

Bayesian Tobit Single Index model

Taha Alshaybawee

taha.alshaybawee@qu.edu.iq

Ahlam Noor Mohammed AL-Shareefi

ahlamnwrn@gmail.com

University of AL-Qadisiyah

Corresponding Author : Ahlam Noor Mohammed AL-Shareefi

Abstract : Single index model (SIMo) has become one of the most important semiparametric to overcome the dimensionality problem. This model work on summarizes the effects of the independent variables within a single variable and refers to them as the index. In this paper, the Bayesian hierarchical model is constructed to estimate the parameters and the unknown nonparametric function for the single index when the response variable is censored. In addition, to get a head start on finding the unknown nonparametric link function, we assume it follows the Gaussian process distribution. The Laplace distribution will be used as a prior distribution for the coefficient vector β when variables are being selected. The performance of the suggested technique is evaluated by applying it to simulation examples and actual data.

Keywords: Bayesian Inference, Gaussian process distribution, Tobit regression, Markov chain Monte Carlo methods; Variable selection.

Introduction: The single index model (SIMo) is a semiparametric model that has the potential to reduce dimensions; this makes it a popular and significant model in statistics and applied quantitative fields like econometrics and psychometrics. A single index model is used to summarize the effects of all the explanatory variables $X = (X_1, \dots, X_p)'$ in one variable called index. For this reason, the SIMo provides an opportunity to widen the scope of both the generalized linear model and the multidimensional regression model. This is done in order to overcome the difficulty of dealing with dimensionality. The SIMo represented by $E(Y | X) = \delta(X'\beta)$, where $\beta = (\beta_1, \dots, \beta_p)'$. Therefore, the fundamental assessment of interest in the SIMo involves two stages: Initially, we calculate the coefficient vector β . Subsequently, we use the index values to estimate the link function δ in a nonparametric manner. Though Bayesian techniques have been effective with several other nonlinear regression models, the Bayesian framework for the SIMo has not been often presented up to this point (Denison, Holmes, Mallick, and Smith, (2002). Using the B-spline approximation for the link function and the Fisher-von Mises prior distribution for the index vector with an additional regularization, the Bayesian approach to the SIMo was recently investigated in a seminal paper by Antoniadis, Grégoire, and McKeague (2004). Wang (2009) most recently examined a framework akin to Antoniadis et al. (2004) with an application to the problem of variable selection. To the best of our knowledge, Bayesian techniques have not been thoroughly explored for SIMs, with the exception of some recent research that mostly used splines. In this context, further research must be done on substitute approaches for the SIM. Using a Gaussian process (GP) prior, the Bayesian technique was used by Choi et al. (2011) and Gramacy and Lian (2012). One benefit of using GPs as the prior for the link function is that the index vector does not need to be normalized to have a unit norm. This simplifies the choice of prior and, as mentioned in Gramacy and Lian (2012), the sampling technique also becomes simpler not at all.

A Bayesian estimation and variable selection strategy is proposed by Sami, Z., and Alshaybawee, T. for the purpose of estimating the parameters and link function of a single index logistic regression model, as well as selecting the variables that are considered to be the most significant. The Laplace distribution is used as the prior for the coefficients vector, whereas a Gaussian process is chosen as the prior for the unknown link function. A hierarchical Bayesian lasso semiparametric logistic regression model is created, and a Markov Chain Monte Carlo (MCMC) approach is used for posterior inference. Regression analysis is a statistical method for characterizing the nature of the connection between the dependent variable and the explanatory variables. When the values of the explanatory components are known, regression analysis may be used to anticipate the dependent variable's value. The work of Tobin and the Tobit models has attracted significant interest in several practical fields of applied research, including econometrics, agriculture, genetics, ecology, and environmental science. It is a great method for assessing the connection between an outcome variable and a set of explanatory factors. The observable response y is expected to be connected to the unobserved latent variable y^* in this model since it is predicted that this may occur.

$$y = \max\{y^*, y_0\}$$

Censoring point y^0 is usually set to 0 without loss of generality. Tobit models have several applications in econometrics, including household health and durable goods spending, and labor force hours. Tobit models provide a versatile approach to addressing the estimated challenges that are associated with these kinds of circumstances. For an excellent analysis of Tobit models, we recommend that the readers check out Amemiya.

Levy and Aradillas-Lopez et al. are two examples of researchers who have researched nonparametric and semiparametric models for Tobit regression. A famous type of semiparametric model, on the other hand, is the single-index model (SIM). This is because it can avoid the problems of the curse of dimensionality and the fact that fully nonparametric methods are hard to understand. The methodology that has been suggested by Zhao, K., & Lian, H. (2015) is an extension of the single-index model (SIM) that is extensively used in the Tobit regression scenario.

In this paper, the Bayesian hierarchical model is constructed to estimate the parameters and the unknown nonparametric function for the single index when the response variable is censored. In addition to this, the Gaussian process distribution is taken into consideration as a prior for the nonparametric link function that is unknown.

2. Formulation of Tobit Single Index Model and assumptions

Consider the following Tobit single index model

$$y_i^* = \delta(X_i' \beta) + e_i; \quad i = 1, 2, \dots, n \quad (1)$$

Where e_i is an identically and independently distributed random variable with zero mean $N(0, \sigma_e^2)$. Therefore, the distribution function for the error term is given as follows:

$$\pi(e_i | \sigma_e^2) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma_e^2}} \exp\left\{-\frac{1}{2\sigma_e^2} e_i^2\right\} \quad (2)$$

So, the likelihood function for y_i^* , $i=1, 2, \dots, n$ will be given as:

$$\begin{aligned} \pi(y_i^* | X, \delta, \beta, \sigma^2) &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left\{-\frac{1}{2\sigma^2} (y_i^* - \delta(X_i' \beta))^2\right\} \\ &= (2\pi\sigma_e^2)^{-\frac{n}{2}} \exp\left\{-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i^* - \delta(X_i' \beta))^2\right\} \end{aligned} \quad (3)$$

Furthermore, $\delta(\cdot)$ is the unknown univariate function, $X_i' = (\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p)'$ is the p -dimensional of explanatory variables, β is the coefficient vector and y_i^* is the latent response variable. The Tobit (censored) single index model has the following structured formulation,

$$y_i = \begin{cases} y_i^* & \text{if } y_i^* > 0 \\ 0 & \text{if } y_i^* \leq 0 \end{cases} \quad (4)$$

Where y_i^* is the unobserved response variable and y_i is the observed random variable.

