

Journal of Mathematical Extension  
Vol. 20, No. 3 (2026) (4) 1-27  
ISSN: 1735-8299  
URL: <http://doi.org/10.30495/JME.2026.3677>  
Original Research Paper

## Robust Asymmetric Censored Regression Model and Outlier Contaminated Data

**O. Moatani**

Marv.C., Islamic Azad University

**K. Zare\***

Marv.C., Islamic Azad University

**M. Maleki**

University of Isfahan

**D. Wraith**

Queensland University of Technology

**Abstract.** This study introduces a robust regression modeling framework for censored observations based on the two-piece scale mixture of normal (TP-SMN) class. Using this distributional family enables exact maximum likelihood inference while simultaneously accommodating skewness, the presence of outliers, and censoring mechanisms. The effectiveness of the proposed methodology is assessed through comprehensive simulation experiments and an empirical analysis using real data, which indicate superior estimation precision and enhanced computational performance. These results highlight the advantages of the proposed modeling censored data with symmetric and asymmetric, as well as light- and heavy-tailed, behavior, offering both methodological flexibility and practical applicability.

---

Received: February 2026; Accepted: April 2026

\*Corresponding Author

**AMS Subject Classification:** 62J05; 62J07

**Keywords and Phrases:** Censored regression model, ECME algorithm, ML estimate, TP-SMN.

## 1 Introduction

Linear regression (LR) models providing a interpretability and flexibility of modeling the responses and covariates (Hastie and Tibshirani [6]). This framework allows certain covariates to have a linear effect, and making LR models widely applicable in economics, biostatistics, and social sciences, particularly when modeling real-world data with heterogeneous effects. In ordinary LR models, the normality assumption is common but often too restrictive in the presence of outliers. To address this issue, Rubio and Genton [12] studied a Bayesian linear regression with skew-symmetric error distributions with applications to survival analysis, using the scale mixture of normal (SMN; proposed by Andrews and Mallows [2]) distributions in LR model to handle light- heavy-tailed errors. The SMN family includes important cases light- heavy-tailed distributions, providing a flexible framework for modeling outliers and achieving robust inference. The asymmetric nature of datasets is another major reason to avoid the normality assumption in statistical models, particularly in LR model. Flexible classes of distributions that are both asymmetric and heavy-tailed can simultaneously address the issues of outliers and asymmetry. In addition to the presence of outliers and asymmetric behavior, which depend on the LR model's error distribution, censored and missing observations are also common challenges in the analysis of real-world data. Censoring is a natural phenomenon that occurs in many datasets for various reasons, such as limitations of measurement instruments that cannot record values beyond certain thresholds. LR model with censored observations are encountered across numerous applied fields such as engineering, economics, biomedical studies, medical surveys, and astronomy. In this context [9] investigated LR model for censored datasets under the presence of outliers by adopting a symmetric heavy-tailed distribution family. Their approach provides robust estimation under such conditions; however, asymmetric behavior in the data was not addressed, which could be incorporated in future

studies to simultaneously handle both outliers and skewness. In this study, we adopt such an approach by considering the highly flexible TP-SMN distribution family, introduced by Maleki and Mahmoudi [8], to address the aforementioned issues. The TP-SMN family also encompasses the symmetric SMN distributions and their asymmetric counterparts, which is an advantage, in addition to covering both light- and heavy-tailed distributions. The remainder of this article is structured as follows. Section 2 outlines key characteristics of the TP-SMN family and develops censored LR (CLR) model derived from this framework, along with an EM-type algorithm for computing maximum likelihood (ML) estimators and the corresponding observed information matrix. Section 3 assesses the performance of the proposed methodology through simulation studies and an empirical example, with comparative analysis against well-known censored regression approaches. Section 4 concludes with some final remarks.

## 2 TP-SMN Censored Linear Regression Models

This section formulates CLR model with TP-SMN error structures and describes the use of an EM-type algorithm to carry out maximization of the complete-data likelihood function. First, we provide a brief review of key aspects of the TP-SMN distribution family. The classical SMN family proposed by Andrews and Mallows [2], denoted as  $SMN(\mu, \sigma, \nu)$ , encompasses the symmetric light-tailed normal (N) as well as several symmetric heavy-tailed alternatives, including the Student-t (T) (with the Cauchy distribution obtained when  $\nu = 1$ ), the Slash (SL), and the contaminated normal (CN) distributions. This SMN framework serves as the foundation for constructing the asymmetric TP-SMN family introduced by Maleki and Mahmoudi [8]. Within this extended class, notable members include the asymmetric light-tailed two-piece normal (TP-N) distribution and a variety of asymmetric heavy-tailed counterparts such as the two-piece Student-t (TP-T), which reduces to the two-piece Cauchy when  $\nu = 1$ , along with the two-piece Slash (TP-SL) and two-piece contaminated normal (TP-CN) distributions. Collectively, these models form a flexible family denoted by  $TP-SMN(\mu, \sigma, \gamma, \nu)$ . It

is worth noting that setting  $\gamma = 0.5$  yields the corresponding symmetric SMN distributions as special cases. For a random variable  $X$  following the TP-SMN( $\mu, \sigma, \gamma, \boldsymbol{\nu}$ ) distribution, its probability density function (pdf) and cumulative distribution function (cdf)—both of which are essential components in the development of censored modeling frameworks—are expressed as follows

$$\mathcal{G}(x|\mu, \sigma, \gamma, \boldsymbol{\nu}) = \begin{cases} 2(1 - \gamma)f_{smn}(x|\mu, \sigma(1 - \gamma), \boldsymbol{\nu}) & x \leq \mu, \\ 2\gamma f_{smn}(x|\mu, \sigma\gamma, \boldsymbol{\nu}) & x > \mu, \end{cases}$$

and

$$\mathcal{G}(x|\mu, \sigma, \gamma, \boldsymbol{\nu}) = \begin{cases} 2(1 - \gamma)F_{smn}(x|\mu, \sigma(1 - \gamma), \boldsymbol{\nu}) & x \leq \mu, \\ 2\gamma F_{smn}(x|\mu, \sigma\gamma, \boldsymbol{\nu}) & x > \mu, \end{cases}$$

where  $\mu \in \mathbb{R}$  denotes the location parameter,  $\sigma > 0$  is the scale parameter,  $0 < \gamma < 1$  controls the degree of skewness. The parameter  $\nu$  corresponds to the degrees of freedom that vary across the members of the family, and  $f_{smn}(\cdot|\cdot)$  and  $F_{smn}(\cdot|\cdot)$  denote the probability density and cumulative distribution functions of the underlying SMN distributions, respectively. A detailed discussion of the properties and construction of this family can be found in Ghasami et al. [5].

## 2.1 The model definition

A TP-SMN-CLRM is specified as

$$y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i, \quad \varepsilon_i \sim TP - SMN(0, \sigma, \gamma, \boldsymbol{\nu}), \quad i = 1, \dots, n, \quad (1)$$

where  $y_i$  denotes the left-censored response from the  $i$ -th experimental unit with censoring threshold  $\kappa_i$  defined by

$$y_i = \begin{cases} \kappa_i & y_i \leq \kappa_i, \\ y_i & y_i > \kappa_i; \end{cases} \quad (2)$$

where  $\mathbf{x}_i$  represents a  $p \times 1$  vector of explanatory variable values,  $\boldsymbol{\beta}$  is a  $p \times 1$  fixed parameter vector and the error terms  $\varepsilon_i$ ,  $i = 1, \dots, n$

