta\beta}^1 &= \frac{1}{\sigma^2} \mathbf{x}^\top \mathbf{x}, \\ I_{\beta\sigma^2}^1 &= \frac{1}{\sigma^4} \mathbf{x}^\top \boldsymbol{\eta}, \\ I_{\sigma^2\sigma^2}^1 &= -\frac{R}{2\sigma^4} + \frac{1}{\sigma^6} \boldsymbol{\eta}^\top \boldsymbol{\eta}, \end{aligned}$$

$$\begin{aligned} I_{\alpha\alpha}^3 &= \frac{m}{\sigma^2}, \\ I_{\alpha\beta}^3 &= \frac{mx_0}{\sigma^2}, \\ I_{\alpha\sigma^2}^3 &= \frac{1}{\sigma^4} \boldsymbol{\eta}_0^\top \mathbf{1}_m, \\ I_{\alpha x_0}^3 &= \frac{m\beta}{\sigma^2}, \\ I_{\beta\beta}^3 &= \frac{mx_0^2}{\sigma^2}, \\ I_{\beta\sigma^2}^3 &= \frac{x_0}{\sigma^4} \boldsymbol{\eta}_0^\top \mathbf{1}_m, \\ I_{\beta x_0}^3 &= \frac{1}{\sigma^2} [x_0\beta \mathbf{1}_m - \boldsymbol{\eta}_0]^\top \mathbf{1}_m, \\ I_{\sigma^2\sigma^2}^3 &= -\frac{m}{2\sigma^4} + \frac{1}{\sigma^6} \boldsymbol{\eta}_0^\top \boldsymbol{\eta}_0, \\ I_{\sigma^2 x_0}^3 &= \frac{\beta}{\sigma^4} \boldsymbol{\eta}_0^\top \mathbf{1}_m, \\ I_{x_0 x_0}^3 &= \frac{m\beta^2}{\sigma^2}, \end{aligned}$$

• *t de Student generalizada*, $Y \sim Gt(\mu, \sigma^2; \nu)$

$$\begin{aligned}
I_{\alpha\alpha}^1 &= -\frac{\nu+1}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V)D(\eta)\eta - V \right]^\top \mathbf{1}_R, \\
I_{\alpha\beta}^1 &= -\frac{\nu+1}{\gamma\sigma^2} \mathbf{x}^\top \left[\frac{2}{\gamma\sigma^2} D^2(V)D(\eta)\eta - V \right], \\
I_{\alpha\sigma^2}^1 &= -\frac{\nu+1}{\gamma\sigma^4} \boldsymbol{\eta}^\top \left[\frac{1}{\gamma\sigma^2} D^2(V)D(\eta)\eta - V \right], \\
I_{\beta\beta}^1 &= -\frac{\nu+1}{\gamma\sigma^2} \mathbf{x}^\top \left[\frac{2}{\gamma\sigma^2} D^2(V)D^2(\eta) - D(V) \right] \mathbf{x}, \\
I_{\beta\sigma^2}^1 &= -\frac{\nu+1}{\gamma\sigma^4} \mathbf{x}^\top D(\eta) \left[\frac{1}{\gamma\sigma^2} D^2(V)D(\eta)\eta - V \right], \\
I_{\sigma^2\sigma^2}^1 &= -\frac{R}{2\sigma^4} + \frac{\nu+1}{\gamma\sigma^6} \boldsymbol{\eta}^\top D(V)\eta - \frac{\nu+1}{2\gamma^2\sigma^8} \boldsymbol{\eta}^\top D^3(\eta)D(V)V,
\end{aligned}$$

$$\begin{aligned}
I_{\alpha\alpha}^3 &= -\frac{\nu+1}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top \mathbf{1}_m, \\
I_{\alpha\beta}^3 &= -\frac{(\nu+1)x_0}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top \mathbf{1}_m, \\
I_{\alpha\sigma^2}^3 &= -\frac{\nu+1}{\gamma\sigma^4} \boldsymbol{\eta}_0^\top \left[\frac{1}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top, \\
I_{\alpha x_0}^3 &= -\frac{(\nu+1)\beta}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top \mathbf{1}_m, \\
I_{\beta\beta}^3 &= -\frac{(\nu+1)x_0^2}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top \mathbf{1}_m, \\
I_{\beta\sigma^2}^3 &= -\frac{(\nu+1)x_0}{\gamma\sigma^4} \boldsymbol{\eta}_0^\top \left[\frac{1}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right], \\
I_{\beta x_0}^3 &= -\frac{\nu+1}{\gamma\sigma^2} V_0^\top \left[\frac{2\beta x_0}{\gamma\sigma^2} D^2(\eta_0)V_0 + \boldsymbol{\eta}_0 - x_0\beta \mathbf{1}_m \right], \\
I_{\sigma^2\sigma^2}^3 &= -\frac{m}{2\sigma^4} + \frac{\nu+1}{\gamma\sigma^6} \boldsymbol{\eta}_0^\top D(V_0)\eta_0 - \frac{\nu+1}{2\gamma^2\sigma^8} \boldsymbol{\eta}_0^\top D^3(\eta_0)D(V_0)V_0, \\
I_{\sigma^2 x_0}^3 &= -\frac{(\nu+1)\beta}{\gamma\sigma^4} \boldsymbol{\eta}_0^\top \left[\frac{1}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right], \\
I_{x_0 x_0}^3 &= -\frac{(\nu+1)\beta^2}{\gamma\sigma^2} \left[\frac{2}{\gamma\sigma^2} D^2(V_0)D(\eta_0)\eta_0 - V_0 \right]^\top \mathbf{1}_m,
\end{aligned}$$

- *Exponencial potência*, $Y \sim EP(\mu, \sigma^2; \nu)$

$$\begin{aligned}
I_{\alpha\alpha}^1 &= \frac{\nu(2\nu-1)}{\sigma^{2\nu}} (\boldsymbol{\eta}^{2\nu-2})^\top \mathbf{1}_R, \\
I_{\alpha\beta}^1 &= \frac{\nu(2\nu-1)}{\sigma^{2\nu}} \mathbf{x}^\top (\boldsymbol{\eta}^{2\nu-2}), \\
I_{\alpha\sigma^2}^1 &= \frac{\nu^2}{\sigma^{2(\nu+1)}} (\boldsymbol{\eta}^{2\nu-1})^\top \mathbf{1}_R, \\
I_{\beta\beta}^1 &= \frac{\nu(2\nu-1)}{\sigma^{2\nu}} \mathbf{x}^\top D (\boldsymbol{\eta}^{2\nu-2}) \mathbf{x}, \\
I_{\beta\sigma^2}^1 &= \frac{\nu^2}{\sigma^{2(\nu+1)}} \mathbf{x}^\top (\boldsymbol{\eta}^{2\nu-1}), \\
I_{\sigma^2\sigma^2}^1 &= -\frac{R}{2\sigma^4} + \frac{\nu(\nu+1)}{2\sigma^{2(\nu+2)}} (\boldsymbol{\eta}^\nu)^\top (\boldsymbol{\eta}^\nu),
\end{aligned}$$

