

Robust Bayesian Analysis for Shape Parameters with Progressive Type-II Censoring

N. Wheed^{1,*}, M. M. Nassar¹, M. Mahmoud¹ and M. Yusuf²

¹Department of Mathematics, Faculty of Science, Ain Shams University, Cairo, Egypt

²Department of Mathematics, Faculty of Science, Helwan University, Ain Helwan, Cairo, Egypt

Received: 17 Jan. 2021, Revised: 27 Feb. 2021, Accepted: 3 Mar. 2021

Published online: 1 May 2022

Abstract: In this paper, we discuss the robustness of bayesian estimation for the exponentiated inverted Weibull distribution and the inverted Kumaraswamy distribution under progressive type-II censored samples. To estimate the unknown shape parameters of the two distributions, we use Lindley approximation to obtain Bayes estimates using the loss functions; square error loss function, linex loss function and general entropy loss function under the non-informative and informative priors for the parameters. A simulation study is conducted to compare different estimates and the actual values of the shape parameters. Our aim is to study the robustness of estimation with regard to the three elements; the distribution model, the prior distribution and the loss function.

Keywords: Inverted Kumaraswamy Distribution (IKUM); Exponentiated Inverted Weibull Distribution (EIW); Bayesian and Non-Bayesian Estimations; Progressive Censoring; Loss Functions; Markov Chain Monte Carlo (MCMC) Methods; Robust Bayesian Analysis.

1 Introduction

Robust Bayesian analysis is concerned with the sensitivity of the results of a Bayesian analysis to the inputs for the analysis. The early 90s was the golden age of robust Bayesian analysis, where many statisticians were highly active in research in this area, and rapid progress was being achieved. Insua et al [1].

Robustness tests were originally introduced to avoid problems in inter laboratory studies and to identify the potentially responsible factors. This means that a robustness test was performed at a late stage in the method validation since inter laboratory studies are performed in the final stage. Thus the robustness test was considered a part of method validation related to the precision (reproducibility) determination of the method. Interest naturally expanded into study of robustness with respect to the likelihood and loss function as well, with the aim of having a general approach to sensitivity towards all the ingredients of the Bayesian paradigm (model/prior/loss). Heyden et al [2].

The inverted distributions have a wide range of applications; in problems related to econometrics, survey sampling, engineering sciences, medical research and life testing problems. Many researchers have been developing various extensions and modified forms of the inverted distributions and its applications; for example, Aljuaid [3] presented exponentiated inverted Weibull distribution Also AL-Fattah et al [4] discussed properties and estimation of parameters of the inverted kumaraswamy distribution as well as a generalization of the Inverse weibull distribution by adding a new shape parameter by exponentiation Khan et al [5].

There are numerous phenomena in life-testing, in reliability experiments, and survival analysis in which the experimenter usually does not have complete control of failure times from the experiments in hand. For example, patients may enroll in a clinic for treatment from a certain disease at random points of time and leave before completion of treatment or die from a cause different from the one under consideration. In an industrial experiment, units may break accidentally. In several situations, however, the removal of units prior to failure is prepared in advance that provides savings in terms of time and cost associated with testing.

* Corresponding author e-mail: nourawheed1994@gmail.com

Hence, the recorded survival times will involve randomly censored data. The scheme of progressive type-II right censoring arises naturally in life-testing experimentation, such that it allows an experimenter to remove survival items from testing at points other than the termination point.

Under this scheme, suppose that n items are placed and only $(m < n)$ are totally observed until failure. The censoring occurs progressively in m stages. These m stages offer failure times of the m totally observed items. At once following the time of the first observed failure (the first stage), a fixed number R_1 of the $n - 1$ surviving items are randomly dropped out (censored purposely) from the test. A fixed number R_2 of the $n - 2 - R_1$ surviving items are dropped out at the time of the second failure (the second stage), and so on. Finally, at the time of the (m^{th}) failure (the m^{th} stage), all the remaining $R_m = n - R_1 - R_2 - \dots - R_{m-1} - m$ surviving items are dropped out.

We denote this as progressive type-II right censoring scheme $R_1, R_2, \dots, R_m$. It seems that $n = m + \sum_i R_i$ the special case when $R_1 = R_2 = \dots = R_{m-1} = 0$ so that $R_m = n - m$ is the case of conventional type-II right censored sampling. Also, when $R_1 = R_2 = \dots = R_{m-1} = 0$ so that $n = m$ the progressively type-II right censoring scheme reduces to the case of no censoring Cramer and Iliopoulos [6].

Kumaraswamy[7] presented a distribution, which has many similarities to the beta distribution, and it is a new probability distribution for variables that are lower and upper bounded. In probability and statistics, the Kumaraswamy's double bounded distribution is a family of continuous probability distributions defined on the interval $[0, 1]$ differing in the values of their two non-negative shape parameters, α and β . Many times the data are modeled by finite range distributions. The inverted Kumaraswamy distribution can be derived from Kumaraswamy (Kum) distribution using the transformation.

In this paper, we discuss the effect of the three elements; the distribution, the prior and the loss function on the robustness of Bayes estimation on distribution, prior and loss function.

In case of the model element, we study the robustness on the models (IKUM and EIW) by choosing values of the shape parameters such that the two distributions are close in shape. Then, generate data from one distribution and use it to estimate the parameters of the other distribution and vice versa.

In Figure 1, we give the plots for the probability density function of two distributions IKUM and EIW.

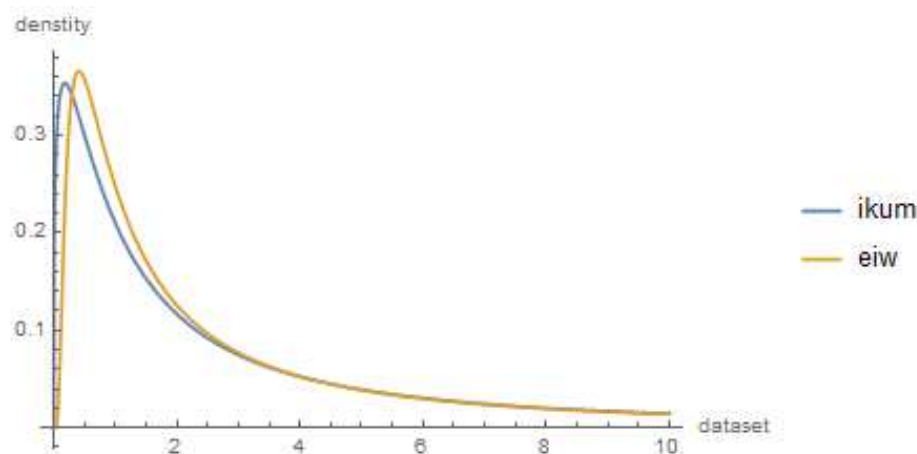


Fig. 1: Plot for the pdf of IKUM and EIW distribution $\alpha_J = 0.7, \beta_J = 1.3, J = 1, 2$

To study the effect of the prior distribution on the estimation of the parameters, we use non-informative and informative gamma prior and fixing the distribution and loss function.