are assumed to be independent and identically distributed according to the TP-SMN family. The model (1) can be re-expressed the matrix form  $\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$  where  $\mathbf{Y} = (y_1, \dots, y_n)^\top$  is an  $n \times 1$  vector of observed responses,  $\mathbf{X}$  is an  $n \times p$  design matrix with rows  $\mathbf{x}_i$ . Finally, the error vector  $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)^\top$  consists of independent random components drawn from the TP-SMN family of distributions. Let  $\boldsymbol{\Theta} = (\boldsymbol{\beta}^\top, \sigma, \gamma, \boldsymbol{\nu})^\top$  denote the complete vector of unknown parameters associated with the TP-SMN-CLR models. Assume that, among the  $n$  observations, a subset of size  $m$  is subject to censoring. Accordingly, the observed response vector  $\mathbf{y}_{obs}$  can be decomposed into two groups consisting of  $m$  censored observations and  $n - m$  uncensored ones, such that  $\mathbf{y}_{obs} = \{\kappa_1, \dots, \kappa_m; y_{m+1}, \dots, y_n\}$ . Under this setup, the log-likelihood function based on the observed data takes the form

$$\ell(\boldsymbol{\Theta}|\mathbf{y}_{obs}) = \log \left[ \prod_{i=1}^n \{\mathcal{G}(\kappa_i|\mu_i, \sigma, \gamma, \boldsymbol{\nu})\}^{I_i} \{\mathcal{G}(\kappa_i|\mu_i, \sigma, \gamma, \boldsymbol{\nu})\}^{1-I_i} \right]$$

where  $\mu_i = \mathbf{x}_i^\top \boldsymbol{\beta}$ , and  $I_i$  is an indicator variable such that  $I_i = 1$  if  $y_i \leq \kappa_i$ , and  $I_i = 0$  otherwise. Owing to the computational complexity of direct optimization, the maximum likelihood (ML) estimators of the TP-SMN-CLR model parameters are obtained through the ECME algorithm (Liu and Rubin, [7]), as direct maximization can be cumbersome.

## 2.2 ML estimation of TP-SMN-CLR model parameters

This section presents the implementation details of the ECME algorithm employed to compute the ML estimators for the parameters of the TP-SMN-CLR model.

### 2.2.1 ECME algorithm

The hierarchical representation of the TP-SMN-CLR model facilitates ML estimation through the ECM-E algorithm, where some CM-steps use the actual likelihood instead of the expected complete-data log-likelihood. By exploiting the stochastic representation of TP-SMN random variables proposed by Bazrafkan et al. [3], one can derive a hierarchical structure for the TP-SMN-CLR model, summarized in the proposition below.

**Proposition 2.1.** *Let the TP-SMN-CLR model be defined as in equations (1) and (2). Then, for  $i = 1, \dots, n$ , the model admits the following hierarchical representation:*

$$\begin{aligned} Y_i | U_i = u_i, \quad W_i = w_i &\stackrel{iid}{\sim} TN(\mathbf{x}_i^\top \boldsymbol{\beta}, u_i^{-1} w_i^2 \sigma^2) I_B(y_i), \\ P(W_i = w_i) &= \gamma^{\frac{w_i + |w_i|}{2w_i}} (1 - \gamma)^{\frac{w_i + |w_i|}{2w_i}}; \quad w_i = \gamma, -(1 - \gamma), \\ U_i &\stackrel{iid}{\sim} H(u_i, \boldsymbol{\nu}), \end{aligned}$$

where  $TN(\cdot) I_B(\cdot)$  denotes the truncated normal distribution on the interval  $B$ . Specifically, the truncation set is defined as  $B = (\kappa_i, \mathbf{x}_i^\top \boldsymbol{\beta}]$  if  $w_i < 0$ , and  $B = (\mathbf{x}_i^\top \boldsymbol{\beta}, +\infty)$  if  $w_i \geq 0$ .

Let  $\mathbf{Z} = (\mathbf{y}_{obs}, \mathbf{y}_M, u_1, \dots, u_n, w_1, \dots, w_n)^\top$  denote the complete data vector, where  $\mathbf{y}_M = \{y_1, \dots, y_m\}$  corresponds to the missed (censored) component of the response. Then, the associated complete-data log-likelihood function, ignoring constants, can be written as

$$\begin{aligned} \ell_c(\boldsymbol{\Theta} | \mathbf{Z}) &= -\frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n \frac{u_i}{w_i^2} e_i^2 + \sum_{i=1}^n \frac{w_i + |w_i|}{2w_i} \log \gamma \\ &\quad + \sum_{i=1}^n \frac{w_i + |w_i|}{2w_i} \log(1 - \gamma), \end{aligned}$$

where  $e_i = y_i - \mathbf{x}_i^\top \boldsymbol{\beta}$ . Throughout the sequel, the notation  $(r)$  is used to denote the value of a parameter estimate at the  $r$ -th iteration of the algorithm. Within the E-step, the corresponding Q-function is constructed as

$$Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = E \left[ \ell_{cp}(\boldsymbol{\Theta} | \mathbf{Z}) | \mathbf{y}_{obs}, \hat{\boldsymbol{\Theta}}^{(r)} \right].$$

Starting from an initial value  $\hat{\boldsymbol{\Theta}}^{(0)}$ , the Q-function is computed using  $\hat{\boldsymbol{\Theta}}^{(r)}$  for  $\boldsymbol{\Theta}$ , where  $\hat{\boldsymbol{\Theta}}^{(r)} = (\hat{\boldsymbol{\beta}}^{(r)}, \hat{\sigma}^{(r)}, \hat{\gamma}^{(r)}, \hat{\boldsymbol{\nu}}^{(r)})$ . For this purpose, define

$$\alpha_{ki}(\hat{\boldsymbol{\Theta}}^{(r)}) = E_{\hat{\boldsymbol{\Theta}}^{(r)}} \left[ U_i W_i^{-2} Y_i^k | \mathbf{y}_{obs(i)} \right], \quad k = 0, 1, 2,$$

and

$$\zeta_i^{(r)} = E_{\hat{\boldsymbol{\Theta}}^{(r)}} \left[ \frac{|W_i|}{W_i} | \mathbf{y}_{obs(i)} \right],$$

for  $i = 1, \dots, n$ . With these quantities, the E-step is completed by expressing the Q-function as

$$\begin{aligned} Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = & -\frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n \left[ \alpha_{0i}(\hat{\boldsymbol{\Theta}}^{(r)}) (\mathbf{x}_i^\top \boldsymbol{\beta})^2 \right. \\ & \left. - 2\alpha_{1i}(\hat{\boldsymbol{\Theta}}^{(r)}) (\mathbf{x}_i^\top \boldsymbol{\beta}) + \alpha_{2i}(\hat{\boldsymbol{\Theta}}^{(r)}) \right] \\ & + \sum_{i=1}^n \frac{1 + \zeta_i(\hat{\boldsymbol{\Theta}}^{(r)})}{2} \log \gamma + \sum_{i=1}^n \frac{1 - \zeta_i(\hat{\boldsymbol{\Theta}}^{(r)})}{2} \log(1 - \gamma) \end{aligned}$$

At each step, for the censored  $i$ th observation ( $i = 1, \dots, m$ ), we have

$$\alpha_{ki}(\hat{\boldsymbol{\Theta}}^{(r)}) = E_{\hat{\boldsymbol{\Theta}}^{(r)}}[U_i W_i^{-2} Y_i^k | Y_i \leq \kappa_i], \quad k = 0, 1, 2,$$

and  $\zeta_i(\hat{\boldsymbol{\Theta}}^{(r)}) = E_{\hat{\boldsymbol{\Theta}}^{(r)}}[|W_i|/W_i | Y_i \leq \kappa_i]$ , which are derived by the truncated TP-SMN distributions.

At follows for uncensored  $i$ th observation ( $i = m + 1, \dots, n$ ), we have  $\mathbf{y}_{obs(i)} = Y_i \sim TP - SMN(\mathbf{x}_i^\top \boldsymbol{\beta}, \sigma, \gamma, \boldsymbol{\nu})$ . Therefore,

$$\alpha_{ki}(\hat{\boldsymbol{\Theta}}^{(r)}) = y_i^k E_{\hat{\boldsymbol{\Theta}}^{(r)}}[U_i W_i^{-2} | y_i], \quad k = 0, 1, 2,$$

where the  $E_{\hat{\boldsymbol{\Theta}}^{(r)}}[U_i W_i^{-2} | y_i]$  and  $\zeta_i(\hat{\boldsymbol{\Theta}}^{(r)}) = E_{\hat{\boldsymbol{\Theta}}^{(r)}}[|W_i|/W_i]$  are derived by the TP-SMN distributions.