Based on Chio et al, (2011) and Gramcy and Lain (2012). A Gaussian process will be used as a prior distribution for the nonparametric function $\delta(\cdot)$. Therefore, the gaussian process with mean zero and square exponential covariance function $GPr(0, C(\cdot, \cdot))$. So, Gaussian process (GP) with zero mean and covariance matrix can be shown as,

$$C(x_i, x_j) = \tau \exp\left\{-\frac{(x_i - x_j)^2}{v}\right\} \quad (5)$$

Where τ and v are hyperparameters. The prior distribution to the nonparametric function will be written as

$$\pi(\delta | \beta, \tau) = \{ \det \det(C_n) \}^{-\frac{1}{2}} \exp\left\{-\frac{1}{2} (\delta' C_n^{-1} \delta)\right\} \quad (6)$$

Inverse gamma distribution will be considered as a prior to the hyperparameter τ . following Park and Casella (2008), symmetric Laplace distribution will be set as a prior distribution to the vector parameters β and can be shown as follows

$$\pi(\lambda, \sigma^2) = \frac{\lambda}{2\sqrt{\sigma^2}} e^{-\frac{\lambda|\beta_j|}{\sqrt{\sigma^2}}} \quad (7)$$

Where $\lambda > 0$ is a shrunken parameter. Gamma and inverse gamma distribution are the prior for λ and σ^2 respectively. Based on all assumptions above the hierarchical prior of the Tobit single index model (TSIMo), will be considered as follows:

$$\begin{aligned} y_i^* &= \delta(X_i' \beta) + e_i \quad (i=1, 2, \dots, n) \\ y_i &= \begin{cases} y_i^* & \text{if } y_i^* > 0 \\ 0 & \text{if } y_i^* \leq 0 \end{cases} \\ y^* | \delta_n, \sigma^2, \beta &\sim N_n(\delta_n, \sigma^2 I_n), \\ \delta_n | \beta, \tau, x_i &\sim GP(0, C_n) \\ \beta &\sim \prod_{j=1}^p \frac{\lambda}{2\sigma^2} \exp\left\{-\frac{\lambda|\beta_j|}{\sigma^2}\right\}, \end{aligned} \quad (8)$$

$$\begin{aligned}\sigma^2 &\sim \text{Inverse Gamma}(a, b), \\ \lambda &\sim \text{Gamma}(a_\lambda, b_\lambda), \\ \tau &\sim \text{Inverse Gamma}(c, d)\end{aligned}$$

3. The Conditional Posterior Distributions

Based on the hierarchical model the full conditional posterior distribution of all parameters is given by $\pi(\beta, \delta, \lambda, \sigma^2, \tau | y^*) \propto$

$$(2\pi\sigma_e^2)^{-\frac{n}{2}} \exp\left\{-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i^* - \delta(X_i' \beta))^2\right\} \times \{\det(C_n)\}^{-\frac{1}{2}} \exp\left\{-\frac{1}{2} (\delta' C_n^{-1} \delta)\right\} \times \prod_{j=1}^p \frac{\lambda}{2\sigma^2} \exp\left\{-\frac{\lambda|\beta_j|}{\sigma^2}\right\} \left(\frac{1}{\tau}\right)^{c+1} \exp\left\{-\frac{d}{\tau}\right\} \times \lambda^{a_\lambda-1} \exp\{-b_\lambda \lambda\} \times (\sigma^2)^{-a-1} e^{-b|\sigma^2}|$$

(9)

From (9), to get the conditional posterior we will use the algorithm of Markov Chain Monte Carlo (MCMC) for all parameters, and can be used as follows:

- 1- Latent variables will be sampled using truncated normal distribution:

$$y^* | \delta_n, \sigma^2, \beta \sim y_i I\{y_i > 0\} + TN_{(-\infty, 0)}(\delta(X_i' \beta), \sigma^2) I\{y_i \leq 0\}$$

- 2- Sample the unknown function δ_n

$$\begin{aligned}\pi(\delta_n | \beta, y^*, \lambda, \sigma^2, \tau) &\propto \pi(y^* | \delta_n, \sigma^2, \beta) \times \pi(\delta_n | \beta, \tau, \mathbf{x}) \\ &\propto \{\det(D)\}^{-1/2} \exp\left\{-\frac{1}{2} (y^* - \delta_n)' D^{-1} (y^* - \delta_n)\right\} \times \{\det(C_n)\}^{-\frac{1}{2}} \exp\left\{-\frac{1}{2} (\delta_n' C_n^{-1} \delta_n)\right\} \\ &\quad \delta_n | \beta, y^*, \lambda, \sigma^2, \tau \sim N(A, B)\end{aligned}$$

Where $A = C_n(D + C_n)^{-1}(y^*)$ and $B = A = C_n(D + C_n)^{-1}D$,

Where $D = \sigma^2 I_n$.

- 3- Sample the coefficient vector β

$$\begin{aligned}\pi(\beta | \delta_n, y^*, \lambda, \sigma^2, \tau) &\propto \pi(y^* | \delta_n, \sigma^2, \beta) \times \pi(\delta_n | \beta, \tau, \mathbf{x}) \times \pi(\beta | \lambda, \tau) \\ &\propto \exp\left\{-\frac{1}{2} (y^*)' (D + C_n)^{-1} (y^*)\right\} \{\det(D + C_n)\}^{-\frac{1}{2}} \times \prod_{j=1}^p \exp\left\{-\frac{\lambda|\beta_j|}{\sigma^2}\right\}\end{aligned}$$

To sample β The Metropolis algorithm will be used.