$$\begin{aligned}
I_{\alpha\alpha}^3 &= \frac{\nu(2\nu-1)}{\sigma^{2\nu}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-2})^\top \mathbf{1}_m, \\
I_{\alpha\beta}^3 &= \frac{\nu(2\nu-1)x_0}{\sigma^{2\nu}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-2})^\top \mathbf{1}_m, \\
I_{\alpha\sigma^2}^3 &= \frac{\nu^2}{\sigma^{2(\nu+1)}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-1})^\top \mathbf{1}_m, \\
I_{\alpha x_0}^3 &= \frac{\nu(2\nu-1)\beta}{\sigma^{2\nu}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-2})^\top \mathbf{1}_m, \\
I_{\beta\beta}^3 &= \frac{\nu(2\nu-1)x_0^2}{\sigma^{2\nu}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-2})^\top \mathbf{1}_m, \\
I_{\beta\sigma^2}^3 &= \frac{\nu^2 x_0}{\sigma^{2(\nu+1)}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-1})^\top \mathbf{1}_m, \\
I_{\beta x_0}^3 &= \frac{\nu}{\sigma^{2\nu}} [\beta x_0(2\nu-1) \boldsymbol{\eta} \mathbf{0}^{2\nu-2} - \boldsymbol{\eta} \mathbf{0}^{2\nu-1}]^\top \mathbf{1}_m, \\
I_{\sigma^2\sigma^2}^3 &= -\frac{m}{2\sigma^4} + \frac{\nu(\nu+1)}{2\sigma^{2(\nu+2)}} (\boldsymbol{\eta} \mathbf{0}^\nu)^\top (\boldsymbol{\eta} \mathbf{0}^\nu), \\
I_{\sigma^2 x_0}^3 &= \frac{\nu^2 \beta}{\sigma^{2(\nu+1)}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-1})^\top \mathbf{1}_m, \\
I_{x_0 x_0}^3 &= \frac{\nu(2\nu-1)\beta^2}{\sigma^{2\nu}} (\boldsymbol{\eta} \mathbf{0}^{2\nu-2})^\top \mathbf{1}_m,
\end{aligned}$$

• *Slash*, $Y \sim SL(\mu, \sigma^2; \mu)$

$$\begin{aligned}
I_{\alpha\alpha}^1 &= -\frac{1}{\sigma^2} \{[d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3), -Ig_1 \odot Ig_3] \odot (Ig_1 \odot Ig_1)\}^\top \mathbf{1}_R, \\
I_{\alpha\beta}^1 &= -\frac{1}{\sigma^2} \{[d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3), -Ig_1 \odot Ig_3] \odot (Ig_1 \odot Ig_1)\}^\top \mathbf{x}, \\
I_{\alpha\sigma^2}^1 &= -\frac{1}{2\sigma^4} \{[d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3), -2Ig_1 \odot Ig_3] \odot (Ig_1 \odot Ig_1)\}^\top \boldsymbol{\eta}, \\
I_{\beta\beta}^1 &= -\frac{1}{\sigma^2} \mathbf{x}^\top D \{[d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3), -Ig_1 \odot Ig_3] \odot (Ig_1 \odot Ig_1)\} \mathbf{x}, \\
I_{\beta\sigma^2}^1 &= -\frac{1}{2\sigma^4} \mathbf{x}^\top D \{[d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3), -2Ig_1 \odot Ig_3] \odot (Ig_1 \odot Ig_1)\} \boldsymbol{\eta}, \\
I_{\sigma^2\sigma^2}^1 &= -\frac{1}{2\sigma^4} \{R + d^\top [(1/2d \odot (Ig_1 \odot Ig_5 - Ig_3 \odot Ig_3) - 2Ig_1 \odot Ig_3) \odot \\
&\quad (Ig_1 \odot Ig_1)]\},
\end{aligned}$$

$$\begin{aligned}
I_{\alpha\alpha}^3 &= -\frac{1}{\sigma^2} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m, \\
I_{\alpha\beta}^3 &= -\frac{x_0}{\sigma^2} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m, \\
I_{\alpha\sigma^2}^3 &= -\frac{1}{2\sigma^4} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - 2Ig_{01} \odot Ig_{03}] \odot \\
&\quad (Ig_{01} \odot Ig_{01})\}^\top \boldsymbol{\eta}_0, \\
I_{\alpha x_0}^3 &= -\frac{\beta}{\sigma^2} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m, \\
I_{\beta\beta}^3 &= -\frac{x_0^2}{\sigma^2} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m, \\
I_{\beta\sigma^2}^3 &= -\frac{x_0}{2\sigma^4} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - 2Ig_{01} \odot Ig_{03}] \odot \\
&\quad (Ig_{01} \odot Ig_{01})\}^\top \boldsymbol{\eta}_0, \\
I_{\beta x_0}^3 &= -\frac{1}{\sigma^2} \{[x_0\beta [d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] + \\
&\quad \boldsymbol{\eta}_0 \odot Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m, \\
I_{\sigma^2\sigma^2}^3 &= -\frac{1}{2\sigma^4} \{m + d_0^\top [(1/2d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - \\
&\quad 2Ig_{01} \odot Ig_{03}) \odot (Ig_{01} \odot Ig_{01})]\}, \\
I_{\sigma^2 x_0}^3 &= -\frac{\beta}{2\sigma^4} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - 2Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \boldsymbol{\eta}_0, \\
I_{x_0 x_0}^3 &= -\frac{\beta^2}{\sigma^2} \{[d_0 \odot (Ig_{01} \odot Ig_{05} - Ig_{03} \odot Ig_{03}) - Ig_{01} \odot Ig_{03}] \odot (Ig_{01} \odot Ig_{01})\}^\top \mathbf{1}_m,
\end{aligned}$$

onde

$$\begin{aligned}
I g_{kij} &= \int_0^1 u^{\nu+k/2-1} e^{-u d_{ij}/2} du = \frac{\Gamma(\nu+k/2)}{(d_{ij}/2)^{\nu+k/2}} \mathbb{P}_1(\nu+k/2, d_{ij}/2), \\
U_{ij,k} &\sim Gama(\nu+k/2, d_{ij}/2), \\
I g_{0ki} &= \int_0^1 u_0^{\nu+k/2-1} e^{-u_0 d_{0i}/2} du_0 = \frac{\Gamma(\nu+k/2)}{(d_{0i}/2)^{\nu+k/2}} \mathbb{P}_1(\nu+k/2, d_{0i}/2), \\
U_{0i,k} &\sim Gama(\nu+k/2, d_{0i}/2) \text{ e}
\end{aligned}$$

$\mathbb{P}_1(a, b)$ é a fda de uma variável aleatória $Gama(a, b)$ avaliada em 1 e $\odot$ e $\oslash$ são o produto e a divisão de Hadamard, respectivamente.