To study the effect of the loss function on robustness of Bayes estimates, we use the different loss functions; symmetric square error loss function, asymmetric linex and general entropy loss functions with fixed distribution and prior.

2 Maximum Likelihood Estimation (MLE)

2.1 Maximum Likelihood (ML) Estimation For Inverted Kumaraswamy Distribution

The probability density function of the inverted kumaraswamy distribution, denoted by IKUM (α_1, β_1) , with shape parameters α_1 and $\beta_1 > 0$ is given by

$$f(x; \alpha_1, \beta_1) = \alpha_1 \beta_1 (1+x)^{-(\alpha_1+1)} (1 - (1+x)^{-\alpha_1})^{\beta_1-1}, x > 0, \alpha_1, \beta_1 > 0 \quad (1)$$

and the corresponding cumulative distribution function (cdf) is given by

$$F(x; \alpha_1, \beta_1) = (1 - (1+x)^{-\alpha_1})^{\beta_1}, x > 0, \alpha_1, \beta_1 > 0 \quad (2)$$

Let $X_1, \dots, X_m$ be progressively type II censored sample, with censoring scheme $(R_1, \dots, R_m)$.

The likelihood function is given by

$$L(\underline{x}; \alpha_1, \beta_1) = C \prod_{i=1}^m f(x_i) [1 - F(x_i)]^{R_i} \quad (3)$$

where $C = n(n - R_1 - 1) \cdots (n - R_1 - \cdots - R_{m-1} - m - 1)$, $f(X)$ and $F(x)$ are given by (1) and (2) respectively.

Then

$$\begin{aligned} L(\underline{x}; \alpha_1, \beta_1) &= C \prod_{i=1}^m [(\alpha_1 \beta_1 (1+x_i)^{-(\alpha_1+1)} (1 - (1+x_i)^{-\alpha_1})^{\beta_1-1} (1 - (1 - (1+x_i)^{-\alpha_1})^{\beta_1}))^{R_i}] \\ &= C (\alpha_1 \beta_1)^m e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1 - (1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1 - (1 - (1+x_i)^{-\alpha_1})^{\beta_1})}. \end{aligned} \quad (4)$$

The log-likelihood function can be written as

$$\begin{aligned} \ell &= \log L(\underline{x}; \alpha_1, \beta_1) \\ &= \log C + m \log \alpha_1 + m \log \beta_1 - (\alpha_1 + 1) \sum_{i=1}^m \ln(1+x_i) + (\beta_1 - 1) \sum_{i=1}^m \ln(1 - (1+x_i)^{-\alpha_1}) + \sum_{i=1}^m R_i \ln(1 - (1 - (1+x_i)^{-\alpha_1})^{\beta_1}). \end{aligned} \quad (5)$$

The MLE of the unknown parameters can be obtained by differentiating the log-likelihood function (5) with respect to the unknown parameters and equating to zero, thus giving

$$\begin{aligned} \frac{\partial \ell}{\partial \alpha_1} &= \frac{m}{\alpha_1} + \sum_{i=1}^m \ln(1+x_i) + (\beta_1 - 1) \sum_{i=1}^m \frac{(1+x_i)^{-\alpha_1} \ln(1+x_i)}{(1 - (1+x_i)^{-\alpha_1})} - \sum_{i=1}^m R_i \frac{\beta_1 (1 - (1+x)^{-\alpha_1})^{\beta_1-1} (1+x_i)^{-\alpha_1} \ln(1+x_i)}{1 - (1 - (1+x)^{-\alpha_1})^{\beta_1}} \\ &= 0. \end{aligned} \quad (6)$$

$$\frac{\partial \ell}{\partial \beta_1} = \frac{m}{\beta_1} + \sum_{i=1}^m \ln(1 - (1+x_i)^{-\alpha_1}) - \sum_{i=1}^m R_i \frac{(1 - (1+x)^{-\alpha_1})^{\beta_1} \ln(1 - (1+x_i)^{-\alpha_1})}{1 - (1 - (1+x)^{-\alpha_1})^{\beta_1}} = 0. \quad (7)$$

The solution of the non-linear equations (6) and (7) are $\hat{\alpha}_1$ and $\hat{\beta}_1$.

2.2 Maximum Likelihood (ML) Estimation For Exponentiated Inverted Weibull Distribution

The probability density function of the exponentiated inverted Weibull distribution, denoted by EIW (α_2, β_2) , with shape parameters α_2 and $\beta_2 > 0$ is given by

$$f(x; \alpha_2, \beta_2) = \alpha_2 \beta_2 x^{-(\alpha_2+1)} (e^{-x^{-\alpha_2}})^{\beta_2}, x > 0, \alpha_2, \beta_2 > 0. \quad (8)$$

and the corresponding cumulative distribution function (cdf) is given by

$$F(x; \alpha_2, \beta_2) = (e^{-x^{-\alpha_2}})^{\beta_2}, x > 0, \alpha_2, \beta_2 > 0. \quad (9)$$

Let $X_1, \dots, X_m$ be progressively type II censored sample, with censoring scheme $(R_1, \dots, R_m)$. The likelihood function is again given by (3) where $f(X)$ and $F(x)$ are given by (8) and (9) respectively.

Then

$$L(\underline{x}; \alpha_2, \beta_2) = C \prod_{i=1}^m [(\alpha_2 \beta_2 x_i^{-(\alpha_2+1)} (e^{-x_i^{-\alpha_2}})^{\beta_2}) (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}] \quad (10)$$

The log-likelihood function can be written as

$$\begin{aligned} \ell &= \log L(\underline{x}; \alpha_2, \beta_2) \\ &= \log C + m \log \alpha_2 + m \log \beta_2 - \beta_2 \sum_{i=1}^m x_i^{-\alpha_2} - (\alpha_2 + 1) \sum_{i=1}^m \log(x_i) + \sum_{i=1}^m R_i \log(1 - (e^{-x_i^{-\alpha_2}})^{\beta_2}) \end{aligned} \quad (11)$$

The MLE of the unknown parameters can be obtained by differentiating the log-likelihood function (11) with respect to the unknown parameters and equating to zero, thus giving

$$\frac{\partial \ell}{\partial \alpha_2} = \frac{m}{\alpha_2} - \sum_{i=1}^m \log(x_i) + \beta_2 \sum_{i=1}^m \log(x_i) x_i^{-\alpha_2} - \sum_{i=1}^m R_i \frac{e^{-\beta_2 x_i^{-\alpha_2}} x_i^{-\alpha_2} \beta_2 \log(x_i)}{(1 - e^{-\beta_2 x_i^{-\alpha_2}})} = 0 \quad (12)$$

$$\frac{\partial \ell}{\partial \beta_2} = \frac{m}{\beta_2} - \sum_{i=1}^m x_i^{-\alpha_2} - \sum_{i=1}^m R_i \frac{e^{-\beta_2 x_i^{-\alpha_2}} x_i^{-\alpha_2}}{(1 - e^{-\beta_2 x_i^{-\alpha_2}})} = 0 \quad (13)$$

The solution of the non-linear equations (12) and (13) are $\hat{\alpha}_2$ and $\hat{\beta}_2$.