**E-step:** Given the current parameter estimates  $\boldsymbol{\Theta} = \hat{\boldsymbol{\Theta}}^{(r)}$ , the conditional expectations  $\alpha_{ki}(\hat{\boldsymbol{\Theta}}^{(r)})$  for  $k = 0, 1, 2$ , and  $i = 1, \dots, n$  are evaluated using the quantities derived above. Subsequently, the maximization of the Q-function is carried out within the EM framework by solving the system of equations obtained from the first-order optimality condition  $\nabla Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = 0$ , where  $\nabla$  denotes the gradient operator. This procedure leads to the following parameter update expressions:

**CM-step:** The parameter vector is updated from  $\hat{\boldsymbol{\Theta}}^{(r)}$  to  $\hat{\boldsymbol{\Theta}}^{(r+1)}$  by maximizing  $Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)})$  with respect to  $\boldsymbol{\Theta}$ . The resulting parameter

updates are obtained as follows:

$$\begin{aligned}\hat{\beta}^{(r+1)} &= \left( \sum_{i=1}^n \alpha_{0i} \hat{\Theta}^{(r)} \mathbf{x}_i \mathbf{x}_i^\top \right)^{-1} \sum_{i=1}^n \mathbf{x}_i \alpha_{1i} \hat{\Theta}^{(r)}, \\ \hat{\sigma}^{(r+1)} &= \sqrt{\frac{\sum_{i=1}^n \left[ \alpha_{0i} \left( \hat{\Theta}^{(r)} \right) \mu_i^{(r+1)^2} - 2\alpha_{1i} \left( \hat{\Theta}^{(r)} \right) \mu_i^{(r+1)} + \alpha_{2i} \left( \hat{\Theta}^{(r)} \right) \right]}{n}}, \\ \hat{\gamma}^{(r+1)} &= \frac{1}{2n} \sum_{i=1}^n \left( 1 + \zeta_i^{(r)} \right),\end{aligned}$$

**CML-step:** In the final step, the degrees-of-freedom parameter  $\nu$  is updated by maximizing the observed-data log-likelihood directly, leading to the following estimate:

$$\hat{\nu}^{(r+1)} = \operatorname{argmax}_{\nu} \ell \left( \hat{\Theta}_{-\nu}^{(r+1)} | \mathbf{y}_{obs} \right) \quad (3)$$

where  $\hat{\Theta}_{-\nu}^{(r+1)}$  denotes the updated  $\Theta$  without  $\nu$  at the  $(r+1)$ -th stage of algorithm. A more computationally efficient implementation of the CML-step (3) can be achieved using standard optimization routines available in R, such as `optimize` or `optimx` (R Core Team [11]). The algorithm proceeds by alternating between the E-step and the CML-step until convergence is attained. Convergence is determined when a suitable measure of change in the observed-data log-likelihood between successive iterations, for example  $\| \ell(\hat{\Theta}^{(r+1)} | \mathbf{Y}, \mathbf{X}) / \ell(\hat{\Theta}^{(r)} | \mathbf{Y}, \mathbf{X}) - 1 \|$ , falls below a pre-specified tolerance threshold

### 2.3 Approximation of standard errors

This section focuses on deriving standard error estimates for the ML estimators associated with the proposed TP-SMN-CLR model. Under the usual regularity assumptions, the asymptotic covariance matrix of the parameter vector  $\Theta$  can be approximated by the inverse of an empirical information matrix, defined as

$$I_0 = \sum_{i=1}^n \mathbf{s}(y_i | \Theta) \mathbf{s}^\top(y_i | \Theta) - n^{-1} \left( \sum_{i=1}^n \mathbf{s}(y_i | \Theta) \right) \left( \sum_{i=1}^n \mathbf{s}(y_i | \Theta) \right)^\top,$$

where  $\mathbf{s}(y_i|\boldsymbol{\Theta})$  denotes the score contribution associated with the  $i$ -th observation and is given by

$$Q\left(\boldsymbol{\Theta}|\hat{\boldsymbol{\Theta}}^{(r)}\right) = \sum_{i=1}^n Q_i\left(\boldsymbol{\Theta}|\hat{\boldsymbol{\Theta}}^{(r)}\right).$$

By replacing  $\boldsymbol{\Theta}$  with its ML estimate  $\hat{\boldsymbol{\Theta}}$  in  $I_0$ , we obtain an approximation in the form of  $\hat{I}_0 = \sum_{i=1}^n \hat{\mathbf{s}}_i \hat{\mathbf{s}}_i^\top$ , where  $\hat{\mathbf{s}}_i = \mathbf{s}(y_i|\hat{\boldsymbol{\Theta}})$  denotes the individual score contribution evaluated at the estimated parameter values. The score vector can be partitioned as  $\hat{\mathbf{s}}_i = (\hat{\mathbf{s}}_{i,\beta}^\top, \hat{\mathbf{s}}_{i,\sigma}, \hat{\mathbf{s}}_{i,\gamma}, \hat{\mathbf{s}}_{i,\nu}^\top)$ , with components given explicitly by

$$\begin{aligned} \hat{\mathbf{s}}_{i,\beta}^\top &= -\frac{1}{\sigma^2} \left[ \alpha_{0i}(\hat{\boldsymbol{\Theta}}) \mathbf{x}_i^\top \boldsymbol{\beta} - \alpha_{1i}(\hat{\boldsymbol{\Theta}}) \right] \mathbf{x}_i^\top, \\ \hat{\mathbf{s}}_{i,\sigma} &= -\frac{1}{\hat{\sigma}} + \frac{1}{\hat{\sigma}^3} \left[ \alpha_{0i}(\hat{\boldsymbol{\Theta}}) (\mathbf{x}_i^\top \boldsymbol{\beta})^2 - 2\alpha_{1i}(\hat{\boldsymbol{\Theta}}) \mathbf{x}_i^\top \hat{\boldsymbol{\beta}} + \alpha_{2i}(\hat{\boldsymbol{\Theta}}) \right], \\ \hat{\mathbf{s}}_{i,\gamma} &= \frac{1 + \zeta_i(\hat{\boldsymbol{\Theta}})}{2\hat{\gamma}} - \frac{1 - \zeta_i(\hat{\boldsymbol{\Theta}})}{2(1 - \hat{\gamma})} \\ \hat{\mathbf{s}}_{i,\nu}^\top &= \frac{\partial}{\partial \boldsymbol{\nu}} E_{\hat{\boldsymbol{\Theta}}} \left[ \log h(u_i; \boldsymbol{\nu}) | \mathbf{y}_{obs(i)} \right]. \end{aligned}$$

Note that the block part of  $\hat{\mathbf{s}}_{i,\nu}^\top$  involves integral expressions that do not admit closed-form solutions and are therefore evaluated numerically. Finally, standard error estimates for  $\hat{\boldsymbol{\Theta}}$  are obtained as the square roots of the diagonal elements of  $\hat{I}_0$ .

## 2.4 Diagnostic analysis

In practice, the estimates from the regression models may be unduly influenced by outliers. After estimating the proposed CLR model, we examine the fitted model to detect outliers and influential observations. To do so, first we obtain the Hessian Matrix of the model, and then use it in a case-deletion approach in Cook [4].