3.2.4 Matriz hessiana

A matriz hessiana das distribuições *Normal Assimétrica*, *t de Student Assimétrica*, *Exponencial Potência Assimétrica* e *Slash Assimétrica* é a matriz oposta da matriz de informação de Fisher observada, cujas expressões de seus elementos foram apresentadas na seção anterior. Mas, de forma geral, os elementos da matriz hessiana podem ser obtidos por meio das expressões a seguir, derivadas de (3.25).

$$\begin{aligned}
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \alpha \partial \alpha} &= -\frac{1}{\sigma^2} \left[(\boldsymbol{\kappa} + \lambda^2 \mathbf{1}_R)^\top \mathbf{1}_R + (\boldsymbol{\kappa}_0 + \lambda^2 \mathbf{1}_m)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \alpha \partial \beta} &= -\frac{1}{\sigma^2} \left[(\boldsymbol{\kappa} + \lambda^2 \mathbf{1}_R)^\top \mathbf{x} + x_0 (\boldsymbol{\kappa}_0 + \lambda^2 \mathbf{1}_m)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \alpha \partial \sigma^2} &= -\frac{1}{\sigma^4} \left[\left(D(\boldsymbol{\kappa})\boldsymbol{\eta} + \lambda^2 \boldsymbol{\eta} - \lambda \widehat{\mathbf{t}} \right)^\top \mathbf{1}_R + \left(D(\boldsymbol{\kappa}_0)\boldsymbol{\eta}_0 + \lambda^2 \boldsymbol{\eta}_0 - \lambda \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \alpha \partial \lambda} &= \frac{1}{\sigma^2} \left[\left(2\lambda \boldsymbol{\eta} - \widehat{\mathbf{t}} \right)^\top \mathbf{1}_R + \left(2\lambda \boldsymbol{\eta}_0 - \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \alpha \partial x_0} &= -\frac{\beta}{\sigma^2} \left[(\boldsymbol{\kappa}_0 + \lambda^2 \mathbf{1}_m)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \beta \partial \beta} &= -\frac{1}{\sigma^2} \left[(D(\boldsymbol{\kappa})\mathbf{x} + \lambda^2 \mathbf{x})^\top \mathbf{x} + x_0^2 (\boldsymbol{\kappa}_0 + \lambda^2 \mathbf{1}_m)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \beta \partial \sigma^2} &= -\frac{1}{\sigma^4} \left[\left(D(\boldsymbol{\kappa})\boldsymbol{\eta} + \lambda^2 \boldsymbol{\eta} - \lambda \widehat{\mathbf{t}} \right)^\top \mathbf{x} + x_0 \left(D(\boldsymbol{\kappa}_0)\boldsymbol{\eta}_0 + \lambda^2 \boldsymbol{\eta}_0 - \lambda \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \beta \partial \lambda} &= \frac{1}{\sigma^2} \left[\left(2\lambda \boldsymbol{\eta} - \widehat{\mathbf{t}} \right)^\top \mathbf{x} + x_0 \left(2\lambda \boldsymbol{\eta}_0 - \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \beta \partial x_0} &= \frac{1}{\sigma^2} \left[\left(D(\boldsymbol{\kappa}_0)\boldsymbol{\eta}_0 - x_0 \beta \boldsymbol{\kappa}_0 + \lambda^2 (\boldsymbol{\eta}_0 - x_0 \beta \mathbf{1}_m) - \lambda \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \sigma^2 \partial \sigma^2} &= \frac{R+m}{\sigma^4} - \frac{1}{\sigma^6} \left[\boldsymbol{\eta}^\top D(\boldsymbol{\kappa})\boldsymbol{\eta} + \lambda^2 \boldsymbol{\eta}^\top \boldsymbol{\eta} + \widehat{\mathbf{t}}^2^\top \mathbf{1}_R - 2\lambda \boldsymbol{\eta}^\top \widehat{\mathbf{t}} \right. \\
&\quad \left. + \boldsymbol{\eta}_0^\top D(\boldsymbol{\kappa}_0)\boldsymbol{\eta}_0 + \lambda^2 \boldsymbol{\eta}_0^\top \boldsymbol{\eta}_0 + \widehat{\mathbf{t}}_0^2^\top \mathbf{1}_m - 2\lambda \boldsymbol{\eta}_0^\top \widehat{\mathbf{t}}_0 \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \sigma^2 \partial \lambda} &= \frac{1}{\sigma^4} \left[\lambda \boldsymbol{\eta}^\top \boldsymbol{\eta} - \boldsymbol{\eta}^\top \widehat{\mathbf{t}} + \lambda \boldsymbol{\eta}_0^\top \boldsymbol{\eta}_0 - \boldsymbol{\eta}_0^\top \widehat{\mathbf{t}}_0 \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \sigma^2 \partial x_0} &= -\frac{\beta}{\sigma^4} \left[\boldsymbol{\kappa}_0^\top \boldsymbol{\eta}_0 + \left(\lambda^2 \boldsymbol{\eta}_0 - \lambda \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \lambda \partial \lambda} &= -\frac{1}{\sigma^2} \left[\boldsymbol{\eta}^\top \boldsymbol{\eta} + \boldsymbol{\eta}_0^\top \boldsymbol{\eta}_0 \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial \lambda \partial x_0} &= \frac{\beta}{\sigma^2} \left[\left(2\lambda \boldsymbol{\eta}_0 - \widehat{\mathbf{t}}_0 \right)^\top \mathbf{1}_m \right], \\
\frac{\partial^2 Q(\boldsymbol{\theta}|\widehat{\boldsymbol{\theta}})}{\partial x_0 \partial x_0} &= \frac{\beta}{\sigma^2} \left[\boldsymbol{\kappa}_0^\top \boldsymbol{\eta}_0 + \lambda^2 \boldsymbol{\eta}_0^\top \mathbf{1}_m \right].
\end{aligned}$$

Estudo de Simulação

As simulações foram realizadas a fim de estudar o comportamento do estimador de máxima verossimilhança de x_0 . Os valores considerados na simulação foram $x_0 = 0,8$, $\alpha = 0,1$, $\beta = 2,0$, $\sigma^2 = 0,1$ e $1,0$, $\lambda = 0,5$ e $2,0$, já o parâmetro ν considerado conhecido, variou de acordo com a distribuição utilizada, com exceção da distribuição *normal assimétrica* que não depende deste parâmetro. Para o caso da distribuição *t de Student assimétrica*, $\nu = 3,0$ e $6,0$, para a *slash assimétrica*, $\nu = 0,5$ e $1,0$ e para a *exponencial potência assimétrica*, $\nu = 0,6$ e $0,8$. Os valores $x_1, \dots, x_n$ são equidistantes, sendo $x_1 = 0$ e $x_n = 2$.