3 Bayes Estimation In Case of The Non-informative Prior

In this section, we consider the Bayesian estimation for the two unknown parameters α_J and β_J , ($J = 1, 2$) of the distributions (IKUM) and (EIW), using progressive type-II censoring samples on the square error loss function, linear-exponential loss function, and general Entropy loss function.

Assume that α_J and β_J are independent random variables, having non-informative prior distributions defined by:

$$\pi(\alpha_J) \propto \frac{1}{\alpha_J}, 0 < \alpha_J < \infty \quad (14)$$

$$\pi(\beta_J) \propto \frac{1}{\beta_J}, 0 < \beta_J < \infty \quad (15)$$

Then the joint prior distribution for α_J and β_J is given by:

$$\pi(\alpha_J, \beta_J) \propto \frac{1}{\alpha_J \beta_J}, \alpha_J, \beta_J > 0 \quad (16)$$

Also, the joint posterior distribution of α_J and β_J is given by:

$$\pi(\alpha_J, \beta_J | \underline{x}) = \frac{\pi(\alpha_J, \beta_J) L(\underline{x} | \alpha_J, \beta_J)}{\int_0^\infty \int_0^\infty \pi(\alpha_J, \beta_J) L(\underline{x} | \alpha_J, \beta_J) d\alpha_J d\beta_J} \quad (17)$$

By using Equation (16), (17) and Equation (4) for (IKUM) distribution we obtain the joint posterior distribution defined by:

$$\pi(\alpha_1, \beta_1 | \underline{x}) = \frac{C(\alpha_1 \beta_1)^{m-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})}}{\int_0^\infty \int_0^\infty C(\alpha_1 \beta_1)^{m-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1} \quad (18)$$

Also, using Equation (16), (17) and Equation (10) for (EIW) distribution we obtain the joint posterior distribution defined by:

$$\pi(\alpha_2, \beta_2 | \underline{x}) = \frac{C(\alpha_2 \beta_2)^{m-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}})^{\prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]} }{\int_0^\infty \int_0^\infty C(\alpha_2 \beta_2)^{m-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}})^{\prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]} d\alpha_2 d\beta_2} \quad (19)$$

3.1 Bayesian Estimators Under Square Error Loss Function

The Bayesian estimator of any function for α_J and β_J say $U(\alpha_J, \beta_J)$, under the squared error loss function, is given by:

$$E(u(\alpha_J, \beta_J)|\underline{x}) = \frac{\int_0^\infty \int_0^\infty u(\alpha_J, \beta_J) \pi(\alpha_J, \beta_J) L(\underline{x}|\alpha_J, \beta_J) d\alpha_J d\beta_J}{\int_0^\infty \int_0^\infty \pi(\alpha_J, \beta_J) L(\underline{x}|\alpha_J, \beta_J) d\alpha_J d\beta_J} \quad (20)$$

The Bayesian estimators for two shape parameters of the (IKUM) distribution are given as follows

$$\begin{aligned} \hat{\alpha}_{1sq} &= E(\alpha_1) \\ &= \frac{\int_0^\infty \int_0^\infty C(\alpha_1)^m (\beta_1)^{m-1} e^{-(\alpha_1+1)\sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1)\sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1+x_i)^{-\alpha_1})^{\beta_1}} d\alpha_1 d\beta_1}{\int_0^\infty \int_0^\infty C(\alpha_1 \beta_1)^{m-1} e^{-(\alpha_1+1)\sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1)\sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1+x_i)^{-\alpha_1})^{\beta_1}} d\alpha_1 d\beta_1} \end{aligned} \quad (21)$$

$$\begin{aligned} \hat{\beta}_{1sq} &= E(\beta_1) \\ &= \frac{\int_0^\infty \int_0^\infty C(\alpha_1)^{m-1} (\beta_1)^m e^{-(\alpha_1+1)\sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1)\sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1+x_i)^{-\alpha_1})^{\beta_1}} d\alpha_1 d\beta_1}{\int_0^\infty \int_0^\infty C(\alpha_1 \beta_1)^{m-1} e^{-(\alpha_1+1)\sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1)\sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1+x_i)^{-\alpha_1})^{\beta_1}} d\alpha_1 d\beta_1} \end{aligned} \quad (22)$$

Also, the Bayesian estimators for two shape parameters of the (EIW) distribution are given as follows

$$\begin{aligned} \hat{\alpha}_{2sq} &= E(\alpha_2) \\ &= \frac{\int_0^\infty \int_0^\infty C(\alpha_2)^m (\beta_2)^{m-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}} \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]) d\alpha_2 d\beta_2}{\int_0^\infty \int_0^\infty C(\alpha_2 \beta_2)^{m-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}} \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]) d\alpha_2 d\beta_2} \end{aligned} \quad (23)$$

$$\begin{aligned} \hat{\beta}_{2sq} &= E(\beta_2) \\ &= \frac{\int_0^\infty \int_0^\infty C(\alpha_2)^{m-1} (\beta_2)^m (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}} \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]) d\alpha_2 d\beta_2}{\int_0^\infty \int_0^\infty C(\alpha_2 \beta_2)^{m-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}} \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}]) d\alpha_2 d\beta_2} \end{aligned} \quad (24)$$

Provided that $E(\alpha_J)$ and $E(\beta_J)$ in both cases of the two distributions exist and are finite. These integrations cannot be solved analytically, so we use Lindley approximation.