- 4- Sample σ^2

$$\begin{aligned}\pi(\sigma^2 | \delta_n, y^*, \lambda, \beta, \tau) &\propto \pi(y^* | \delta_n, \sigma^2, \beta) \times \pi(\beta | \lambda, \tau) \times \pi(\sigma^2) \\ &\propto (\sigma^2)^{-\frac{n}{2}} \exp\left\{-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i^* - \delta(X_i' \beta))^2\right\} \times (\sigma^2)^{-p} \exp\left\{-\sum_{j=1}^p \frac{\lambda|\beta_j|}{\sigma^2}\right\} \times (\sigma^2)^{-a-1} e^{-b|\sigma^2|} \\ &\quad \pi(\sigma^2 | \delta_n, y^*, \lambda, \beta, \tau) \sim IG\left(\frac{n}{2} + p + a, \frac{1}{2} \sum_{i=1}^n (y_i^* - \delta(X_i' \beta))^2 + \sum_{j=1}^p \lambda|\beta_j| + b\right)\end{aligned}$$

- 5- Sample λ

$$\begin{aligned}\pi(\lambda | \sigma^2, \beta) &\propto \pi(\beta | \lambda, \tau) \times \pi(\lambda) \\ &\propto \prod_{j=1}^p \frac{\lambda}{2\sigma^2} \exp\left\{-\frac{\lambda|\beta_j|}{\sigma^2}\right\} \times \lambda^{a_\lambda-1} \exp\{-b_\lambda \lambda\} \\ &\propto \lambda^{p+a_\lambda-1} \exp\left\{-\lambda\left(\sum_{j=1}^p \frac{|\beta_j|}{\sigma^2} + b_\lambda\right)\right\} \\ &\quad \pi(\lambda | \sigma^2, \beta) \sim \text{Gamma}\left(p + a_\lambda, \sum_{j=1}^p \frac{|\beta_j|}{\sigma^2} + b_\lambda\right)\end{aligned}$$

- 6- Sample τ

$$\begin{aligned}\pi(\tau | \delta_n, y^*, \lambda, \beta, \sigma^2) &\propto \pi(y^* | \delta_n, \sigma^2, \beta) \times \pi(\delta_n | \beta, \tau, \mathbf{x}) \times \pi(\tau) \\ &\propto \exp\left\{-\frac{1}{2} (y^*)' (D + C_n)^{-1} (y^*)\right\} \{\det(D + C_n)\}^{-\frac{1}{2}} \times \left(\frac{1}{\tau}\right)^{c+1} \exp\left\{-\frac{d}{\tau}\right\}\end{aligned}$$

To sample τ Metropolis algorithm will be used.

4. Simulation Study

Here, we compare the proposed Bayesian Tobit single index model (BTSI) with two existing methods, Bayesian Tobit regression (BTR) and Tobit regression (TR), and show how well BTSI performs in estimation tasks. Two simulation examples have been conducted via MCMC algorithms that have been implemented in R, we run 10,000 iterations and the first 2,000 iterations have burned in for the sake of the stability of the generating process. These simulations have been replicated 150 times and covered different criteria of estimation to evaluate the performance of index vector; Standard deviation (SD), bias of estimation, mean square error (MSE), $MSE = \frac{1}{n} \sum_{i=1}^n \{y_i^{*true} - \hat{\delta}(\hat{\beta}^T x_i)\}^2$ and mean absolute error (MAE) $MAE = \frac{1}{n} \sum_{i=1}^n |y_i^{*true} - \hat{\delta}(\hat{\beta}^T x_i)|$. Also, the two simulation examples considered the homoscedastic of the error term in the studied models. We used different sample sizes (n=50,150,250,450) for both simulation examples. Also, we plotted the standard deviation values and the MSE and MAE criterion values. It is worth mentioning that these simulation examples have been studied by Alkenani and Yu (2013), and Turki (2021).

Example 1: As an example, we take into consideration the data that were produced by the Tobit single index model that is described below.

$$y_i = \begin{cases} y_i^* & \text{if } y_i^* > 0 \\ 0 & \text{if } y_i^* \leq 0 \end{cases}$$

Where the index $u = X'\beta$ is a linear combination of a 5-dimensional covariates $X = (X_1, X_2, \dots, X_5)'$ with known true index β -vector, $\beta = \frac{1}{\sqrt{5}}(1, 2, 0, 0, 0)^T$, X_i are identically and independently distributed (i.i.d) random variables from uniform (0,1) with $i=1, 2, \dots, 5$. Also, the error term $e \sim \exp(-\frac{1}{2})$ with X_i 's and e are mutually independent, and $\delta(u) = 5 \cos \cos(u) + \exp \exp(-u^2)$.

The following table contains the standard deviations (SD) values of the estimated index vector of parameters.

Table 1. summary of standard deviations example 1

Sample Size	Methods	SD1	SD2	SD3	SD4	SD5
50	BTSI	0.20129	0.21465	0.21976	0.23136	0.16102
	BTR	2.87533	2.62696	1.93637	1.94501	2.13816
	TR	2.46095	1.99408	1.26021	1.79600	1.61949
150	BTSI	0.22900	0.20536	0.24828	0.24812	0.28632
	BTR	1.25723	1.17180	1.29688	0.92770	1.39097
	TR	1.14325	1.06430	1.19594	0.83459	1.30132
250	BTSI	0.22231	0.20238	0.23499	0.21630	0.22296
	BTR	2.46687	3.49695	3.51514	3.50894	2.15289
	TR	2.27025	2.49954	3.64709	3.89295	2.05003
450	BTSI	0.23786	0.20184	0.21959	0.23570	0.23155
	BTR	1.14702	2.50436	1.42572	1.32941	1.07238
	TR	1.12280	2.38808	1.43253	1.32727	1.04702

Table-1 clearly shows that the proposed model (BTSI) has the smallest standard deviation comparing with the other methods and that indicates the (BTSI) method is outperformed and comparable to other methods. Also, we can see that the proposed Bayesian method worked better than other methods as the sample size grew. It means the proposed Bayesian single index model provides asymptotically consistent estimators. Consequently, the values in Table 1 indicate that the proposed Bayesian single index model gives accurate estimates of single index parameters homogeneity error term. The following Figure plot the values of standard deviations in Table-1.