Utilizou-se o *software R*, versão 2.14.1, para gerar 1000 amostras *Monte Carlo* de cada distribuição por meio do pacote '*mixsmn*', cujos tamanhos amostrais considerados no primeiro (n) e segundo (m) estágios foram 5 e 20 com $r = 10$ repetições para cada amostra. As estimativas de máxima verossimilhança dos parâmetros foram obtidas utilizando o *algoritmo EM*, no qual o critério de convergência adotado foi, tal que, para cada parâmetro do vetor θ , $|\theta^{(p+1)} - \theta^p| < 10^{-5}$.

As medidas empíricas analisadas no estudo de simulação foram a média, o viés, o erro quadrático médio e a variância, obtidos, respectivamente, por

$$\begin{aligned}\widehat{x_0} &= \sum_{i=1}^{1000} \frac{\widehat{x_{0i}}}{1000}, \\ \text{viés} &= \sum_{i=1}^{1000} \frac{(x_0 - \widehat{x_{0i}})}{1000}, \\ \text{EQM} &= \sum_{i=1}^{1000} \frac{(x_0 - \widehat{x_{0i}})^2}{1000} e \\ \widehat{\text{var}}(\widehat{x_0}) &= \sum_{i=1}^{1000} \frac{\widehat{\text{var}}(\widehat{x_{0i}})}{1000}\end{aligned}$$

onde $\widehat{\text{var}}(\widehat{x_{0i}})$ é um elemento da matriz inversa da matriz de informação de Fisher observada.

Tabela 4.1: Distribuição t de Student Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	3	0,5	0,8064	0,0064	0,0097	0,0079
			6		0,8014	0,0014	0,0070	0,0064
			3	2,0	0,8116	0,0116	0,0063	0,0046
			6		0,8014	0,0014	0,0039	0,0034
5	10	20	3	0,5	0,7981	-0,0019	0,0025	0,0025
			6		0,7971	-0,0029	0,0021	0,0021
			3	2,0	0,8018	0,0018	0,0015	0,0014
			6		0,7994	-0,0006	0,0012	0,0011
20	10	5	3	0,5	0,8028	0,0028	0,0082	0,0074
			6		0,8009	0,0009	0,0060	0,0057
			3	2,0	0,8094	0,0094	0,0102	0,0042
			6		0,8057	0,0057	0,0034	0,0033
20	10	20	3	0,5	0,8029	0,0029	0,0019	0,0019
			6		0,7984	-0,0016	0,0016	0,0016
			3	2,0	0,8008	0,0008	0,0014	0,0010
			6		0,8022	0,0022	0,0009	0,0009

Tabela 4.2: Distribuição t de Student Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	3	0,5	0,8048	0,0048	0,0966	0,0777
			6		0,7023	-0,0977	0,0691	0,0631
			3	2,0	0,8251	0,0251	0,0470	0,0397
			6		0,7943	-0,0057	0,0322	0,0301
5	10	20	3	0,5	0,7886	-0,0114	0,0231	0,0250
			6		0,7631	-0,0369	0,0218	0,0209
			3	2,0	0,7957	-0,0043	0,0129	0,0131
			6		0,7956	-0,0044	0,0103	0,0100
20	10	5	3	0,5	0,8191	0,0191	0,0840	0,0705
			6		0,8021	0,0021	0,0593	0,0586
			3	2,0	0,8188	0,0188	0,0409	0,0367
			6		0,8189	0,0189	0,0305	0,0284
20	10	20	3	0,5	0,8078	0,0078	0,0194	0,0189
			6		0,7873	-0,0127	0,0154	0,0158
			3	2,0	0,8046	0,0046	0,0097	0,0096
			6		0,8063	0,0063	0,0084	0,0078

Tabela 4.3: Distribuição Normal Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$

n	r	m	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,5	0,8026	0,0026	0,0050	0,0049
			2,0	0,8006	0,0006	0,0027	0,0027
5	10	20	0,5	0,7998	-0,0002	0,0016	0,0016
			2,0	0,7992	-0,0008	0,0007	0,0008
20	10	5	0,5	0,7985	-0,0014	0,0046	0,0044
			2,0	0,8015	0,0015	0,0024	0,0024
20	10	20	0,5	0,7998	-0,0002	0,0012	0,0011
			2,0	0,8003	0,0003	0,0006	0,0006

Tabela 4.4: Distribuição Normal Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$

n	r	m	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,5	0,7831	-0,0169	0,0612	0,0470
			2,0	0,7984	-0,0016	0,0255	0,0224
5	10	20	0,5	0,7716	-0,0284	0,0212	0,0117
			2,0	0,8006	0,0006	0,0080	0,0072
20	10	5	0,5	0,7994	-0,0006	0,0435	0,0437
			2,0	0,8096	0,0096	0,0023	0,0022
20	10	20	0,5	0,8007	0,0007	0,0124	0,0120
			2,0	0,7999	-0,0001	0,0065	0,0060

Tabela 4.5: Distribuição Slash Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,5	0,5	0,8099	0,0099	0,0795	0,0244
			1,0		0,7997	-0,0003	0,0250	0,0128
			0,5	2,0	0,8173	0,0173	0,0467	0,0135
			1,0		0,8096	0,0096	0,0090	0,0067
5	10	20	0,5	0,5	0,7914	-0,0086	0,0092	0,0072
			1,0		0,7979	-0,0021	0,0043	0,0038
			0,5	2,0	0,7924	-0,0076	0,0045	0,0036
			1,0		0,7989	-0,0011	0,0024	0,0021
20	10	5	0,5	0,5	0,8036	0,0036	0,0580	0,0240
			1,0		0,7970	-0,0030	0,0854	0,0150
			0,5	2,0	0,8173	0,0173	0,0246	0,0131
			1,0		0,8054	0,0054	0,0081	0,0070
20	10	20	0,5	0,5	0,8012	0,0012	0,0062	0,0061
			1,0		0,8011	0,0011	0,0035	0,0032
			0,5	2,0	0,8020	0,0020	0,0036	0,0034
			1,0		0,8007	0,0007	0,0020	0,0018

Tabela 4.6: Distribuição Slash Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,5	0,5	0,7594	-0,0406	0,8260	0,3679
			1,0		0,7147	-0,0853	0,1638	0,1544
			0,5	2,0	0,8711	0,0711	0,2659	0,1566
			1,0		0,7671	-0,0329	0,0820	0,0741
5	10	20	0,5	0,5	0,7915	-0,0085	0,1017	0,1158
			1,0		0,7163	-0,0837	0,0510	0,0510
			0,5	2,0	0,7850	-0,0150	0,0511	0,0597
			1,0		0,7525	-0,0475	0,0252	0,0258
20	10	5	0,5	0,5	0,8569	0,0569	0,6712	0,2520
			1,0		0,8087	0,0087	0,1404	0,1275
			0,5	2,0	0,8597	0,0597	0,4014	0,1386
			1,0		0,8268	0,0268	0,0709	0,0675
20	10	20	0,5	0,5	0,8153	0,0153	0,0662	0,0647
			1,0		0,7756	-0,0244	0,0336	0,0349
			0,5	2,0	0,8169	0,0169	0,0334	0,0353
			1,0		0,7866	-0,0134	0,0154	0,0174