As given by Lindley [8], Lindley approximation for any function u of parameter θ , $\theta = (\alpha_J, \beta_J)$, is defined by

$$E(u(\theta)|\underline{x}) \approx (u(\alpha_J, \beta_J) + \frac{1}{2} [\sum_r \sum_t (u_{rt} + 2u_{rt} Q_t) \sigma_{rt} + \sum_r \sum_t \sum_k \sum_w L_{rtkw} u_{rw} \sigma_{rt} \sigma_{kw}]) \quad (25)$$

where $r, t, k, w = 1, 2$

Then

$$\begin{aligned} E(u(\theta)|\underline{x}) &= u + Q_1(u_1 \sigma_{11} + u_2 \sigma_{21}) + Q_2(u_1 \sigma_{12} + u_2 \sigma_{22}) + \frac{1}{2} [u_{11} \sigma_{11} + 2u_{21} \sigma_{12} + u_{22} \sigma_{22}] \\ &\quad + \frac{1}{2} [L_{111}(\sigma_{11}^2 u_1 + u_2 \sigma_{11} \sigma_{12}) + L_{112}(u_2(\sigma_{11} \sigma_{22} + 2\sigma_{12}^2) + 3u_1 \sigma_{21} \sigma_{11}) \\ &\quad + L_{122}(u_1(\sigma_{11} \sigma_{22} + 2\sigma_{21}^2) + 3u_2 \sigma_{12} \sigma_{22}) + L_{222}(u_1 \sigma_{22} \sigma_{21} + u_2 \sigma_{22}^2)] \end{aligned} \quad (26)$$

where

$$\begin{aligned} Q(\alpha_J, \beta_J) &= \ln(\pi(\alpha_J, \beta_J)), \quad Q_1 = \frac{\partial Q(\alpha_J, \beta_J)}{\partial \alpha_J}, \quad Q_2 = \frac{\partial Q(\alpha_J, \beta_J)}{\partial \beta_J}, \\ u_1 &= \frac{\partial u(\alpha_J, \beta_J)}{\partial \alpha_J}, \quad u_2 = \frac{\partial u(\alpha_J, \beta_J)}{\partial \beta_J}, \quad u_{11} = \frac{\partial^2 u(\alpha_J, \beta_J)}{\partial \alpha_J^2}, \quad u_{22} = \frac{\partial^2 u(\alpha_J, \beta_J)}{\partial \beta_J^2}, \\ u_{12} &= \frac{\partial^2 u(\alpha_J, \beta_J)}{\partial \alpha_J \partial \beta_J}, \quad L_{11} = \frac{\partial^2 \ell}{\partial \alpha_J^2}, \quad L_{12} = \frac{\partial^2 \ell}{\partial \alpha_J \partial \beta_J}, \quad L_{22} = \frac{\partial^2 \ell}{\partial \beta_J^2}, \\ L_{111} &= \frac{\partial^3 \ell}{\partial \alpha_J^3}, \quad L_{222} = \frac{\partial^3 \ell}{\partial \beta_J^3}, \quad L_{112} = \frac{\partial^3 \ell}{\partial \alpha_J^2 \partial \beta_J} \quad \text{and} \quad L_{122} = \frac{\partial^3 \ell}{\partial \alpha_J \partial \beta_J^2}. \end{aligned} \quad (27)$$

and $\sigma_{rt} = (r, t)$ th element in the inverse of the matrix $[-L_{rt}]$; $r, t = 1, 2$ for the two distributions is

$$\sigma_{rt} = \begin{pmatrix} -L_{11} & -L_{12} \\ -L_{21} & -L_{22} \end{pmatrix}^{-1} \quad (28)$$

Using Mathematica program we can calculate the inverse matrix, to find the values of σ_{rt} .

Substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \alpha_J$ in Equation (26) gives the Bayesian estimator for the shape parameter α_J for the two distributions to be

$$\hat{\alpha}_{J(sq)} = E(\alpha_J) = \alpha_J + \left[\left(\frac{-1}{\alpha_J} \right) \sigma_{11} + \left(\frac{-1}{\beta_J} \right) \sigma_{12} \right] + \frac{1}{2} [L_{111} \sigma_{11}^2 + 3L_{112} \sigma_{21} \sigma_{11} + L_{122} (\sigma_{11} \sigma_{22} + 2\sigma_{21}^2) + L_{222} \sigma_{22} \sigma_{21}] \quad (29)$$

Also substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \beta_J$ in Equation (26) gives the Bayesian estimator for the shape parameter β_J for the two distributions to be

$$\hat{\beta}_{J(sq)} = E(\beta_J) = \beta_J + \left[\left(\frac{-1}{\alpha_J} \right) \sigma_{21} + \left(\frac{-1}{\beta_J} \right) \sigma_{22} \right] + \frac{1}{2} [L_{111} \sigma_{11} \sigma_{12} + 3L_{122} \sigma_{12} \sigma_{22} + L_{112} (\sigma_{11} \sigma_{22} + 2\sigma_{12}^2) + L_{222} \sigma_{22}^2] \quad (30)$$

3.2 Bayesian Estimators Under Linear-Exponential Loss Function (LINEX)

The Bayes estimator of α_J denoted by $\hat{\alpha}_{J(L,c_j)}$ using to the Linex loss function is defined as:

$$\hat{\alpha}_{J(L,c_j)} = \frac{-1}{c_j} \log[E(e^{-\alpha_J c_j})]$$

and substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = e^{-\alpha_J c_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter α_J for the two distributions given by

$$\begin{aligned} \hat{\alpha}_{J(L,c_j)} = & \left(\frac{-1}{c_j} * \log[e^{-c_j \alpha_J} + \left[\left(\frac{-1}{\alpha_J} \right) \sigma_{11} + \left(\frac{-1}{\beta_J} \right) \sigma_{12} \right] c_j * e^{-c_j \alpha_J} - \frac{1}{2} * c_j * e^{-c_j \alpha_J} [L_{111} \sigma_{11}^2 \right. \right. \\ & \left. \left. + 3L_{112} \sigma_{21} \sigma_{11} + L_{122} (\sigma_{11} \sigma_{22} + 2\sigma_{21}^2) + L_{222} \sigma_{22} \sigma_{21}] + \frac{1}{2} c_j^2 e^{-\alpha_J c_j} \sigma_{11}] \right) \end{aligned} \quad (31)$$

However, substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = e^{-\beta_J c_j}$ in Equation (26) we obtain the Bayesian estimation for the shape parameter β_J for the two distributions given by

$$\begin{aligned} \hat{\beta}_{J(L,c_j)} = & \left(\frac{-1}{c_j} * \log[e^{-c_j \beta_J} + \left[\left(\frac{-1}{\alpha_J} \right) \sigma_{21} + \left(\frac{-1}{\beta_J} \right) \sigma_{22} \right] c_j * e^{-c_j \beta_J} - \frac{1}{2} * c_j * e^{-c_j \beta_J} [L_{111} \sigma_{11} \sigma_{12} \right. \right. \\ & \left. \left. + 3L_{122} \sigma_{12} \sigma_{22} + L_{112} (\sigma_{11} \sigma_{22} + 2\sigma_{12}^2) + L_{222} \sigma_{22}^2] + \frac{1}{2} c_j^2 e^{-\beta_J c_j} \sigma_{22}] \right) \end{aligned} \quad (32)$$

where c_j is constant and $j = 1, 2$

3.3 Bayesian Estimators Under General Entropy Loss Function

The Bayes estimator of α_J denoted by $\hat{\alpha}_{J(G,q_j)}$ using to the general entropy loss function is defined as:

$$\hat{\alpha}_{J(G,q_j)} = [E(\alpha_J^{-q_j})]^{\frac{-1}{q_j}}$$

and substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \alpha_J^{-q_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter α_J for the two distributions given by

$$\hat{\alpha}_{J(G,q_j)} = (\alpha_J^{-q_j} + (\frac{1}{\alpha_J})q_j\alpha_J^{-q_j-1}\sigma_{11} + (\frac{1}{\beta_J})q_j\alpha_J^{-q_j-1}\sigma_{12} - \frac{1}{2} * q_j * \alpha_J^{-q_j-1}[L_{111}\sigma_{11}^2 + 3L_{112}\sigma_{21}\sigma_{11} + L_{122}(\sigma_{11}\sigma_{22} + 2\sigma_{21}^2) + L_{222}\sigma_{22}\sigma_{21}] + \frac{1}{2}q_j(q_j+1)\alpha_J^{-q_j-2}\sigma_{11})^{\frac{-1}{q_j}} \quad (33)$$

Also substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \beta_J^{-q_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter β_J for the two distributions given by

$$\hat{\beta}_{J(G,q_j)} = (\beta_J^{-q_j} + (\frac{1}{\alpha_J})q_j\alpha_J^{-q_j-1}\sigma_{21} + (\frac{1}{\beta_J})q_j\alpha_J^{-q_j-1}\sigma_{22} - \frac{1}{2} * q_j * \alpha_J^{-q_j-1}[L_{111}\sigma_{11}\sigma_{12} + 3L_{122}\sigma_{12}\sigma_{22} + L_{112}(\sigma_{11}\sigma_{22} + 2\sigma_{12}^2) + L_{222}\sigma_{22}^2] + \frac{1}{2}q_j(q_j+1)\beta_J^{-q_j-2}\sigma_{22})^{\frac{-1}{q_j}} \quad (34)$$

where q_j is constant and $j = 1, 2$

4 Bayes Estimation In Case of The Informative Prior

In this section, we calculate the Bayesian estimators for the two unknown parameters α_J and β_J , ($J = 1, 2$) of the distributions (IKUM) and (EIW), using progressive type-II censoring samples, based on the square error loss function, linear-exponential loss function, and general Entropy loss function.

Assume that α_J and β_J are independent random variables, and consider an informative prior distribution for α_J and β_J following gamma distributions defined by:

$$\pi(\alpha_J) = \frac{a^b e^{-\alpha_J a} \alpha_J^{b-1}}{\Gamma(b)}, 0 < \alpha_J < \infty \quad (35)$$

$$\pi(\beta_J) = \frac{g^d e^{-\beta_J g} \beta_J^{d-1}}{\Gamma(d)}, 0 < \beta_J < \infty \quad (36)$$

Then the joint prior distribution for α_J and β_J is defined by:

$$\pi(\alpha_J, \beta_J) \propto \frac{a^b g^d e^{-(\alpha_J a + \beta_J g)} \alpha_J^{b-1} \beta_J^{d-1}}{\Gamma(b)\Gamma(d)}, \alpha_J, \beta_J > 0 \quad (37)$$

Figure 2 show plot of the joint prior for α_J and β_J following gamma distributions with values ($a = 5$, $b = 5$, $g = 2.25$, $d = 1.5$).

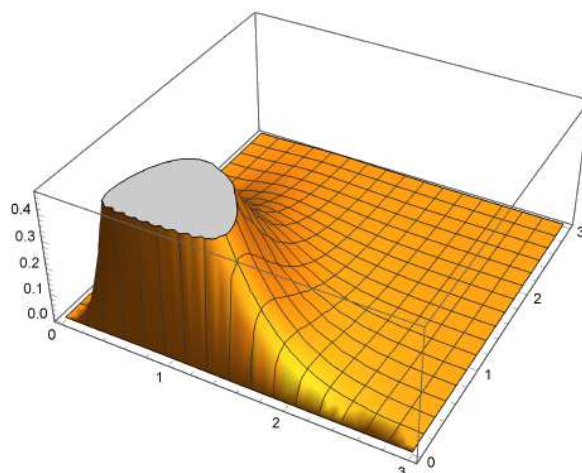


Fig. 2: Plot of the joint prior for α_J and β_J following gamma distributions with values ($a = 5, b = 5, g = 2.25, d = 1.5$)

Using Equation (37), (17) and Equation (4) for (IKUM) distribution, we get the joint posterior distribution defined by:

$$\pi(\alpha_1, \beta_1 | \underline{x}) = \frac{C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b-1} \beta_1^{m+d-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})}}{\int_0^\infty \int_0^\infty C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b-1} \beta_1^{m+d-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1} \quad (38)$$

Again, using Equation (37), (17) and Equation (10) for (EIW) distribution, we get the joint posterior distribution defined by:

$$\pi(\alpha_2, \beta_2 | \underline{x}) = \frac{C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b-1} \beta_2^{m+d-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}})^{\prod_{i=1}^m [X_i^{-(\alpha_2+1)}] (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}}}{\int_0^\infty \int_0^\infty C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b-1} \beta_2^{m+d-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}})^{\prod_{i=1}^m [X_i^{-(\alpha_2+1)}] (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}} d\alpha_2 d\beta_2} \quad (39)$$

4.1 Bayesian Estimators Under Square Error Loss Function

The Bayesian estimator of any function for α_J and β_J say $U(\alpha_J, \beta_J)$, under the squared error loss function, is given by:

$$E(u(\alpha_J, \beta_J) | \underline{x}) = \frac{\int_0^\infty \int_0^\infty u(\alpha_J, \beta_J) \pi(\alpha_J, \beta_J) L(\underline{x} | \alpha_J, \beta_J) d\alpha_J d\beta_J}{\int_0^\infty \int_0^\infty \pi(\alpha_J, \beta_J) L(\underline{x} | \alpha_J, \beta_J) d\alpha_J d\beta_J} \quad (40)$$