### 2.4.1 Hessian matrix

By taking derivatives of Q-function at  $\Theta = \hat{\Theta}$ , we obtain the Hessian matrix

$$\ddot{Q}(\hat{\Theta}|\hat{\Theta}^{(r)}) = \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \Theta \partial \Theta^\top} \Big|_{\Theta=\hat{\Theta}}$$

with following non-zero elements:

$$\begin{aligned} \ddot{Q}_{\beta\beta}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \beta \partial \beta^\top} \Big|_{\Theta=\hat{\Theta}} = -\frac{1}{\hat{\sigma}^2} \sum_{i=1}^n \alpha_{0i}(\hat{\Theta}) \mathbf{x}_i \mathbf{x}_i^\top, \\ \ddot{Q}_{\nu\nu}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \nu \partial \nu^\top} \Big|_{\Theta=\hat{\Theta}} = \sum_{i=1}^n \frac{\partial^2}{\partial \nu \partial \nu^\top} E_{\hat{\Theta}} [\log h(u_i; \nu) | \mathbf{y}_{obs(i)}], \\ \ddot{Q}_{\sigma\sigma}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \sigma \partial \sigma} \Big|_{\Theta=\hat{\Theta}} = \frac{n}{\hat{\sigma}^2} - \frac{3}{\hat{\sigma}^4} \\ &\quad \sum_{i=1}^n [\alpha_{0i}(\hat{\Theta})(\mathbf{x}_i^\top \hat{\beta})^2 - 2\alpha_{1i}(\hat{\Theta})(\mathbf{x}_i^\top \hat{\beta}) + \alpha_{2i}(\hat{\Theta})], \\ \ddot{Q}_{\gamma\gamma}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \gamma \partial \gamma} \Big|_{\Theta=\hat{\Theta}} = \sum_{i=1}^n \frac{1 - \zeta_i(\hat{\Theta})}{2(1 - \hat{\gamma})^2} - \sum_{i=1}^n \frac{1 + \zeta_i(\hat{\Theta})}{2\hat{\gamma}^2}, \\ \ddot{Q}_{\beta\sigma}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial^2 Q(\Theta|\hat{\Theta}^{(r)})}{\partial \beta \partial \sigma} \Big|_{\Theta=\hat{\Theta}} \\ &= \frac{2}{\hat{\sigma}^3} \sum_{i=1}^n [\alpha_{0i}(\hat{\Theta}) \mathbf{x}_i (\mathbf{x}_i^\top \hat{\beta} + \mathbf{n}_i^\top \hat{\mathbf{f}}) - \alpha_{1i}(\hat{\Theta}) \mathbf{x}_i]. \end{aligned}$$

### 2.4.2 Case deletion diagnostic method

Case-deletion is a classical diagnostic technique extensively employed to assess the sensitivity of inferential conclusions to the removal of individual observations. To formalize this procedure, the notation  $[-i]$  is used to indicate quantities computed from the dataset with the  $i$ -th observation excluded. For instance, the vector  $\mathbf{y}_{[-i]} = (y_1, \dots, y_{n-1})^\top$  denotes the  $(n-1) \times 1$  response vector obtained after deleting the  $i$ -th case. Let  $\ell_c(\Theta | \mathbf{Z}_{[-i]})$  represent the complete-data log-likelihood evaluated on the reduced dataset. The corresponding ML estimator, obtained

by maximizing the reduced Q-function, is denoted by

$$\hat{\Theta}_{[-i]} = \left( \hat{\beta}_{[-i]}^\top, \hat{\sigma}_{[-i]}, \hat{\gamma}_{[-i]}, \hat{\nu}_{[-i]}^\top \right)^\top$$

where

$$Q_{[-i]} \left( \Theta | \hat{\Theta}^{(r)} \right) = E \left( \ell_c(\Theta | \mathbf{Z}_{[-i]}) | \mathbf{y}_{obs[-i]}, \hat{\Theta}^{(r)} \right).$$

Here  $\hat{\Theta} = \left( \hat{\beta}^\top, \hat{\sigma}, \hat{\gamma}, \hat{\nu}^\top \right)^\top$  denotes the ML estimate of  $\Theta$  obtained from the full dataset via the ECME algorithm. The difference between  $\hat{\Theta}$  and  $\hat{\Theta}_{[-i]}$  serves as a measure of the influence exerted by the  $i$ -th observation in the ML estimates of  $\Theta$ , such that large values indicate that the  $i$ -th case is potentially influential which needs special attention. However, in the case of large data, the procedure requires a large computational effort and raises difficulties. Zhu et al. [13] proposed a computationally efficient one-step pseudo case-deletion approximation to the leave-one-out estimator  $\hat{\Theta}_{[-i]}$ , constructed directly from the full-sample ML estimate  $\hat{\Theta}$ . Specifically, the approximation is defined as

$$\hat{\Theta}_{[-i]}^* = \hat{\Theta} + \left[ \ddot{Q} \left( \hat{\Theta} | \hat{\Theta}^{(r)} \right) \right]^{-1} \dot{Q}_{[-i]} \left( \hat{\Theta} | \hat{\Theta}^{(r)} \right)$$

where  $\ddot{Q}(\hat{\Theta} | \hat{\Theta}^{(r)})$  denotes the Hessian matrix of the Q-function with respect to  $\hat{\Theta}$ , evaluated at the ML estimate, and

$$\dot{Q}_{[-i]} \left( \hat{\Theta} | \hat{\Theta}^{(r)} \right) = \left. \frac{\partial Q_{[-i]}(\Theta | \hat{\Theta})}{\partial \Theta} \right|_{\Theta = \hat{\Theta}}$$

is the individual score vector associated with the reduced dataset. Note that the elements of the individual score vector evaluated at  $\hat{\Theta}$  are

obtained as follows:

$$\begin{aligned}
\dot{Q}_{[-i]\beta}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial Q_{[-i]}(\Theta|\hat{\Theta})}{\partial \beta} \Big|_{\Theta=\hat{\Theta}} \\
&= \frac{1}{\hat{\sigma}^2} \sum_{j \neq i} \left[ \alpha_{1j}(\hat{\Theta}^{(r)}) \mathbf{x}_j - \alpha_{0j}(\hat{\Theta}^{(r)}) (\mathbf{x}_j^\top \hat{\beta}) \mathbf{x}_j \right], \\
\dot{Q}_{[-i]\sigma}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial Q_{[-i]}(\Theta|\hat{\Theta})}{\partial \sigma} \Big|_{\Theta=\hat{\Theta}} \\
&= -\frac{n-1}{\hat{\sigma}} + \frac{1}{\hat{\sigma}^3} \sum_{j \neq i} \left[ \alpha_{1j}(\hat{\Theta}^{(r)}) \mathbf{x}_j - \alpha_{0j}(\hat{\Theta}^{(r)}) (\mathbf{x}_j^\top \hat{\beta}) \mathbf{x}_j \right], \\
\dot{Q}_{[-i]\gamma}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial Q_{[-i]}(\Theta|\hat{\Theta})}{\partial \gamma} \Big|_{\Theta=\hat{\Theta}} = \sum_{j \neq i} \frac{1 - 2\hat{\gamma} + \zeta_i(\hat{\Theta}^{(r)})}{2\hat{\gamma}(1 - \hat{\gamma})}, \\
\dot{Q}_{[-i]\nu}(\hat{\Theta}|\hat{\Theta}^{(r)}) &= \frac{\partial Q_{[-i]}(\Theta|\hat{\Theta})}{\partial \nu} \Big|_{\Theta=\hat{\Theta}} = \sum_{j \neq i} \frac{\partial}{\partial \nu} E_{\hat{\Theta}}[\log h(u_j; \nu) | \mathbf{y}_{obs[j]}].
\end{aligned}$$

A simple measure of the distance between  $\hat{\Theta}$  and  $\hat{\Theta}_{[-i]}$  for detecting influential observations is the **Q-displacement** (Zhu et al., [13]):

$$QD_i = 2 \left\{ Q(\hat{\Theta}|\hat{\Theta}^{(r)}) - Q(\hat{\Theta}_{[-i]}|\hat{\Theta}^{(r)}) \right\}, \quad i = 1, \dots, n.$$

Zhu et al. [13] also introduced the generalized Cook's distance:

$$GD_i = (\hat{\Theta} - \hat{\Theta}_{[-i]})^\top \left[ -\ddot{Q}(\hat{\Theta}|\hat{\Theta}^{(r)}) \right] (\hat{\Theta} - \hat{\Theta}_{[-i]}), \quad i = 1, \dots, n.$$

which can be approximated by (for  $i = 1, \dots, n$ ):

$$AGD_i = \dot{Q}_{[-i]}(\hat{\Theta}|\hat{\Theta}^{(r)})^\top \left[ -\ddot{Q}(\hat{\Theta}|\hat{\Theta}^{(r)}) \right]^{(-1)} \dot{Q}_{[-i]}(\hat{\Theta}|\hat{\Theta}^{(r)}).$$

### 3 Numerical Results

This section examines the empirical behavior of the proposed TP-SMN-CLR model through an extensive Monte Carlo simulation study, complemented by an analysis of a real dataset previously investigated in the

literature. All computational procedures, including the implementation of the ECME algorithm, were carried out in the R computing environment (R Core Team, [11]), using a workstation equipped with an Intel Core i7-760 processor operating at 2.8 GHz.