Tabela 4.7: Distribuição Exponencial Potência Assimétrica, $\sigma^2 = 0,1$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,6	0,5	0,8150	0,0150	0,0033	0,0023
			0,8		0,8058	0,0058	0,0042	0,0042
			0,6	2,0	0,8041	0,0041	0,0027	0,0023
			0,8		0,8033	0,0033	0,0027	0,0023
5	10	20	0,6	0,5	0,8050	0,0050	0,0003	0,0006
			0,8		0,8033	0,0033	0,0012	0,0013
			0,6	2,0	0,8007	0,0007	0,0012	0,0034
			0,8		0,8020	0,0020	0,0007	0,0007
20	10	5	0,6	0,5	0,8149	0,0149	0,0027	0,0017
			0,8		0,8085	0,0085	0,0040	0,0038
			0,6	2,0	0,7970	-0,0030	0,0039	0,0028
			0,8		0,8022	0,0022	0,0030	0,0029
20	10	20	0,6	0,5	0,8043	0,0043	0,0002	0,0003
			0,8		0,8065	0,0065	0,0009	0,0010
			0,6	2,0	0,8008	0,0008	0,0009	0,0001
			0,8		0,8015	0,0015	0,0010	0,0011

Tabela 4.8: Distribuição Exponencial Potência Assimétrica, $\sigma^2 = 1,0$ e $x_0 = 0,8$

n	r	m	ν	λ	$\hat{x}_0$	viés	EQM	$\widehat{\text{var}}(\hat{x}_0)$
5	10	5	0,6	0,5	0,8086	0,0086	0,0220	0,0212
			0,8		0,8074	0,0074	0,0380	0,0377
			0,6	2,0	0,8259	0,0259	0,0094	0,0096
			0,8		0,8390	0,0390	0,0233	0,0193
5	10	20	0,6	0,5	0,8012	0,0012	0,0036	0,0062
			0,8		0,7985	-0,0015	0,0121	0,0123
			0,6	2,0	0,8030	0,0030	0,0028	0,0036
			0,8		0,8026	0,0026	0,0059	0,0060
20	10	5	0,6	0,5	0,8118	0,0118	0,0199	0,0206
			0,8		0,8182	0,0182	0,0408	0,0362
			0,6	2,0	0,8162	0,0162	0,0083	0,0065
			0,8		0,8317	0,0317	0,0188	0,0138
20	10	20	0,6	0,5	0,8022	0,0022	0,0020	0,0032
			0,8		0,8001	0,0001	0,0097	0,0096
			0,6	2,0	0,7999	-0,0001	0,0028	0,0022
			0,8		0,8047	0,0047	0,0044	0,0040

Observa-se nas Tabelas 4.1 a 4.8 que o viés, o EQM e a variância das estimativas de x_0 diminuem com o aumento do tamanho da amostra tanto no primeiro quanto no segundo estágio, n e m , respectivamente. Observa-se também que essas medidas são menores quando $\sigma^2 = 0, 1$, comparado com os resultados considerando $\sigma^2 = 1, 0$.

Verifica-se também nas Tabelas 4.1 a 4.8, que o EQM, na maioria dos cenários, é menor quando considera-se o parâmetro $\lambda = 2, 0$ comparado com $\lambda = 0, 5$.

Aplicações

As aplicações apresentadas neste Capítulo foram realizadas em conjuntos de dados reais obtidos de Neto *et al.* (2007) e Oliveira e Aguiar (2009).

Os critérios de informação considerados neste trabalho são AIC (Akaike Information Criterion) e BIC (Bayesian Information Criterion) os quais são úteis para a seleção de modelos. Quanto menor é o valor das estatísticas AIC e BIC, segundo as expressões abaixo, mais adequado é o modelo.

$$\text{AIC} = -2\ell(\hat{\boldsymbol{\theta}}) + 2k$$

$$\text{BIC} = -2\ell(\hat{\boldsymbol{\theta}}) + k \log(R + m)$$

em que $\ell(\hat{\boldsymbol{\theta}})$ é o valor máximo da função de log-verossimilhança (3.43), k é o número de parâmetros desconhecidos no modelo, R e m são os números de observações da amostra no primeiro e segundo estágios, respectivamente.

Para escolha do parâmetro ν nas distribuições estudadas que dependem desse parâmetro, foi utilizado o critérios de informação $\text{AIC} = -2\ell(\hat{\boldsymbol{\theta}}) + 2k$ como critério de seleção. Assim, a escolha de ν foi feita segundo o menor valor da estatística AIC.

Nas aplicações calculou-se as estimativas dos parâmetro e seus respectivos erros padrões, obtidos por meio da matriz inversa da matriz de informação de Fisher observada.

5.1 Concentração de zinco

O objetivo deste experimento é estimar a concentração de zinco por meio do modelo de calibração proposto relacionando a resposta do instrumento, neste caso a absorbância (y), com a concentração (x) do elemento zinco na amostra. Na Tabela 5.1 são apresentadas as concentrações de soluções aquosas contendo íons de zinco e as respectivas absorbâncias, obtidas em triplicata utilizando a técnica de *espectrometria de absorção atômica* (Neto *et al.*, 2007).

Tabela 5.1: Concentração (mg/l) e absorbâncias das soluções-padrão de zinco.

Concentração de zinco	Respostas de medidas repetidas		
x_i	y_{i1}	y_{i2}	y_{i3}
0,0	0,696	0,696	0,706
0,5	7,632	7,688	7,603
1,0	14,804	14,861	14,731
2,0	28,895	29,156	29,322
3,0	43,993	43,574	44,699

Ajustou-se quatro modelos de calibração para as seguintes distribuições dos erros: *t de Student assimétrica*, *normal assimétrica*, *slash assimétrica* e *exponencial potência assimétrica*. Para efeito de aplicação, utiliza-se $x_3 = 1,0$, como valor desconhecido de x_0 e os

respectivos valores de y_{3j} referentes a essa concentração serão tomados como a observação de y_0 para compor os dados do segundo estágio da calibração. Nesta aplicação $R = 4 \times 3 = 12$.

Tabela 5.2: Estimativas dos parâmetros para as distribuições tA , NA , SLA e EPA .

Distribuição	Parâmetros				
	α	β	σ^2	λ	x_0
tA ($\nu = 2$)	0,4968 (0,0186)	14,1949 (0,0113)	0,0410 (0,0256)	40,1850 (48,2166)	1,0020 (0,0015)
NA	0,2997 (0,2241)	14,2904 (0,1121)	0,2698 (0,1014)	33,2747 (57,8539)	1,0035 (0,0227)
SLA ($\nu = 1, 5$)	0,4905 (0,0016)	14,1980 (0,0068)	0,0527 (0,0013)	40,1364 (48,0320)	1,0019 (0,0021)
EPA ($\nu = 0, 8$)	0,3121 (0,2505)	14,2852 (0,1249)	0,1587 (0,0723)	33,0014 (55,2220)	1,0006 (0,0221)

Na Tabela 5.2 apresenta-se as estimativas dos parâmetros do modelo calibração proposto segundo as distribuições $tA(2)$, NA , $SLA(1,5)$ e $EPA(0,8)$, em que observa-se que a distribuição *exponencial potência assimétrica* apresenta o menor viés e a distribuição *t de Student assimétrica* apresenta o menor erro padrão.