The Bayesian estimators for the two shape parameters of the (IKUM) distribution are given as follows

$$\hat{\alpha}_{1(sq)} = E(\alpha_1) = \frac{\int_0^\infty \int_0^\infty C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b} \beta_1^{m+d-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1}{\int_0^\infty \int_0^\infty C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b-1} \beta_1^{m+d-1} (\alpha_1 \beta_1)^{m-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1} \quad (41)$$

$$\hat{\beta}_{1(sq)} = E(\beta_1) = \frac{\int_0^\infty \int_0^\infty C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b-1} \beta_1^{m+d} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1}{\int_0^\infty \int_0^\infty C e^{-(\alpha_1 a + \beta_1 g)} \alpha_1^{m+b-1} \beta_1^{m+d-1} e^{-(\alpha_1+1) \sum_{i=1}^m \ln(1+x_i)} e^{(\beta_1-1) \sum_{i=1}^m \ln(1-(1+x_i)^{-\alpha_1})} e^{\sum_{i=1}^m R_i \ln(1-(1-(1+x_i)^{-\alpha_1})^{\beta_1})} d\alpha_1 d\beta_1} \quad (42)$$

Also, the Bayesian estimators for the two shape parameters of the (EIW) distribution are given as follows

$$\begin{aligned}\hat{\alpha}_{2(sq)} &= E(\alpha_2) \\ &= \frac{\int_0^\infty \int_0^\infty C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b} \beta_2^{m+d-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}}) \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}] d\alpha_2 d\beta_2}{\int_0^\infty \int_0^\infty C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b-1} \beta_2^{m+d-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}}) \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}] d\alpha_2 d\beta_2}\end{aligned}\quad (43)$$

$$\begin{aligned}\hat{\beta}_{2(sq)} &= E(\beta_2) \\ &= \frac{\int_0^\infty \int_0^\infty C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b-1} \beta_2^{m+d} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}}) \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}] d\alpha_2 d\beta_2}{\int_0^\infty \int_0^\infty C e^{-(\alpha_2 a + \beta_2 g)} \alpha_2^{m+b-1} \beta_2^{m+d-1} (e^{-\beta_2 \sum_{i=1}^m x_i^{-\alpha_2}}) \prod_{i=1}^m [X_i^{-(\alpha_2+1)} (1 - (e^{-x_i^{-\alpha_2}})^{\beta_2})^{R_i}] d\alpha_2 d\beta_2}\end{aligned}\quad (44)$$

Provided that $E(\alpha_J)$, $E(\beta_J)$ for the two distributions exist and are finite. These integrations can not be solved analytically, so we use Lindley Bayes approximation.

Substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \alpha_J$ in Equation (26) we obtain the Bayesian estimator for the shape parameter α_J for the two distributions given by

$$\begin{aligned}\hat{\alpha}_{J(sq)} &= E(\alpha_J) \\ &= \alpha_J + [(-a + \frac{-1+b}{\alpha_J})\sigma_{11} + (-g + \frac{-1+d}{\beta_J})\sigma_{12}] + \frac{1}{2}[L_{111}\sigma_{11}^2 + 3L_{112}\sigma_{21}\sigma_{11} + L_{122}(\sigma_{11}\sigma_{22} + 2\sigma_{21}^2) + L_{222}\sigma_{22}\sigma_{21}]\end{aligned}\quad (45)$$

Also, substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \beta_J$ in Equation (26) we obtain the Bayesian estimator for the shape parameter β_J for the two distributions given by

$$\begin{aligned}\hat{\beta}_{J(sq)} &= E(\beta_J) \\ &= \beta_J + [(-a + \frac{-1+b}{\alpha_J})\sigma_{21} + (-g + \frac{-1+d}{\beta_J})\sigma_{22}] + \frac{1}{2}[L_{111}\sigma_{11}\sigma_{12} + 3L_{122}\sigma_{12}\sigma_{22} + L_{112}(\sigma_{11}\sigma_{22} + 2\sigma_{12}^2) + L_{222}\sigma_{22}^2]\end{aligned}\quad (46)$$

4.2 Bayesian Estimators Under Linear-Exponential Loss Function (LINEX)

The Bayes estimator of α_J denoted by $\hat{\alpha}_{J(L,c_j)}$ using to the Linex loss function is defined as:

$$\hat{\alpha}_{J(L,c_j)} = \frac{-1}{c_j} \log[E(e^{-\alpha_J c_j})]$$

Substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = e^{-\alpha_J c_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter α_J for the two distributions given by

$$\begin{aligned}\hat{\alpha}_{J(L,c_j)} &= \left(\frac{-1}{c_j} * \log[e^{-c_j \alpha_J} - [(-a + \frac{-1+b}{\alpha_J})\sigma_{11} + (-g + \frac{-1+d}{\beta_J})\sigma_{12}]c_j * e^{-c_j \alpha_J} - \frac{1}{2} * c_j * e^{-c_j \alpha_J} [L_{111}\sigma_{11}^2 \right. \\ &\quad \left. + 3L_{112}\sigma_{21}\sigma_{11} + L_{122}(\sigma_{11}\sigma_{22} + 2\sigma_{21}^2) + L_{222}\sigma_{22}\sigma_{21}] + \frac{1}{2}c_j^2 e^{-\alpha_J c_j} \sigma_{11}]\right)\end{aligned}\quad (47)$$

Also substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = e^{-\beta_J c_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter β_J for the two distributions given by

$$\begin{aligned}\hat{\beta}_{J(L,c_j)} &= \left(\frac{-1}{c_j} * \log[e^{-c_j \beta_J} - [(-a + \frac{-1+b}{\alpha_J})\sigma_{21} + (-g + \frac{-1+d}{\beta_J})\sigma_{22}]c_j * e^{-c_j \beta_J} - \frac{1}{2} * c_j * e^{-c_j \beta_J} [L_{111}\sigma_{11}\sigma_{12} \right. \\ &\quad \left. + 3L_{122}\sigma_{12}\sigma_{22} + L_{112}(\sigma_{11}\sigma_{22} + 2\sigma_{12}^2) + L_{222}\sigma_{22}^2] + \frac{1}{2}c_j^2 e^{-\beta_J c_j} \sigma_{22}]\right)\end{aligned}\quad (48)$$

where c_j is constant and $j = 1, 2$.