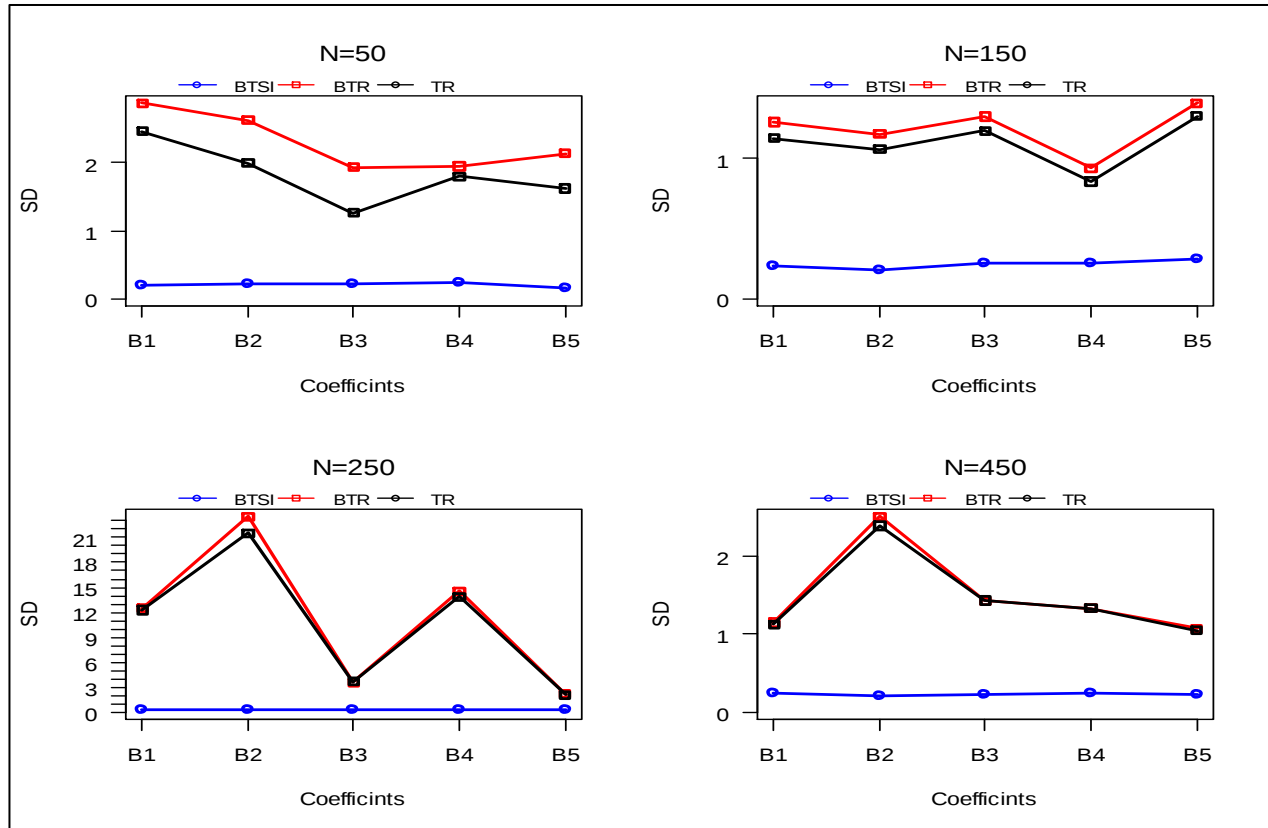


Figure 1. Plots of standard deviation in example 1.

From Figure 1, we can see that the proposed method (BTSI) outperformed the other methods, where the standard deviation values are much closed to zero.

Moreover, we write down the bias values of the estimators of index vectors with different sample sizes.

Table 2. Summary of bias values of example 1.

Sample Size	Methods	Bias 1	Bias 2	Bias 3	Bias 4	Bias 5
50	BTSI	0.02364	0.45626	0.31485	0.39654	0.21282
	BTR	0.96274	3.97818	0.87350	0.60319	0.59182
	TR	0.87449	3.10589	0.60893	0.55148	0.46997
150	BTSI	0.02685	0.53743	0.39867	0.30412	0.35109
	BTR	1.88083	3.38415	0.50572	0.73187	0.88489
	TR	1.69963	3.05456	0.51859	0.61311	0.58576
250	BTSI	0.09228	0.41861	0.46750	0.33381	0.35182
	BTR	2.88769	2.55796	1.06167	1.86229	0.86071
	TR	2.71844	1.76222	0.92355	1.84627	0.65078
450	BTSI	0.08850	0.49450	0.41400	0.42201	0.38138
	BTR	1.74754	3.99682	0.77085	1.35395	0.70216
	TR	1.69589	3.87674	0.68257	1.15256	0.60027

Table 2. Summarizes the bias of index vector of unknown estimator that estimated by the proposed method (BTSI) and the other methods. From table 2, we can conclude that the way that was suggested (BTSI) obtained the smallest bias and that indicates the (BTSI) gives more accurate estimates and outperformed the other method over different sample sizes. Lastly, table 3 shows the values of MSE and MAE criteria over different sample sizes.

Table 3. Values of MSE and MAE in example 1

Sample Size	Methods	MSE	MAE
50	BTSI	1.09939	0.90554
	BTR	6.78016	2.45047
	TR	4.50335	1.99510
150	BTSI	3.84834	1.08252
	BTR	11.02493	2.78152
	TR	9.63156	2.52907
250	BTSI	1.00134	0.88807
	BTR	11.04128	2.55601
	TR	10.26098	2.04082
450	BTSI	0.82249	0.79673
	BTR	8.20087	2.68879
	TR	7.73958	2.61042

It is clear from table 3; that when compared to other approaches, the suggested method (BTSI) yields the lowest values of mean squared error (MSE) and mean absolute error (MAE). This demonstrates that our proposed method provides the highest quality estimators. Additionally, we present the values of the MSE and MAE criterion in the graph.