Tabela 5.3: Critérios de informação - Concentração de zinco.

Distribuição	AIC	BIC
tA ($\nu = 2$)	-80,35	-76,80
NA	-7,99	-4,45
SLA ($\nu = 1, 5$)	-70,47	-66,93
EPA ($\nu = 0, 8$)	-9,03	-5,49

Segundo os critérios de informação apresentados na Tabela 5.3, a distribuição mais adequada para esta aplicação é a distribuição *t de Student assimétrica* por apresentar o menor valor para as estatísticas AIC e BIC.

5.2 Concentração de benzatona

Nesta aplicação o objetivo é estimar a concentração cromatográfica da benzatona em diferentes concentrações por meio do modelo de calibração linear proposto, relacionando a altura do pico como sendo o (y), com a concentração (x). Na Tabela 5.4 são apresentadas as concentrações em mgL^{-1} e as respectivas alturas dos picos em centímetros (Oliveira e Aguiar, 2009).

Tabela 5.4: Concentração (mgL^{-1}) e altura do pico (cm) das soluções-padrão de benzonato.

Concentração mgL^{-1}	Respostas de medidas repetidas				
x_i	y_{i1}	y_{i2}	y_{i3}	y_{i4}	y_{i5}
0,0133	0,1836	0,1787	0,1837	0,1806	0,1861
0,0665	0,9373	0,9177	0,9224	-	-
0,3325	4,6227	4,6280	4,6256	-	-
0,6650	9,6905	9,9405	9,5754	-	-
0,9975	14,7607	15,0113	14,9641	-	-
1,3300	21,3001	20,2700	20,5719	20,3540	-

Nesta aplicação ajustou-se também quatro modelos de calibração com as seguintes distribuições dos erros: *t de Student assimétrica*, *normal assimétrica*, *slash assimétrica* e *exponencial potência assimétrica*. Para efeito de aplicação, toma-se $x_5 = 0,9975$, como valor desconhecido de x_0 e os respectivos valores de y_{5j} referentes a essa concentração serão tomados como a observação de y_0 para compor os dados do segundo estágio da calibração. Nesta aplicação $R = \sum_{k=1}^5 r_k = 18$.

Tabela 5.5: Estimativas dos parâmetros para as distribuições tA , NA , SLA e EPA .

Distribuição	Parâmetros				
	α	β	σ^2	λ	x_0
tA ($\nu = 2$)	-0,4912 (0,3683)	15,1377 (0,5306)	0,1415 (0,0937)	16,1885 (26,6169)	1,0049 (0,0112)
NA	-0,5434 (0,1471)	15,2283 (0,2021)	0,3030 (0,0911)	16,0483 (18,1250)	1,0005 (0,0017)
SLA ($\nu = 3$)	-0,5310 (0,0424)	15,1988 (0,0707)	0,1940 (0,0184)	16,1997 (18,7150)	1,0024 (0,0035)
EPA ($\nu = 0,6$)	-0,4976 (0,3629)	15,1514 (0,5426)	0,0848 (0,0433)	16,2786 (36,9024)	1,0055 (0,0110)

Na Tabela 5.5 apresenta-se as estimativas dos parâmetros do modelo calibração proposto segundo as distribuições $tA(2)$, NA , $SLA(3)$ e $EPA(0,6)$, em que observa-se que a distribuição *normal assimétrica* apresenta o menor viés na estimação do parâmetro x_0 , seguido da distribuição *slash assimétrica*. As distribuições *normal assimétrica* e *slash assimétrica* apresentam os menores erros padrões.

Tabela 5.6: Critérios de informação - Concentração de benzetona.

Distribuição	AIC	BIC
tA ($\nu = 2$)	-77,00	-71,77
NA	-11,91	-6,68
SLA ($\nu = 3$)	-97,04	-91,81
EPA ($\nu = 0,6$)	-12,64	-7,42

Segundo os valores dos critérios de informação apresentados na Tabela 5.6, a distribuição mais adequada para ser utilizada é a distribuição *slash assimétrica* por apresentar o menor valor para as estatísticas AIC e BIC.

Conclusões

6.1 Considerações finais

Os resultados obtidos nesse trabalho mostram uma gama de distribuições simétricas e assimétricas pertencentes à família *MENA* para a distribuição dos erros no modelo de calibração linear proposto. Assim, o modelo de calibração linear proposto é um caso geral de modelos de calibração linear presentes na literatura. Observa-se que o modelo proposto pode ser utilizado quando a distribuição dos erros é assimétrica ou quando existir observações destoantes em dados reais.

Observou-se nos resultados das simulações apresentados nas Tabelas 4.1 a 4.8 que o estimador de máxima verossimilhança via *algoritmo EM* mostrou-se consistente, pois o viés e o EQM do estimador do parâmetro x_0 diminuem, na maioria dos cenários, com o aumento do tamanho das amostras, tanto no primeiro quanto no segundo estágios.

A utilização do *algoritmo EM* neste trabalho foi de grande importância, pois facilitou os cálculos e viabilizou o desenvolvimento do modelo de calibração linear proposto.

Apesar da versatilidade do modelo de calibração linear proposto, é preciso ressaltar que a utilização deste modelo requer um conhecimento estatístico refinado e bom suporte computacional para a implementação do *algoritmo EM*.

6.2 Perspectivas para pesquisas futuras

Para realização de trabalhos futuros sugere-se o estudo da extensão deste modelo para o modelo com erros de medida nas variáveis independentes, além do estudo e desenvolvimento da matriz de informação de Fisher esperada, da correção de viés dos estimadores, de um modelo com a estimação do parâmetro ν e de modelos de calibração polinomial.

Sugere-se também a realização de estudo para obtenção de outras distribuições pertencentes a família de distribuições *MENA*.

Referências Bibliográficas

- [1] Andrews, D. R. & Mallows, C. L. Scale mixtures of normal distributions. *Journal of the Royal Statistical Society, Series B*, 36, p. 99102, 1974.
- [2] Arellano-Valle, R. B; Bolfarine, H.; Lachos, V. H. Skew-normal linear mixed models. *Journal of Data Science*, 2005.
- [3] Blas, B. G. *Modelos de regressão e calibração com erros de medida*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 2010.
- [4] Blas, B. G.; Sandoval, M. C. Heteroscedastic controlled calibration model applied to analytical chemistry. *Journal of Chemometrics*. 24, 2010.
- [5] Blas, B. G.; Sandoval, M. C.; Yoshida, O. S. Homoscedastic controlled calibration model. *Journal of Chemometrics*. 21, 2007.
- [6] Blas, B. G. *Calibração controlada aplicada na química analítica*. Dissertação de Mestrado em Estatística, Universidade de São Paulo, São Paulo, 2005.
- [7] Bolfarine, H.; Lima, C. R. O. P.; Sandoval, M. C. Linear calibration in functional regression models. *Communications in Statistics: Theory and Methods*. 26, 1997.
- [8] Branco, M. D.; Dey, D. K. A general class of multivariate skew-elliptical distributions. *Journal of Multivariate Analysis*, 79, p. 99-113, 2001.