### 3.1 Simulation study

The finite-sample properties of the proposed TP-SMN-CLR model and the associated estimation procedure are investigated through a comprehensive Monte Carlo simulation framework. Specifically, three simulation experiments are conducted under the data-generating mechanism defined in model (1). Left-censored samples are generated across multiple experimental settings, considering censoring proportions of 0%, 10%, 20%, and 30%, together with sample sizes  $n \in \{200, 300, 400, 600\}$ . For each combination of sample size and censoring rate, 500 independent datasets are simulated from the TP-SMN-CLR model under four distinct distributional configurations for the error term, namely: the two-piece normal CLR (TP-N-CLR), the two-piece Student-t CLR with  $(\nu = 4)$  (TP-T-CLR), the two-piece slash CLR with  $(\nu = 4)$  (TP-SL-CLR), and the two-piece contaminated-normal CLR with contamination parameters  $(\boldsymbol{\nu} = (0.2, 0.2)^\top)$ , (TP-CN-CLR) models. We performed all Monte Carlo (MC) simulations by setting  $(\boldsymbol{\beta} = (1, 5)^\top)$ ,  $\gamma = 0.45$ , (weak asymmetry) and  $\gamma = 0.9$  (strong asymmetry) and  $\sigma = 1$ , with  $\mathbf{x}_i = (x_{i1}, x_{i2})^\top$  for which  $x_{i1} \sim U(0, 1)$  and  $x_{i2} \sim U(1, 2)$  are generated independently.

#### 3.1.1 Recovery of parameters and asymptotic properties

This section summarizes the results of the Monte Carlo simulations conducted for the TP-SMN-CLR model with different censored levels of the observations in each dataset. In each generated sample  $k = (1, \dots, 500)$  with length 300, we obtain the proposed ML estimates of  $\theta$  (denoting  $\boldsymbol{\beta}, \sigma, \gamma$  and  $\boldsymbol{\nu}$ ) denoted by  $\hat{\theta}_k$ . To assess the empirical behavior of the estimators, we compute the Monte Carlo mean and dispersion measures

across replications. Specifically, the Monte Carlo average is defined as

$$\text{MC-Mean} = \hat{\theta} = \frac{1}{500} \sum_{k=1}^{500} \hat{\theta}_k,$$

while the associated Monte Carlo variability is quantified through the empirical standard deviation

$$\text{SD} = \sqrt{\frac{1}{500} \sum_{k=1}^{500} (\hat{\theta}_k - \hat{\theta})^2}.$$

The numerical summaries corresponding to the various TP-SMN-CLR model specifications are reported in Tables 1 and 2. These results highlight the estimation accuracy and robustness of the proposed models across scenarios exhibiting both weak and strong skewness, as well as across different sample sizes  $n = 200, 400$ , and  $600$ .

### 3.1.2 Asymptotic properties

This part of the simulation study is devoted to assessing the asymptotic properties of the proposed estimation procedure for the TP-SMN-CLR model. We generated left-censored samples from the proposed TP-SMN-CLR model with at various censoring levels and moderate asymmetry  $\gamma = 0.75$ , with several sample sizes ( $n = 200, 300, 400$  and  $600$ ), over the 500 replications. For each replication, maximum likelihood estimates of the regression coefficients  $\beta = (\beta_1, \beta_2)^\top$  are obtained. The bias and mean squared error (MSE) are then employed as summary measures of estimator accuracy. Specifically, for  $i = 1, 2$ , these quantities are defined as

$$\text{Bias}(\beta_i) = \frac{1}{500} \sum_{k=1}^{500} (\beta_i - \hat{\beta}_{ik}), \quad \text{MSE}(\beta_i) = \frac{1}{500} \sum_{k=1}^{500} (\beta_i - \hat{\beta}_{ik})^2, \quad i = 1, 2.$$

where  $\hat{\beta}_{ik}$  denotes the ML estimate of  $\beta_i$  obtained from the  $k$ -th simulated sample.

According to Figures 1 and 2, as the sample size increased, it is naturally observed that, in all of the censoring levels, the Bias and MSE of

ML estimates approach to zero. This convergence toward zero provides empirical support for the consistency and favorable asymptotic behavior of the proposed ML estimators within the TP-SMN-CLR framework.

### 3.1.3 Robustness of estimates to the outlier

It is important to highlight the robustness property of the TP-SMN-CLR model for attenuating the effects of the outlier observations. We generate 300 samples of TP-N-CLR model from the proposed TP-SMN-CLR model with sample sizes  $n = 250$  with moderate censoring levels 10% and 20% and moderate asymmetry  $\gamma = 0.75$ . In order to detect the outliers, in each sample, we ignored observation  $Y_{150}$  and considered its replacement by  $Y_{150}^* = Y_{150} + \pi$  with  $\pi = 1, 2, \dots, 10$ . Then, for each original and modified dataset, the proposed TP-SMN-CLR model were fitted and the ML estimates of  $\theta$  (i.e.  $\beta, \sigma$  and  $\gamma$ ) in generated sample  $k = (1, \dots, 500)$  were denoted by  $\hat{\theta}_k$  and  $\hat{\theta}_k(\pi)$  average of relative changes (ARC), given by

$$\text{ARC}_{\pi}(\theta) = \frac{1}{500} \sum_{k=1}^{500} \left| \frac{\hat{\theta}_k(\pi) - \hat{\theta}_k}{\hat{\theta}_k} \right|,$$

was evaluated, this measure increases for larger outliers (See Figure 3). In other words, the TP-T-CLR, TP-SL-CLR and TP-CN-CLR models (as heavy-tailed distributions) are more robust and reasonable than TP-N-CLR model to accommodate atypical observations. Note that we checked the robustness for various values of censoring and skewness and found similar results.

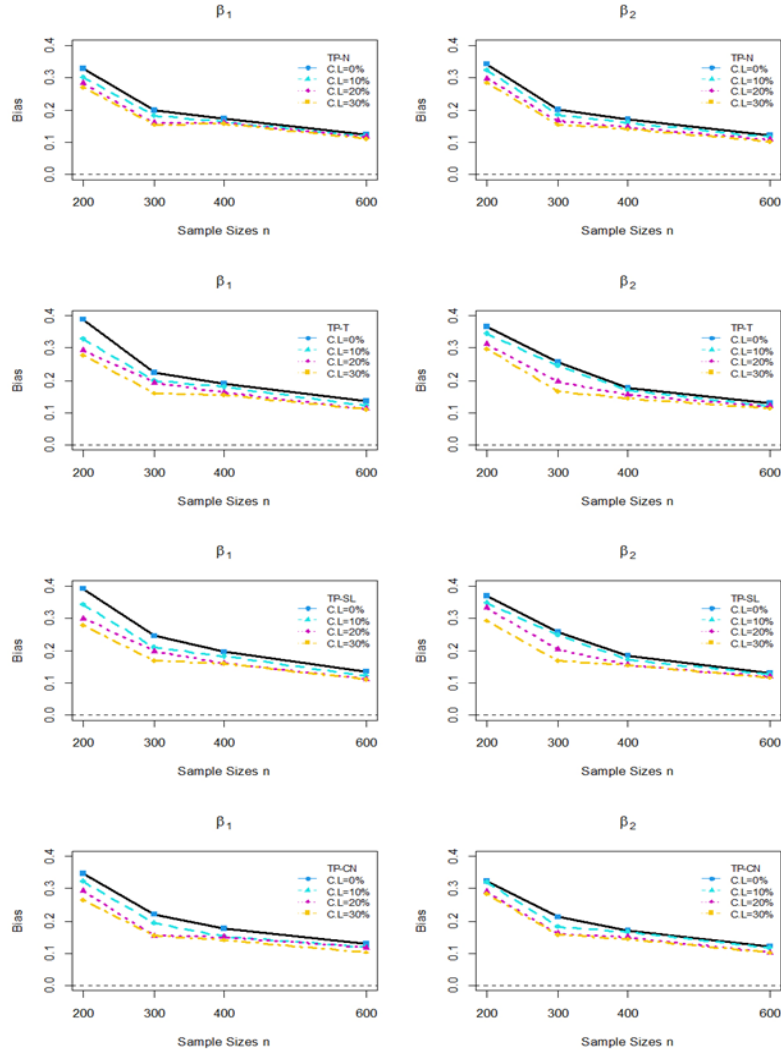
## 3.2 Real application: Wage rate dataset