4.3 Bayesian Estimators Under General Entropy Loss Function

The Bayes estimator of α_J denoted by $\hat{\alpha}_{J(G,q_j)}$ relative to the general entropy loss function is defined as:

$$\hat{\alpha}_{J(G,q_j)} = [E(\alpha_J^{-q_j})]^{-\frac{1}{q_j}}$$

Substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \alpha_J^{-q_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter α_J for the two distributions given by

$$\begin{aligned} \hat{\alpha}_{J(G,q_j)} = & (\alpha_J^{-q_j} - (-a + \frac{-1+b}{\alpha_J})q_j\alpha_J^{-q_j-1}\sigma_{11} - (-g + \frac{-1+d}{\beta_J})q_j\alpha_J^{-q_j-1}\sigma_{12} - \frac{1}{2} * q_j * \alpha_J^{-q_j-1} [L_{111}\sigma_{11}^2 \\ & + 3L_{112}\sigma_{21}\sigma_{11} + L_{122}(\sigma_{11}\sigma_{22} + 2\sigma_{21}^2) + L_{222}\sigma_{22}\sigma_{21}] + \frac{1}{2}q_j(q_j+1)\alpha_J^{-q_j-2}\sigma_{11})^{-\frac{1}{q_j}} \end{aligned} \quad (49)$$

Also substituting $Q_1 = \frac{-1}{\alpha_J}$, $Q_2 = \frac{-1}{\beta_J}$, $U = \beta_J^{-q_j}$ in Equation (26) we obtain the Bayesian estimator for the shape parameter β_J for the two distributions given by

$$\begin{aligned} \hat{\beta}_{J(G,q_j)} = & (\beta_J^{-q_j} - (-a + \frac{-1+b}{\alpha_J})q_j\alpha_J^{-q_j-1}\sigma_{21} - (-g + \frac{-1+d}{\beta_J})q_j\alpha_J^{-q_j-1}\sigma_{22} - \frac{1}{2} * q_j * \alpha_J^{-q_j-1} [L_{111}\sigma_{11}\sigma_{12} \\ & + 3L_{122}\sigma_{12}\sigma_{22} + L_{112}(\sigma_{11}\sigma_{22} + 2\sigma_{12}^2) + L_{222}\sigma_{22}^2] + \frac{1}{2}q_j(q_j+1)\beta_J^{-q_j-2}\sigma_{22})^{-\frac{1}{q_j}} \end{aligned} \quad (50)$$

where q_j is constant and $j = 1, 2$.

5 Simulation Studies

To demonstrate the importance of the results obtained in the preceding sections, simulation studies are conducted. For this purpose, we use Monte Carlo method.

The following algorithm is used to generate a sample following progressive type-II censoring scheme, based on any continuous cdf F .

- 1-Generate m independent Uniform (0, 1) observations $W_1, \dots, W_m$.
- 2-Set $V_i = W_i^{1/\gamma_i}$, $\gamma_i = \left(i + \sum_{j=m-i+1}^m R_j\right)$ for $i = 1, 2, \dots, m$.
- 3- $U_i = 1 - V_m V_{m-1} \dots V_{m-i+1}$, $i = 1, 2, \dots, m$.
- 4-Set $X_i = F^{-1}(U_i)$, then X_i , for $i = 1, 2, \dots, m$, are the elements of the sample based on the progressive type-II censoring scheme following the cdf F .
- 5-Repeat steps 1, 2, 3 and 4 (10000) times, for different values of n and m .

$$\text{Estimation average} = \frac{\sum_{i=1}^{10000} \hat{\theta}_i}{10000}, \text{ mean square error} = \frac{\sum_{i=1}^{10000} (\hat{\theta}_i - \theta_i)^2}{10000}, \text{ where, } \theta \text{ is the parameter and } \hat{\theta} \text{ is the estimator.}$$

All the computations are prepared by Mathematica 11.

Since the non-linear equations are not solvable analytically, numerical methods can be used, such as Newton Raphson method with initial values close to the real values of the parameters.

Throughout this section we will use the following abbreviations:

P_1	: non- informative prior,
P_2	: informative prior gamma distribution in case $(a = 5, b = 5, g = 2.25, d = 1.5)$,
L_1	: square error loss function,
L_2	: linear exponential loss function(linex) has two values $(c_1 = 3, (c_2 = -3)$,
L_3	: general entropy loss function has two values $(q_1 = 3, q_2 = -3)$,
MSE	: mean square error,
MLE	: maximum likelihood estimation,
IKUM	: Inverted Kumaraswamy distribution,
EIW	: (Exponentiated inverted Weibull distribution),
IKUM TO EIW	: generating from the distribution (IKUM)to the distribution (EIW),
EIW TO IKUM	: generating from the distribution (EIW)to the distribution (IKUM),
$\hat{\alpha}_{J(MLE)}$	: ML-estimates of α_J ,
$\hat{\beta}_{J(MLE)}$	: ML-estimates of β_J ,
$\hat{\alpha}_{J(sq)}$	: estimate of α_J in case of square error loss function,
$\hat{\beta}_{J(sq)}$	: estimate of β_J in case of square error loss function,
$\hat{\alpha}_{J(L,c_1)}$	: estimate of α_J in case of linex loss function when $c_1 = 3$,
$\hat{\alpha}_{J(L,c_2)}$	: estimate of α_J in case of linex loss function when $c_2 = -3$,
$\hat{\beta}_{J(L,c_1)}$	: estimate of β_J in case of linex loss function when $c_1 = 3$,
$\hat{\beta}_{J(L,c_2)}$	: estimate of β_J in case of linex loss function when $c_2 = -3$,
$\hat{\alpha}_{J(G,q_1)}$	: estimate of α_J in case of general loss function when $q_1 = 3$,
$\hat{\alpha}_{J(G,q_2)}$	: estimate of α_J in case of general loss function when $q_2 = -3$,
$\hat{\beta}_{J(G,q_1)}$	: estimate of β_J in case of general loss function when $q_1 = 3$,
$\hat{\beta}_{J(G,q_2)}$	: estimate of β_J in case of general loss function when $q_2 = -3$.