- [9] Branco, M. D. *Calibração: uma abordagem bayesiana*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 1997.
- [10] Brereton, R. G. *Chemometrics: data analysis for the laboratory and chemical plant*. John Wiley & Sons Ltd. London, 2003.
- [11] Casella, G.; Berguer, R. L. *Inferência Estatística*. Tradução da 2ª edição norte-americana. CENGAGE Learning. São Paulo, 2011.
- [12] Dempster, A. P.; Laird, N. M.; Rubin, D. B. Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society. Series B*, 39, p. 1-38, 1977.
- [13] Dempster, A. P.; Laird, N. M.; Rubin, D. B. Iteratively reweighted least squares for linear regression when errors are Normal/Independent distributed. Em P. R. Krishnaiah (Ed). *Multivariate Analysis V*, p. 35-57. North-Holland. 1980.
- [14] Eisenhart, C. The interpretation of certain regression methods, and their use in biological and industrial research. *Annals of Mathematical Statistics*. 10, 1939.
- [15] Ferreira, C. S., Bolfarine, H., Lachos, V. H. Skew scale mixtures of normal distributions: Properties and estimation. *Statistical Methodology*, 2011.
- [16] Ferreira, C. S. *Influência e diagnóstico em modelos assimétricos*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 2008.
- [17] Figueiredo, C. C., Bolfarine, H., Sandoval, M. C., C. R. O. P. Lima. On the skew-normal calibration model. *Journal of Applied Statistics*, 2010.
- [18] Figueiredo, C. C. *Calibração linear assimétrica*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 2009.
- [19] Francisconi, C. N. *Comparação de instrumentos de medição usando calibração comparativa*. Dissertação de Mestrado em Estatística, Universidade de São Paulo, São Paulo, 1996.

- [20] Freitas, L. A. *Modelos de regressão com erros normais assimétricos: uma abordagem bayesiana*. Dissertação de Mestrado em Estatística, Universidade Federal de São Carlos, São Carlos, 2005.
- [21] Galea-Rojas, M. *Calibração comparativa estrutural e funcional*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 1995.
- [22] Hartley, H. O. Maximum likelihood estimation from incomplete data. *Biometrics*. 14, 1958.
- [23] Henze, N. A probabilistic representation of the 'Skew-normal' distribution. *Scandinavian Journal of Statistics*. 13, 1986.
- [24] Krutchkoff, R. Classical and inverse regression methods of calibration. *Technometrics*. 9, 1967.
- [25] Lange, K. L.; Little, J.A.; Taylor, M.G. Robust statistical modeling using the t distribution. *Journal of the American Statistical Association*, 84, p. 881896, 1989.
- [26] Lima, C. R. O. P. *Calibração absoluta com erros nas variáveis*. Tese de Doutorado em Estatística, Universidade de São Paulo, São Paulo, 1996.
- [27] Marciano, F. W. P. *Modelo de calibração com erros simétricos*. Dissertação de Mestrado em Estatística, Universidade Federal de Pernambuco, Recife, 2012.
- [28] Neto, B. B.; Scarminio, I. S.; Bruns, R. E. *Como fazer experimentos: pesquisa e desenvolvimento na ciência e na indústria*. Editora da Unicamp, Campinas. 2007.
- [29] Oliveira, E. C.; Aguiar, P. F. Validação da Metodologia da Avaliação de Incertezas em Curvas de Calibração Melhor Ajustadas por Polinômios de 2º Grau. *Química Nova*. 32, 2009.
- [30] Salgado, F. A. O. *Diagnóstico de influência em modelos elípticos com efeitos mistos*. Tese de Doutorado, Universidade de São Paulo, São Paulo, 2006.
- [31] Shukla G. K. On the problem of calibration. *Technometrics*, 1972.

- [32] Zeller, C. B.; Lachos, V. H.; Vilca-Labra, F. E. Local influence analysis for regression models with scale mixtures of skew-normal distributions. *Journal of Applied Statistics*, 2010.
- [33] Zeller, C. B. *Distribuições misturas de escala skew-normal: estimação e diagnóstico em modelos lineares*. Tese de Doutorado, Universidade Estadual de Campinas, São Paulo, 2006.

Apêndice - Código em *R*

Neste apêndice está disponibilizado o código usado no *software R* para a geração de amostras artificiais com a utilização do pacote '*mixsmsn*' e estimação dos parâmetros da calibração para um modelo cujos erros seguem a distribuição *t de Student assimétrica* via *algoritmo EM*.

Para alterar a distribuição dos erros na geração da amostra e na estimação dos parâmetros, basta alterar o campo `family` na função `rmix` e as equações dos vetores `kap` e `kap0`, conforme equações (3.33) a (3.42).

```
#Calibração linear com misturas de escala normal assimétrica
#Estimação dos parâmetros pelo algoritmo EM
#Autor: Marcos Antonio Alves Pereira - Mestrando em Estatística UFPE / 2013

#Geração de dados artificiais (t de Student assimétrica)
#Instalar o pacote "mixsmsn"
```

```

library(mixsmsn)

xi = c(25.83, 34.31, 42.50, 46.75, 48.29, 48.77, 49.65, 50, 54.33, 54.87,
       56.46, 58.83, 59.13, 60.73, 61.12, 63.10, 65.96, 66.40, 70.42, 70.48,
       71.98, 72.00, 72.23, 72.95, 73.44, 74.25, 74.77, 76.33, 81.02, 81.85,
       82.56, 83.33, 83.40, 91.81, 92.10, 92.96, 95.17, 101.4, 104.1, 106.5)

n = length(xi)

#Parâmetros da simulação

mureal = 0
alphareal = 2
betareal = 4
sig2real = 0.5
lbdreal = 2
nu = 4

arg1 = c(mureal, sig2real, lbdreal, nu)

erro = rmix(n, p = 1, family = "Skew.t", arg = list(arg1))
y = alphareal + betareal*xi + erro

x = xi[-c(8)] # valor real de x0 (50.0)
y0 = y[8]
y = y[-c(8)]

n = length(x)          # dimensão de x e y