In this section, we analyze a real data set which has previously been used in the literature (Mroz [10]) to examine models allowing for censoring. The empirical illustration is based on data drawn from the Panel Study of Income Dynamics (PSID), conducted by the University of Michigan. The sample comprises 753 married White women observed in 1975, whose ages range from 30 to 60 years. Among these individuals, 428 reported positive labor market participation during the reference year.

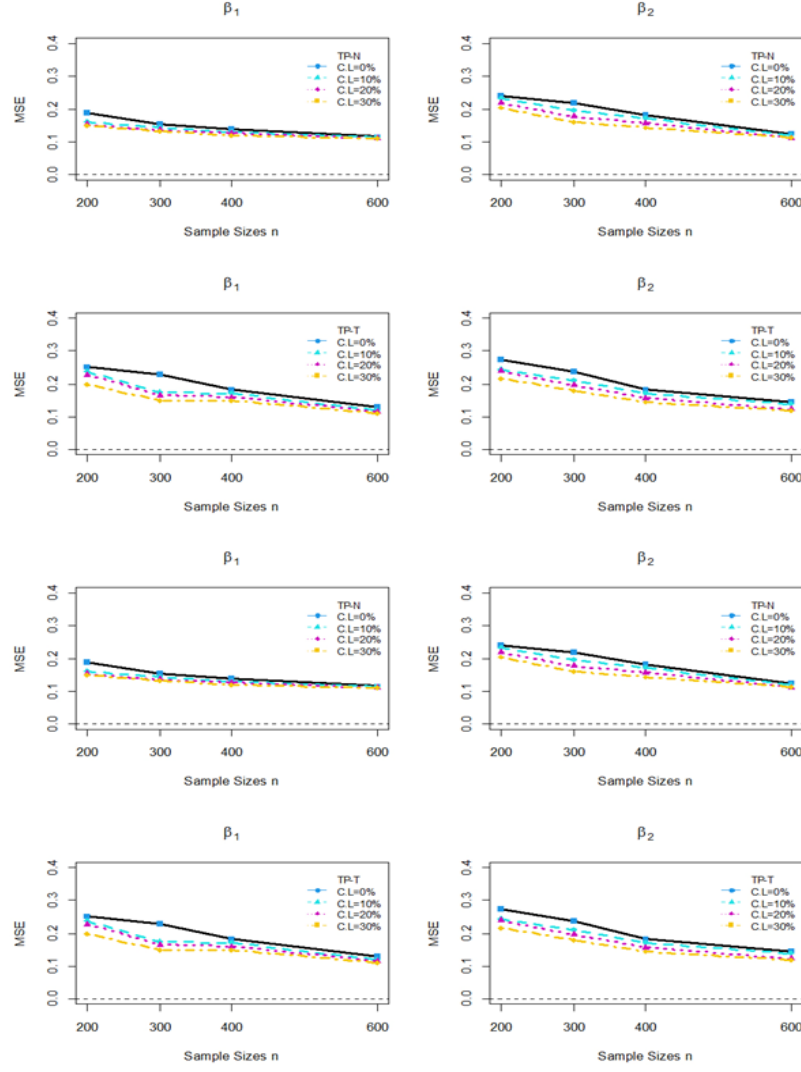


**Table 2:** Monte-Carlo average (MC-Mean) and standard deviation (S.D.) of the ML estimates TP-SMN-CLR model (with  $\gamma = 0.9$ ; strongly skewed ), in 500 simulated samples for different censored levels (C.L.)

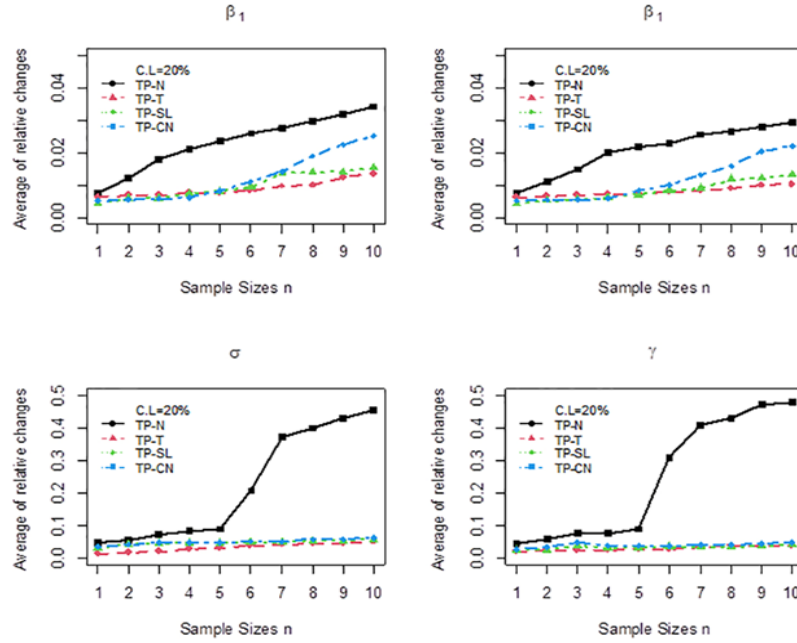
[illegible]



**Figure 1:** Bias for  $\beta_1$  and  $\beta_2$  for different sample sizes and levels of censoring in TP-SMN-CLR model.



**Figure 2:** MSE for  $\beta_1$  and  $\beta_2$  for different sample sizes and levels of censoring in TP-SMN-CLR model.



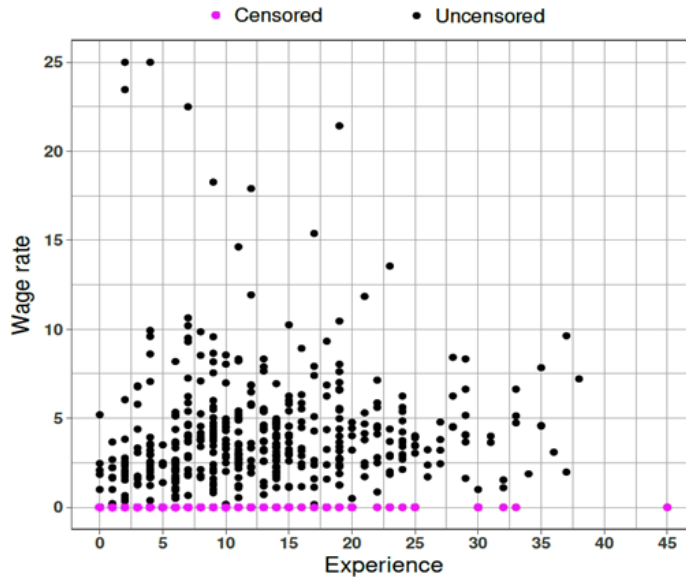
**Figure 3:** Average relative changes of the ML estimates of the TP-SMN-CLR model parameters at censored level 20%.