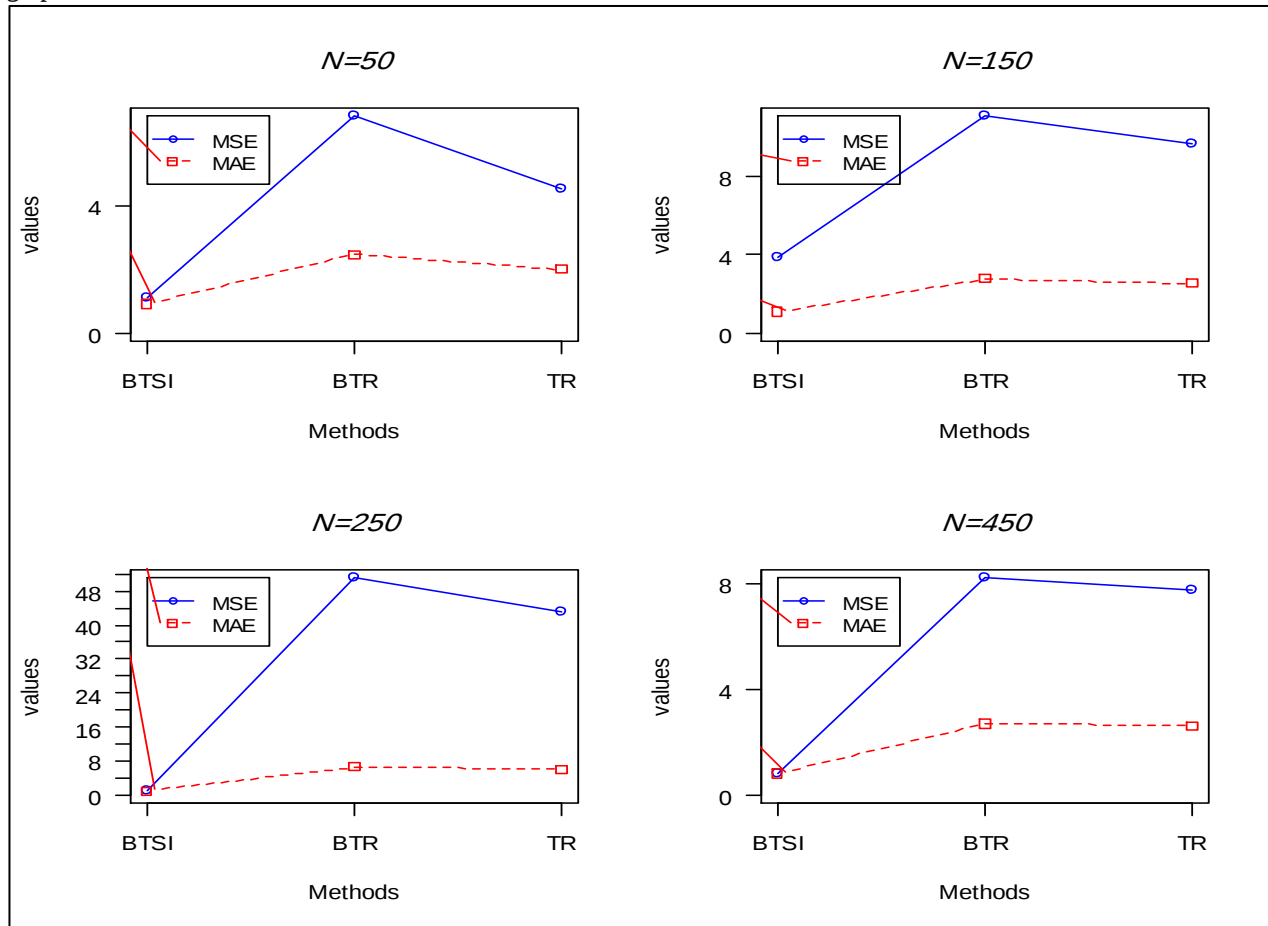


Figure 2. Plots of MSE versus MAE in example 1.

Again from Figure 2 Across a range of sample sizes, it is clear that the MSE and MAE values that our suggested technique produces are the lowest possible values.

Example 2: In this example, the data have been generated by the following Tobit single index model,

$$y_i^* = \delta(X_i' \beta) + e_i, \\ y_i = \begin{cases} y_i^* & \text{if } y_i^* > 0 \\ 0 & \text{if } y_i^* \leq 0 \end{cases}$$

Where, $\delta(X_i' \beta) = \exp(X_i' \beta)$, $X = (X_1, X_2, \dots, X_5)'$, and $\beta = \frac{(1, 1, 0, 0, 1)^T}{\sqrt{3}}$, X_i are (i.i.d) random variables from $N(0, 1)$, $i = 1, 2, \dots, 5$. Moreover, $e_i \sim N(0, 1)$ with X_i 's and e are mutually independent.

The following table summarizes the standard deviations (SD) values of the estimated index vector of parameters.

Table 4. summary of standard deviations example 2

Sample Size	Methods	SD1	SD2	SD3	SD4	SD5
50	BTSI	0.25629	0.19195	0.24573	0.27953	0.25565
	BTR	6.19686	6.29982	4.81735	7.71345	10.6521
	TR	4.31430	3.19714	3.22672	2.99126	3.08995
150	BTSI	0.21280	0.21265	0.24093	0.23695	0.22318
	BTR	2.88133	3.08537	3.09024	2.39133	2.16713
	TR	2.54108	2.71730	2.69340	2.06742	1.93983
250	BTSI	0.25281	0.23497	0.22814	0.19193	0.23654
	BTR	1.69850	1.60784	1.98318	2.15691	1.86284
	TR	1.61268	1.53265	1.87761	2.04345	1.75551
450	BTSI	0.25759	0.23510	0.21292	0.24999	0.24076
	BTR	1.19534	1.63787	1.51426	1.25880	1.22582
	TR	1.14947	1.57479	1.46025	1.21956	1.18102

From Table 4, we can see that the proposed model (BTSI) has the smallest standard deviation comparing with the frequentist methods and other Bayesian method and that indicates the (BTSI) method performed better than the other methods in terms of index estimation. The following Figure shows the values of standard deviations in Table-4.