```

```

m = length(y0)          # dimensão de y0

vet1n = cbind(rep(1,n))
vet1m = cbind(rep(1,m))

x1 = cbind(1, x)

# Valores iniciais para cálculo das esperanças condicionais

betas = solve(t(x1)%*%x1)%*%t(x1)%*%y
sig2 = (1/(n + m))*(t(y - x1)%*%betas)%*%(y - x1)%*%betas) +
        t(y0 - mean(y0))%*%(y0 - mean(y0)))
lbd = sum(((y - x1)%*%betas)/sig2[1])^3)/n
x0 = (mean(y0) - betas[1])/betas[2]

w = kap = a = as.vector(vet1n)
w0 = kap0 = a0 = as.vector(vet1m)

#Passo-E

d = ((y - x1)%*%betas)^2/sig2[1]
d0 = ((y0 - betas[1]*vet1m - betas[2]*x0*vet1m)^2)/sig2[1]

d = as.vector(d)
d0 = as.vector(d0)

for(j in 1:n){
  kap[j] = (2 + 1)/(2 + d[j])    #Mistura t Student
}

```

```

for(j in 1:m){
  kap0[j] = (2 + 1)/(2 + d0[j])    #Mistura t Student
}

a = (lbd/sqrt(sig2[1]))*(y - x1%%betas)
a0 = (lbd/sqrt(sig2[1]))*(y0 - (betas[1] + betas[2]*x0)*vet1m)

for (k in 1:n){
  phi = dnorm(a[k], mean = 0, sd = 1)
  PHI = pnorm(a[k], mean = 0, sd = 1)
  w[k] = phi/PHI
}

for (k in 1:m){
  phi0 = dnorm(a0[k], mean = 0, sd = 1)
  PHI0 = pnorm(a0[k], mean = 0, sd = 1)
  w0[k] = phi0/PHI0
}

t = lbd*(y - x1%%betas) + sqrt(sig2[1])*w
t2 = (lbd^2)*(y - x1%%betas)^2 + sig2[1]*vet1n +
      lbd*sqrt(sig2[1))*(y - x1%%betas)*w

t0 = lbd*(y0 - betas[1]*vet1m - betas[2]*x0*vet1m) + sqrt(sig2[1])*w0
t20 = (lbd^2)*((y0 - betas[1]*vet1m - betas[2]*x0*vet1m)^2) + sig2[1]*vet1m +
      lbd*sqrt(sig2[1))*(y0 - betas[1]*vet1m - betas[2]*x0*vet1m)*w0

# Valores iniciais dos parâmetros para início das iterações

```

```

alpha = betas[1]
beta = betas[2]
sig2 = sig2[1]

q = y - alpha*vet1n - beta*x
q0 = y0 - (alpha + beta*x0)*vet1m
q = as.vector(q)
q0 = as.vector(q0)

#Loop do algoritmo EM

marcos = 0
i = 0
crit = 0.00001 # Critério de parada
itmax = 10000 # Número máximo de iterações

while (marcos == 0){

  #Passo-M

  alphai = alpha
  betai = beta
  sig2i = sig2
  lbdi = lbd
  x0i = x0

  ti = t
  t2i = t2

```

```

t0i = t0
t20i = t20

kapi = kap
kap0i = kap0
qi = q
q0i = q0

alpha1 = (t(y) - betai*t(x))**kapi + lbdi*(lbdi*t(y) - t(ti) -
        lbdi*betai*t(x))**vet1n + t(y0)**kap0i + ((lbdi^2)*t(y0) -
lbdi*t(t0i) - betai*x0i*t(kap0i))**vet1m - m*betai*x0i*(lbdi^2)
alpha2 = t(kapi)**vet1n + t(kap0i)**vet1m + (lbdi^2)*(n + m)

alpha = alpha1[1]/alpha2[1]

beta1 = t(x)**(diag(kapi)**y - alphai*kapi + (lbdi^2)*(y - alphai*vet1n) -
        lbdi*ti) + x0i*(t(y0)**kap0i + ((lbdi^2)*t(y0) - alphai*t(kap0i) -
lbdi*t(t0i))**vet1m - m*alphai*(lbdi^2))
beta2 = t(x)**(diag(kapi) + (lbdi^2)*diag(n))**x +
        (x0i^2)*(t(kap0i)**vet1m + m*(lbdi^2))

beta = beta1[1]/beta2[1]

sig21 = (t(qi)**diag(kapi) + (lbdi^2)*t(qi) - 2*lbdi*t(ti))**qi +
        t(t2i)**vet1n +
        (t(q0i)**diag(kap0i) + (lbdi^2)*t(q0i) - 2*lbdi*t(t0i))**q0i +
        t(t20i)**vet1m

sig2 = sig21[1]/(2*(n + m))

```

```

lbd1 = t(ti)%%qi + t(t0i)%%q0i
lbd2 = t(qi)%%qi + t(q0i)%%q0i

lbd = lbd1[1]/lbd2[1]

x01 = t(y0)%%kap0i - (alphai*t(kap0i) + lbdi*t(t0i))%%vet1m +
      (lbdi^2)*(t(y0)%%vet1m - alphai*m)
x02 = betai*(t(kap0i)%%vet1m + m*(lbdi^2))

x0 = x01[1]/x02[1]

if((is.na(alpha)) | (is.na(beta)) | (is.na(sig2)) | (is.na(lbd)) |
    (is.na(x0)) | (i > itmax) | (sig2 <= 0)){
  marcos = 1
  print("0 processo iterativo falhou")
}
else{
  q = y - alpha*vet1n - beta*x
  q0 = y0 - (alpha + beta*x0)*vet1m
  q = as.vector(q)
  q0 = as.vector(q0)

  #Passo-E

  d = (diag(q)%%q)/sig2
  d0 = (q0^2)/sig2

  for(j in 1:n){

```

```

    kap[j] = (nu + 1)/(nu + d[j])      #Mistura t Student
}

for(j in 1:m){
    kap0[j] = (nu + 1)/(nu + d0[j])    #Mistura t Student
}

a = (lbd/sqrt(sig2))*q
a0 = (lbd/sqrt(sig2))*q0

for (k in 1:n){
    phi = dnorm(a[k], mean = 0, sd = 1)
    PHI = pnorm(a[k], mean = 0, sd = 1)
    w[k] = phi/PHI
}

for (k in 1:m){
    phi0 = dnorm(a0[k], mean = 0, sd = 1)
    PHI0 = pnorm(a0[k], mean = 0, sd = 1)
    w0[k] = phi0/PHI0
}

t = lbd*q + sqrt(sig2)*w
t2 = (lbd^2)*diag(q)%*%q + sig2*vet1n + lbd*(sqrt(sig2))*q*w

t0 = lbd*q0 + sqrt(sig2)*w0
t20 = (lbd^2)*q0^2 + sig2*vet1m + lbd*(sqrt(sig2))*q0*w0

convergiu = ((abs(alpha - alphai) < crit) & (abs(beta - betai) < crit) &

```

```

        (abs(sig2 - sig2i) < crit)    & (abs(lbd - lbdi) < crit)    &
        (abs(x0 - x0i) < crit))

    }

    if(convergiu == "TRUE"){
        marcos = 1

        resultado = list(alpha=alpha, beta=beta, sigma2=sig2, lambda=lbd, x0=x0)
        iter = list(iterações=i+1)
        print(iter)
        print(resultado)
    }

    i = i + 1
}

```