The outcome of interest is the wage rate, measured as average hourly earnings. For women who did not participate in paid work during 1975, the recorded wage equals zero, resulting in a naturally left-censored response variable. In addition to the wage outcome, several socioeconomic and demographic covariates are considered, including:

- $y$  : Wage rates (response variable);
- $x_1$  : Wife's age;
- $x_2$  : Years of schooling;
- $x_3$  : Husband's hours worked;
- $x_4$  : Husband's wage in dollars;

- $x_5$  : Tax rate faced by the wife;
- $x_6$  : Number of children younger than six years old in the household;
- $x_7$  : Number of children between the ages of six and nineteen;

The relation between the wage rate and the number of years the wife worked since age 18 exhibits can be seen in Figure 4.



**Figure 4:** Wage rates vs. number of years the wife worked (Experience).

### 3.2.1 Analysis of model fitting

We fitted robust and symmetrically/asymmetrically models to this data with the proposed estimates and corresponding standard errors calculated via the empirical information matrix, provided in Tables 3 and 4. The  $\ell(\hat{\Theta})$  and AIC ([?]) values in Tables 3 and 4, show that asymmetrical members of the TP-SMN-CLR model are preferred compared to

**Table 3:** ML estimates of fitted SMN–CLR model to the wage rate data.

| Parameter            | N–CLR   |          | T–CLR   |          | SL–CLR         |          | CN–CLR  |          |
|----------------------|---------|----------|---------|----------|----------------|----------|---------|----------|
|                      | ML est. | SE.      | ML est. | SE.      | ML est.        | SE.      | ML est. | SE.      |
| $\beta_1$            | 0.7512  | (0.0800) | 0.6691  | (0.0654) | 0.6492         | (0.0731) | 0.6621  | (0.0728) |
| $\beta_2$            | -0.0599 | (0.0211) | -0.0711 | (0.0211) | -0.0728        | (0.0199) | -0.0707 | (0.0200) |
| $\beta_3$            | -0.0009 | (0.0004) | -0.0005 | (0.0003) | -0.0004        | (0.0003) | -0.0004 | (0.0003) |
| $\beta_4$            | -0.1796 | (0.0713) | -0.1422 | (0.0539) | -0.1501        | (0.0600) | -0.1601 | (0.0443) |
| $\beta_5$            | -8.2911 | (2.9736) | -5.6532 | (3.1029) | -5.1082        | (3.2212) | -6.2221 | (3.1001) |
| $\beta_6$            | -1.8409 | (0.3726) | -1.7042 | (0.3025) | -1.7019        | (0.3208) | -1.8376 | (0.3211) |
| $\beta_7$            | 0.3401  | (0.1594) | 0.1998  | (0.1019) | 0.1938         | (0.1071) | 0.2101  | (0.1037) |
| $\sigma$             | 3.9842  | (1.2401) | 2.3432  | (0.6301) | 2.0012         | (0.3028) | 2.6102  | (0.6809) |
| $\nu$                | -       | -        | 2.8651  | (0.4213) | 1.1231         | (0.1265) | 0.1101  | (0.0031) |
| $\tau$               | -       | -        | -       | -        | -              | -        | 0.1201  | (0.0110) |
| $\ell(\hat{\Theta})$ | -1346.5 |          | -1299.1 |          | <b>-1293.4</b> |          | -1294.6 |          |
| AIC                  | 2709    |          | 2616.2  |          | <b>2604.8</b>  |          | 2609.2  |          |

the corresponding symmetrical members (i.e., SMN–CLR model). The proposed TP–T–CLR model has the best performance between all of the proposed symmetrical and asymmetrical models highlighting that it is important to consider models that allow for both heavy tailed and asymmetric behavior.

### 3.2.2 Diagnostic checking

To assess influential observations, we apply the diagnostic measures from Subsection 2.4 to the dataset. In the ECME algorithm, the scale mixers are estimated as  $\alpha_{o_i}(\hat{\Theta}) = \hat{u}_i \hat{w}_i^{-2}$ , and plotted against the Mahalanobis distance  $d_i^2 = (y_i - \mathbf{x}_i^\top \hat{\beta})^2 / \hat{\sigma}^2$  for  $i = 1, \dots, 753$  (See Figure 5). For the light-tailed Normal–CLR model, the weights remain close to 1 (red segments in Figure 5). The figure shows that the weights decrease as the Mahalanobis distance increases; large  $d_i^2$  produce small weights. Thus, the heavier-tailed TP–T–CLR assigns smaller weights to potentially influential observations, reducing their impact on parameter estimation. In terms of detecting influential observations, the index plots for the approximate generalized Cook’s distance  $AGD_i$ ,  $i = 1, \dots, 753$ , are shown in Figure 6. High values of  $AGD_i$  suggest that the  $i$ -th observation has

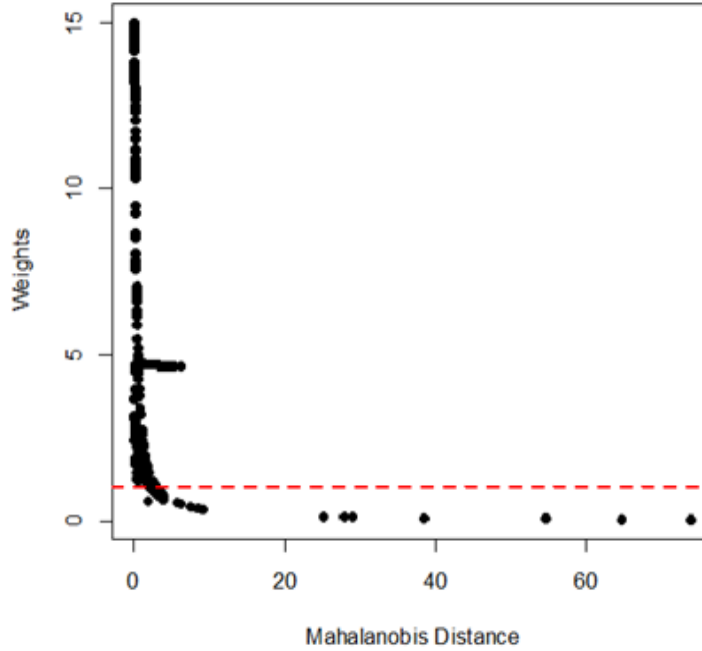
**Table 4:** ML estimates of fitted TP-SMN-CLR model to the wage rate data.

| Parameter            | TP-N-CLR |          | TP-T-CLR      |          | TP-SL-CLR |          | TP-CN-CLR |          |
|----------------------|----------|----------|---------------|----------|-----------|----------|-----------|----------|
|                      | ML est.  | SE.      | ML est.       | SE.      | ML est.   | SE.      | ML est.   | SE.      |
| $\beta_1$            | 0.5908   | (0.0733) | 0.4653        | (0.0987) | 0.4421    | (0.0701) | 0.4701    | (0.0726) |
| $\beta_2$            | 0.0034   | (0.0008) | 0.0020        | (0.0005) | 0.0021    | (0.0008) | 0.0022    | (0.0007) |
| $\beta_3$            | 0.0009   | (0.0000) | 0.0002        | (0.0000) | 0.0002    | (0.0000) | 0.0002    | (0.0000) |
| $\beta_4$            | -0.0043  | (0.0007) | -0.0029       | (0.0007) | -0.0031   | (0.0006) | -0.0028   | (0.0007) |
| $\beta_5$            | -5.1310  | (1.0209) | -3.5647       | (0.8635) | -3.7252   | (0.9817) | -3.7625   | (0.8999) |
| $\beta_6$            | -0.0312  | (0.0076) | -0.0290       | (0.0061) | -0.0254   | (0.0051) | -0.0276   | (0.0061) |
| $\beta_7$            | 0.4107   | (0.0332) | 0.1208        | (0.0123) | 0.1209    | (0.0141) | 0.12130   | (0.0142) |
| $\sigma$             | 4.0726   | (1.0716) | 2.7827        | (0.8262) | 2.9723    | (0.8010) | 2.8972    | (0.9001) |
| $\gamma$             | 0.3526   | (0.0091) | 0.3016        | (0.0065) | 0.3100    | (0.0098) | 0.3122    | (0.0099) |
| $\nu$                | -        | -        | 2.0811        | (0.3901) | 1.1198    | (0.4001) | 0.1002    | (0.0098) |
| $\tau$               | -        | -        | -             | -        | -         | -        | 0.1187    | (0.0102) |
| $\ell(\hat{\Theta})$ | -1098.4  |          | <b>-976.3</b> |          | -997.2    |          | -1001.3   |          |
| AIC                  | 2214.8   |          | <b>1972.6</b> |          | 2014.4    |          | 2024.6    |          |

an impact on the ML estimates. According to the top left panel for the TP-N-CLR model, observations #185, #210, #349, #357, #366, #369, #394, #408 and #692 are potentially influential in the ML estimates, but for distributions with heavy tails, i.e., TP-T-CLR, TP-SL-CLR and TP-CN-CLR models, these observations are no longer influential.