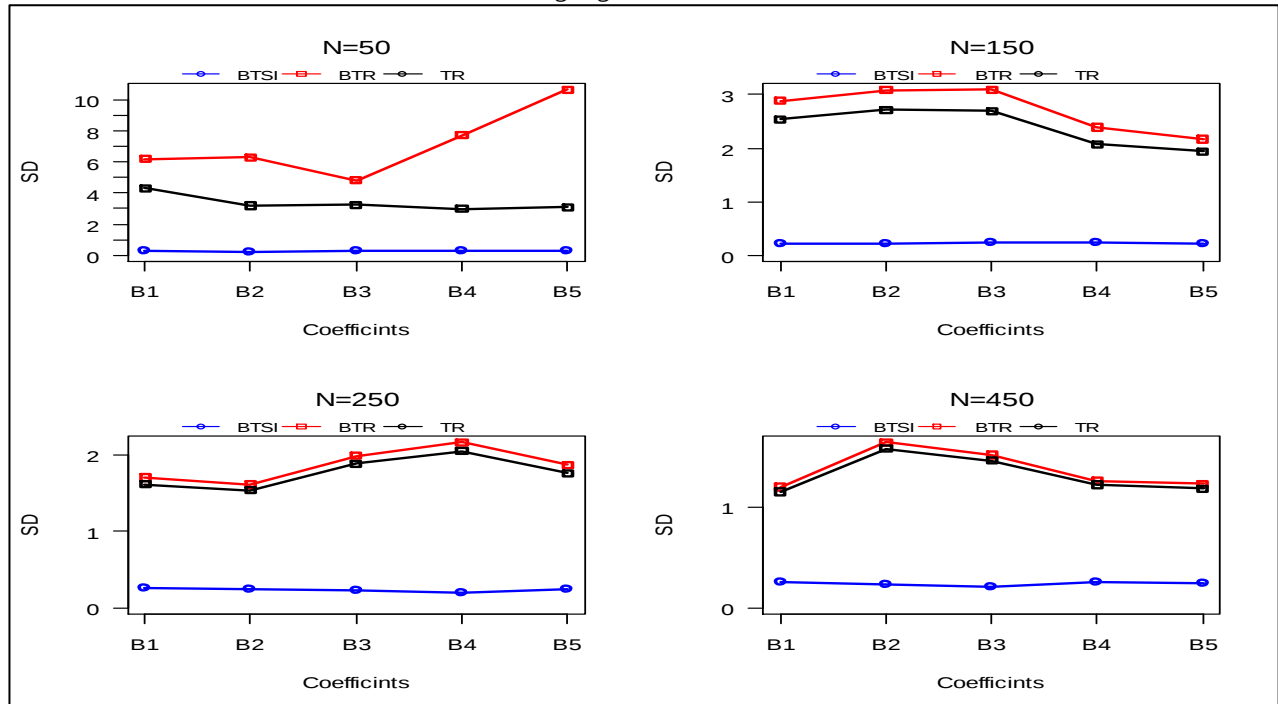


Figure 3. Plots of standard deviation in example 2.

Figure 3, demonstrates that the proposed method (BTSI) performed better than the other methods, where the standard deviation values are much closer to zero.

Also, the bias values of the estimators of the index vector are illustrated in the next table.

Table 5. Summary of bias values of example 2.

Sample Size	Methods	Bias 1	Bias 2	Bias 3	Bias 4	Bias 5
50	BTSI	0.19971	0.17348	0.34788	0.36378	0.22590
	BTR	0.42048	0.61129	0.43492	0.54865	1.27544
	TR	0.50073	0.30781	0.53893	0.74214	0.91673
150	BTSI	0.20006	0.19318	0.40987	0.35768	0.16490
	BTR	0.39089	0.70235	0.50567	0.80244	0.27899
	TR	0.35825	0.60884	0.62536	0.51571	0.28875
250	BTSI	0.18431	0.19701	0.37285	0.38598	0.16386
	BTR	0.46105	0.36824	0.52831	0.76741	0.39379
	TR	0.47122	0.49247	0.41175	0.55017	0.41554
450	BTSI	0.21559	0.15446	0.38274	0.41449	0.25719
	BTR	0.34967	0.27403	0.50917	0.82878	0.40834
	TR	0.36887	0.28522	0.49569	0.72278	0.41971

Table 5, demonstrated that the proposed method (BTSI) has the smallest bias between the other estimation methods and that indicates the (BTSI) gives more accurate estimates and performed better than other method over different sample sizes. Lastly, table 5 contains the values of MSE and MAE criteria.

Table 6. Values of MSE and MAE in example 2

Sample Size	Methods	MSE	MAE
50	BTSI	5.63118	1.57073
	BTR	7.87073	1.91729
	TR	6.97730	1.85462
150	BTSI	11.89335	2.12277
	BTR	13.99262	2.67651
	TR	13.94707	2.57485
250	BTSI	8.84507	2.04155
	BTR	10.53085	2.74698
	TR	10.62452	2.34729
450	BTSI	7.27120	1.30272
	BTR	8.88127	1.81648
	TR	8.91642	1.61804

From table 6, the smallest values of MSE and MAE are obtained by the proposed method (BTSI) comparing to other methods and that indicates the high performance of our proposed method. The following plot summarizes the values of the MSE and MAE criteria in simulation example 2.

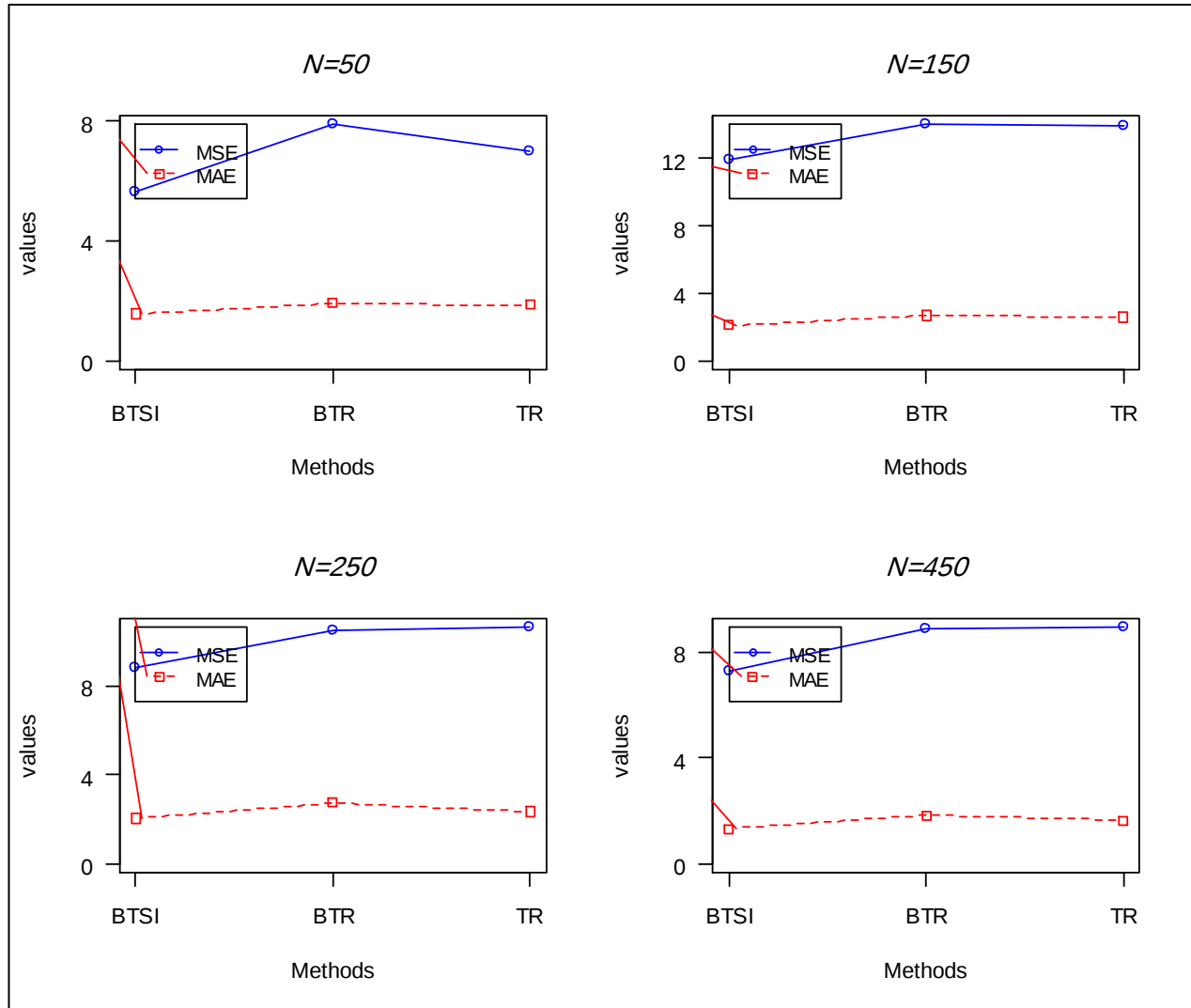


Figure 4. Plots of MSE versus MAE in example 2.

Figure 4 illustrates that the proposed method (BTSI) has the smallest values of MSE and MAE over all different sample sizes.

5. Real Data Analysis

In 1970, Prof. Ray Fair collected and analyzed the responses to a question about extramarital affairs by Tobit model for sample (601) observations of married men and women for the first time. In this section, we reanalyze this data by using the Bayesian Tobit single index. The study observations were based on 8 predictor variables, but we do not use three of them and used just the following variables:

y = no. of affairs in the past year {0, 1, 2, 3, 4–10 coded as 7}, {monthly, weekly, or daily,” coded as 12},

x_1 = sex = {0 if female 1 if male ,

x_2 = age,

x_3 = no. of years married,

x_4 = children = {0 if no 1 if yes ,

x_5 = self rating of marriage

(1 = Very unhappy, 2 = Pretty unhappy, 3 = Mildly happy, 4 = Pretty happy, 5 = Very happy).

(<http://pages.stern.nyu.edu/~wgreene/Text/tables/tablelist5.htm>).

The following table presents the index parameter estimates.

Table 7. Index parameters estimates of real data

Methods	β_1	β_2	β_3	β_4	β_5
BTSI	0.04505	0.63206	-0.66042	0.44596	-0.26072
BTR	-1.69653	3.15557	-2.03711	0.60406	-2.58131
TR	-1.66578	3.08729	-1.96868	0.59324	-2.52074

From Table 7, BTSI gives an estimate of $\beta_1 = 0.04505$, which indicates that the male effects the value of the response variable with %45 of times, but the effects were in negative with the other methods. $\beta_2 = 0.632$ in BTSI method, which indicates that the effects of age on response variable is positive, which is logically make no sense. $\beta_3 = -0.6602$ in BTSI method, have negative influence on the response variable, so the number of ears married have negative impact on the number of intercourses. Also, $\beta_4 = 0.445$ in BTSI method has positive impact of the response variable, that indicates in case of the family having children that will positively impacting on y. Finally, $\beta_5 = -0.26072$, the self-rating marriage have negative impact of the response variable. The following table includes the values of MSE and MAE .

Table 8. Real data MSE and MAE values

Methods	MSE	MAE
BTSI	14.41656	2.18976
BTR	19.81551	3.53365
TR	19.23183	3.47316

It can be seen that from table 8. The proposed method has the smallest values of MSE and MAE criteria and that indicates the outperformance of our proposed method comparing with other methods. Also, we plot the values of table 8 to support our understanding of the performance of the estimation method.

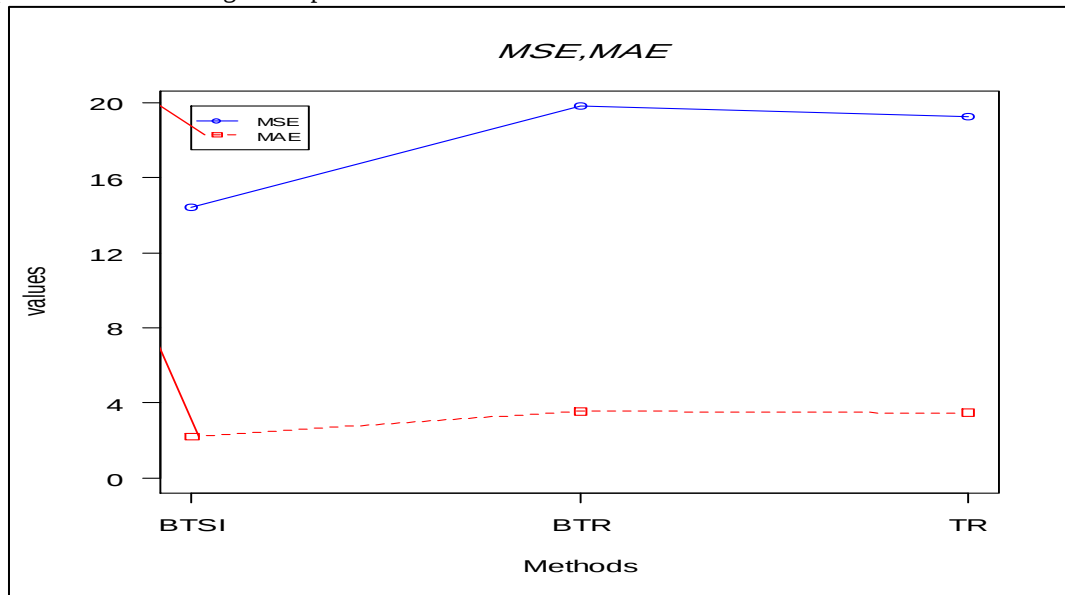


Figure 5. Plot of the MSE and MAE value

From figure 5 it can be seen that the proposed method has the lowest tract of the MSE and MAE values. Moreover, to check the convergence of the MCMC algorithm that has proposed to generate the interested estimates of the index parameters, we plot the trace plot for each generating process from the target posterior distribution.