## 4 Conclusion

This work introduces a class of CLR models for censored responses under the TP-SMN family of distributions, which provides a highly adaptable framework capable of accommodating both symmetric and asymmetric behaviors, as well as light- and heavy-tailed error structures. For parameter estimation of the proposed TP-SMN-CLR model, an ECME-type algorithm is carried out that yields analytical forms for parameter estimation and greater accuracy which is a main advantage of the proposed models. The empirical performance of the proposed models is demonstrated through extensive Monte Carlo simulations and an application to real data, highlighting their robustness and practical effectiveness across a variety of scenarios. An R package is currently in development for the



**Figure 5:** Estimated weights  $u_i w_i^{-2}$  versus Mahalanobis distance  $d_i^2$  for fitted TP-T-CLR model on the Wage rate data.

proposed models to allow users to model censored data with greater accuracy and without approximations for the estimation of the model parameters. Future work will extend this framework by formulating a Bayesian version of the TP-SMN-CLR model, with the aim of further increasing modeling flexibility while offering potential gains in computational efficiency and inferential richness.

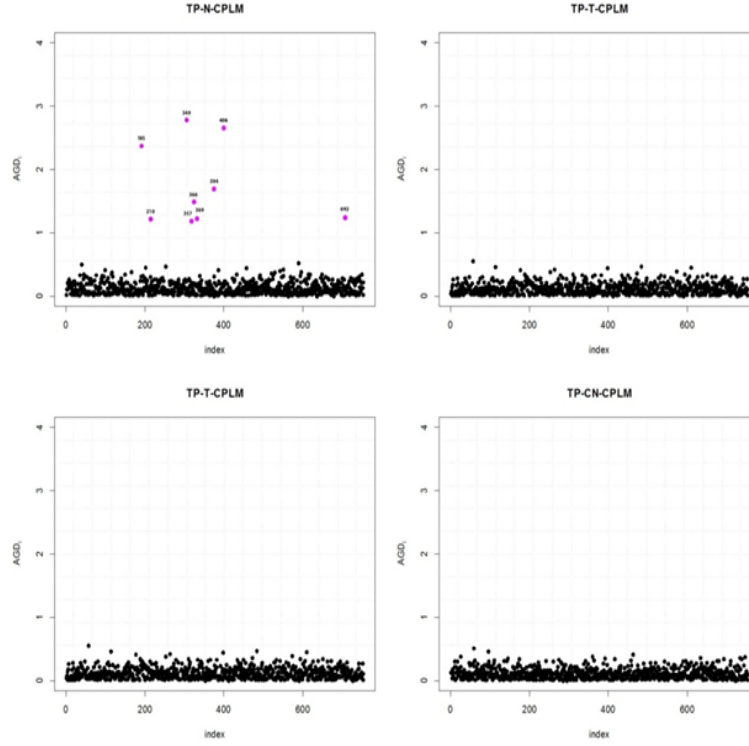
#### **Compliance with Ethical Standards**

#### **Funding:**

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

#### **Conflict of Interest:**

The authors declare that they have no conflict of interest.



**Figure 6:** Approximate generalized Cook's distance ( $AGD_i$ ) for several TP-SMN-CLR model.

### **Ethical Conduct:**

This article does not contain any studies with human participants or animals performed by any of the authors.

### **Data Availability Statement:**

The datasets analyzed in the current study are publicly available within R software packages. Specifically, the data are included in the R packages **AER** (Mroz), **sampleSelection** (Mroz87), **wooldridge** (mroz), and **survival**, all of which are freely accessible through the Comprehensive R Archive Network (CRAN) at <https://cran.r-project.org>.

## References

- [1] H. Akaike, A new look at the statistical model identification, *IEEE Trans. Autom. Control*, 19 (1974), 716–723.
- [2] D. F. Andrews and C. L. Mallows, Scale mixtures of normal distributions, *J. R. Stat. Soc. Ser. B*, 36 (1974), 99–102.
- [3] M. Bazrafkan, K. Zare, M. Maleki and Z. Khodadadi, Partially linear models based on heavy-tailed and asymmetrical distributions, *Stoch. Environ. Res. Risk Assess.*, 36 (2022), 1243–1253.
- [4] R. D. Cook, Detection of influential observation in linear regression, *Technometrics*, 19 (1977), 15–18.
- [5] S. Ghasami, M. Maleki and Z. Khodadadi, Leptokurtic and platykurtic class of robust symmetrical and asymmetrical time series models, *J. Comput. Appl. Math.*, 376 (2020), 112806.
- [6] T. J. Hastie and R. J. Tibshirani, *Generalized Additive Models*, Monographs on Statistics and Applied Probability, vol. 43. Chapman and Hall, London, 1990.
- [7] C. Liu and D. B. Rubin, The ECME algorithm: A simple extension of EM and ECM with faster monotone convergence, *Biometrika*, 81 (1994), 633–648.
- [8] M. Maleki and M. R. Mahmoudi, Two-piece location-scale distributions based on scale mixtures of normal family, *Commun. Stat. Theory Methods*, 46 (2017), 12356–12369.
- [9] M. B. Massuia, A. M. Garay, C. R. B. Cabral, and V. H. Lachos, Bayesian analysis of censored linear regression models with scale mixtures of skew-normal distributions, *Statistics and Its Interface*, 10 (2017), 425–439.
- [10] T. A. Mroz, The sensitivity of an empirical model of married women’s hours of work to economic and statistical assumptions, *Econometrica*, 55 (1987), 765–799.

- [11] R Core Team, *R: A Language and Environment for Statistical Computing*, Version 4.3.3. R Foundation for Statistical Computing, Vienna, Austria, 2024. Available online: <https://www.R-project.org/> (accessed on 29 February 2024).
- [12] F. J. Rubio and M. G. Genton, Bayesian linear regression with skew-symmetric error distributions with applications to survival analysis, *Stat. Med.*, 35 (2026), 2441–2454.
- [13] H. Zhu, S. Y. Lee, B. C. Wei and J. Zhou, Case-deletion measures for models with incomplete data, *Biometrika*, 88 (2001), 727–737.

**Omolbanin Moatani**

Department of Statistics,  
Ph.D. Student of Statistics,  
Marv.C., Islamic Azad University, Marvdasht, Iran.  
E-mail: om.moatani@iau.ac.ir

**Karim Zare**

Department of Statistics,  
Associate Professor of Statistics,  
Marv.C., Islamic Azad University, Marvdasht, Iran.  
E-mail: karim.zare@iau.ac.ir

**Mohsen Maleki**

Department of Statistics, Faculty of Mathematics and Statistics,  
Assistant Professor of Statistics,  
University of Isfahan, Isfahan, Iran.  
E-mail: m.maleki.stat@gmail.com

**Darren Wraith**

Institute of Health and Biomedical Innovation,  
Associate Professor of Statistics,  
Queensland University of Technology, Queensland, Australia.  
E-mail: d.wraith@qut.edu.au