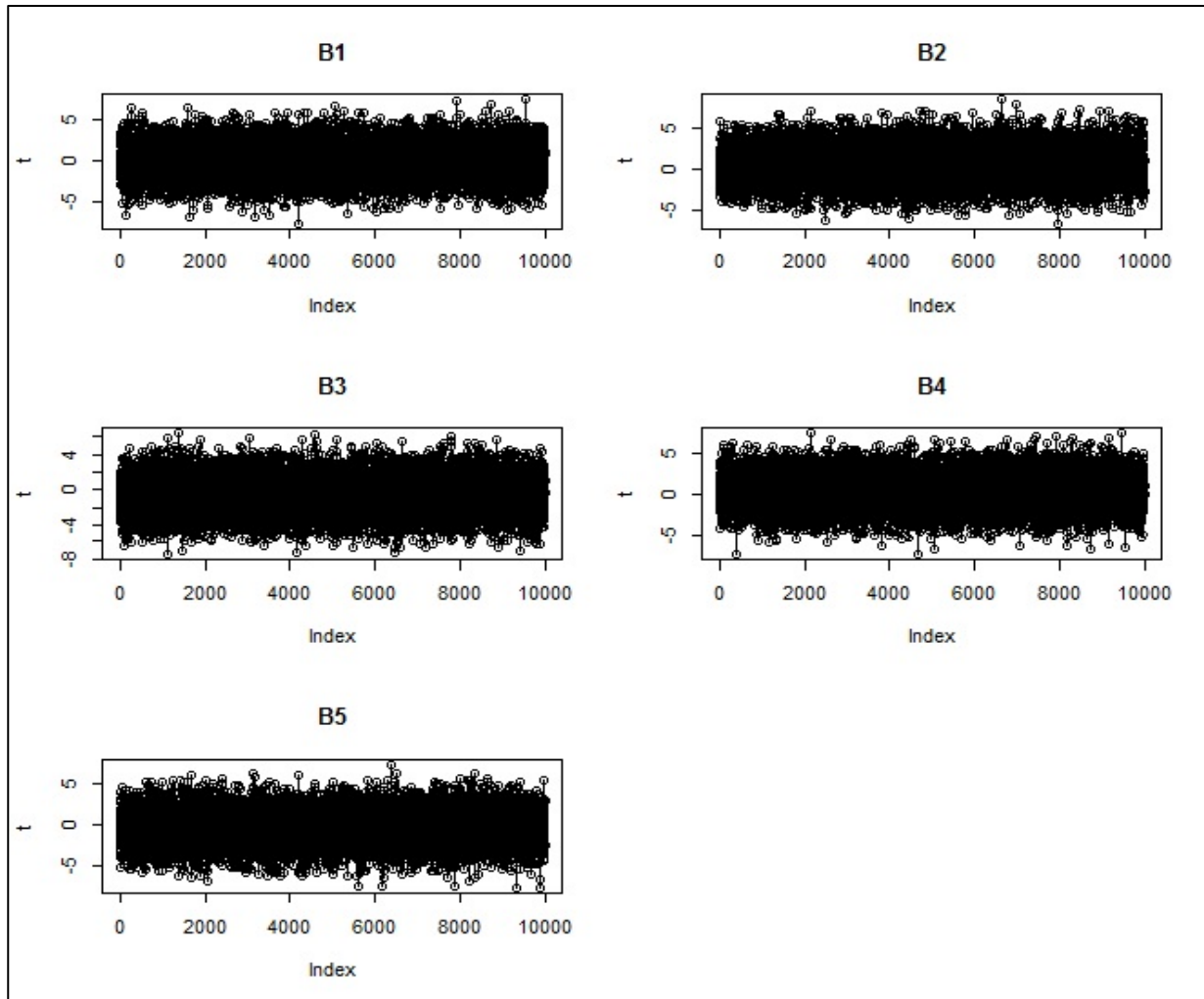


Figure 6. The trace plot of the index parameters $\beta_1 - \beta_5$.

From figure 6, we can see that the Gibbs sampling algorithm converge to the target distribution and have not suffer from slow mixing.

5. Conclusion

In this paper, we have constructed a new Bayesian hierarchical model to estimate the unknown nonparametric link function and the index parameters for the Tobit single index model. The double exponential distribution is set as a prior distribution to the parameter vector and the Gaussian process is considered as a prior distribution to the nonparametric function.

To examine the performance of our proposed method, we have considered some simulation examples and also used real data. Our proposed method Bayesian Tobit single index (BTISI) compares with two other methods Tobit regression (TR) and Bayesian Tobit regression (BTR). The results listed in the tables and figures showed that the proposed method had better performance than other methods.

References:

- 1) Alkenani, A., & Yu, K. (2013). Penalized single-index quantile regression. *International Journal of Statistics and Probability*, 2(3), 12
- 2) Amemiya T. (1984). Tobit models: A survey. *J Econ.*; 24:3–61.
- 3) Antoniadis, A., Grégoire, G., and McKeague, I. (2004), ‘Bayesian Estimation in Single-index Models’, *Statistica Sinica*, 14, 1147–1164.
- 4) Aradillas-Lopez A, Honoré B, Powell J. (2007). Pairwise difference estimation with nonparametric control variables. *Int Econ Rev*; 48(4):1119–1158.
- 5) Choi, T., Shi, J., Wang, B., (2011). A Gaussian process regression approach to a single-index model. *Journal of Nonparametric Statistics* 23 (1), 21–36.

- 6) Denison, D.G., Holmoes, C.C., Mallick, B.K., and Smith, A. (2002), Bayesian Methods for Nonlinear Classification and Regression, New York: Wiley.
- 7) Gramacy, R., Lian, H., (2012). Gaussian process single-index models as emulators for computer experiments. *Technometrics* to Appear.
- 8) Levy A.A., (2003). Generalized additive Tobit model. *Empir Econ.*;28(1):3–22.
- 9) Park T, Casella G (2008) The Bayesian lasso. *J Am Stat Assoc* 103:681–686.
- 10) Sami, Z., & Alshaybawee, T. (2021). Bayesian Variable Selection for Semiparametric Logistic Regression. *Al-Qadisiyah Journal of Pure Science*, 26.57-44 ,(5)
- 11) Tobin J. (1958). Estimation of relationships for limited dependent variables. *Econometrica.*; 26:24–36.
- 12) Turki, S.Z., Alshaybawee, H.T. (2021). Bayesian Estimation for Semiparametric Logistic Regression with an Application. M.Sc./Thesis. The University of Al-Qadisiyah. Iraq.
- 13) Wang, H.B. (2009), ‘Bayesian Estimation and Variable Selection for Single Index Models’, *Computational Statistics and Data Analysis*, 53, 2617–2627.
- Zhao, K., & Lian, H. (2015). Bayesian Tobit quantile regression with single-index models. *Journal of Statistical Computation and Simulation* 85(6), 1247-1263.