Table 1: Case of the non- informative prior for the exact value $(\alpha_J = 0.7, \beta_J = 1.3)$ and loss function L_1, L_2 (for values $(c_1 = 3, c_2 = -3)$), L_3 (for values $(q_1 = 3, q_2 = -3)$), $n = 50, m = 25$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(MLE)}$	0. 782354	0. 737692	0. 541503	1. 00816
$\hat{\beta}_{J(MLE)}$	1. 45637	1. 31583	0. 755294	2. 08722
$MSE\hat{\alpha}_{J(MLE)}$	0. 062028	0. 0149798	0. 0481071	0. 173807
$MSE\hat{\beta}_{J(MLE)}$	0. 194132	0. 046699	176. 609	0. 923247
$\hat{\alpha}_{J(sq)}$	0. 764883	0. 732697	0. 52658	0. 996774
$\hat{\beta}_{J(sq)}$	1. 44118	1. 31878	0. 804967	2. 1009
$MSE\hat{\alpha}_{J(sq)}$	0. 00423622	0. 00107059	0. 127473	0. 0899175
$MSE\hat{\beta}_{J(sq)}$	0. 0200994	0. 000385695	2. 5285	0. 655527
$\hat{\alpha}_{J(L,c_1)}$	0. 708723	0. 714892	0. 525683	0. 914762
$\hat{\beta}_{J(L,c_1)}$	1. 29612	1. 25496	0. 734936	1. 78561
$MSE\hat{\alpha}_{J(L,c_1)}$	0. 000120162	0. 000226206	0. 0316456	0. 0453779
$MSE\hat{\beta}_{J(L,c_1)}$	0. 00027242	0. 00210892	0. 320993	0. 237664
$\hat{\alpha}_{J(L,c_2)}$	0. 826532	0. 751022	0. 547507	1. 08103
$\hat{\beta}_{J(L,c_2)}$	1. 59626	1. 38124	0. 778978	2. 39032
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 0161808	0. 00261278	0. 0234182	0. 146873
$MSE\hat{\beta}_{J(L,c_2)}$	0. 0886893	0. 00689066	0. 274157	1. 1955
$\hat{\alpha}_{J(G,q_1)}$	0. 684945	0. 702702	0. 514458	0. 898726
$\hat{\beta}_{J(G,q_1)}$	1. 63246	1. 29664	0. 745555	2. 48684
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 00031214	0. 0000188415	0. 0345852	0. 0349025
$MSE\hat{\beta}_{J(G,q_1)}$	0. 119318	0. 0000902827	0. 307877	1. 41099
$\hat{\alpha}_{J(G,q_2)}$	0. 81839	0. 74932	0. 549107	1. 05627
$\hat{\beta}_{J(G,q_2)}$	1. 77569	1. 39687	0. 802013	2. 649
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 014167	0. 00244081	0. 01984	0. 128792
$MSE\hat{\beta}_{J(G,q_2)}$	0. 230306	0. 00959913	0. 259316	1. 85824

Table 2: Case of the informative prior ($a = 5, b = 5, g = 2.25, d = 1.5$) for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 50, m = 25$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(sq)}$	0. 753343	0. 764108	0. 554958	0. 802727
$\hat{\beta}_{J(sq)}$	1. 36756	1. 25209	0. 791329	1. 57406
$MSE\hat{\alpha}_{J(sq)}$	0. 00286425	0. 00416668	0. 130768	0. 0136397
$MSE\hat{\beta}_{J(sq)}$	0. 00459385	0. 00256437	2. 82841	0. 0881187
$\hat{\alpha}_{J(L,c_1)}$	0. 699718	0. 745664	0. 554624	0. 795932
$\hat{\beta}_{J(L,c_1)}$	1. 2537	1. 20422	0. 719845	1. 62584
$MSE\hat{\alpha}_{J(L,c_1)}$	0. 000124332	0. 00213503	0. 0225257	0. 0110182
$MSE\hat{\beta}_{J(L,c_1)}$	0. 00253476	0. 00964854	0. 33844	0. 110407
$\hat{\alpha}_{J(L,c_2)}$	0. 816435	0. 779835	0. 57479	0. 873607
$\hat{\beta}_{J(L,c_2)}$	1. 54447	1. 3219	0. 763297	2. 07732
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 0136238	0. 0064332	0. 0158967	0. 0320401
$MSE\hat{\beta}_{J(L,c_2)}$	0. 0605235	0. 00048639	0. 29078	0. 607928
$\hat{\alpha}_{J(G,q_1)}$	0. 678508	0. 730727	0. 540461	0. 804616
$\hat{\beta}_{J(G,q_1)}$	1. 26074	1. 20437	0. 708979	1. 61009
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 000614417	0. 000988006	0. 0256152	0. 0122006
$MSE\hat{\beta}_{J(G,q_1)}$	0. 00189076	0. 00959962	0. 348135	0. 101375
$\hat{\alpha}_{J(G,q_2)}$	0. 807832	0. 778537	0. 575862	0. 836466
$\hat{\beta}_{J(G,q_2)}$	1. 45485	1. 2868	0. 762798	1. 64623
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 0116787	0. 00622805	0. 0124284	0. 0218309
$MSE\hat{\beta}_{J(G,q_2)}$	0.	0.	0.	0.

Table 3: Case of the non- informative prior for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 100, m = 50$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(MLE)}$	0. 741824	0. 726905	0. 508555	0. 966764
$\hat{\beta}_{J(MLE)}$	1. 38127	1. 30142	1. 16632	1. 97296
$MSE\hat{\alpha}_{J(MLE)}$	0. 0229283	0. 00715946	0. 0496857	0. 106791
$MSE\hat{\beta}_{J(MLE)}$	0. 0720253	0. 0233523	0. 0337782	0. 576451
$\hat{\alpha}_{J(sq)}$	0. 732833	0. 72423	0. 505874	0. 957314
$\hat{\beta}_{J(sq)}$	1. 37255	1. 30244	1. 17126	1. 9678
$MSE\hat{\alpha}_{J(sq)}$	0. 00107954	0. 000587277	0. 0377195	0. 066227
$MSE\hat{\beta}_{J(sq)}$	0. 0052747	8. 30705*10 ⁻⁶	0. 0194067	0. 446123
$\hat{\alpha}_{J(L,c_1)}$	0. 705025	0. 715596	0. 502063	0. 918371
$\hat{\beta}_{J(L,c_1)}$	1. 30039	1. 2713	1. 1412	1. 81001
$MSE\hat{\alpha}_{J(L,c_1)}$	0. 0000309385	0. 000243753	0. 0392029	0. 0477088
$MSE\hat{\beta}_{J(L,c_1)}$	0. 0000369162	0. 000829364	0. 0220974	0. 260453
$\hat{\alpha}_{J(L,c_2)}$	0. 762102	0. 733	0. 509715	0. 998285
$\hat{\beta}_{J(L,c_2)}$	1. 44806	1. 33337	1. 19633	2. 12855
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 00387323	0. 0010902	0. 0362661	0. 0890698
$MSE\hat{\beta}_{J(L,c_2)}$	0. 0220413	0. 00113912	0. 0113387	0. 687438
$\hat{\alpha}_{J(G,q_1)}$	0. 687982	0. 708964	0. 496315	0. 90685
$\hat{\beta}_{J(G,q_1)}$	1. 44931	1. 29221	1. 15715	2. 08921
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 000159694	0. 0000819137	0. 0415032	0. 0428278
$MSE\hat{\beta}_{J(G,q_1)}$	0. 0225951	0. 0000688465	0. 0200852	0. 626154
$\hat{\alpha}_{J(G,q_2)}$	0. 759291	0. 732285	0. 510851	0. 986271
$\hat{\beta}_{J(G,q_2)}$	1. 54276	1. 34087	1. 2054	2. 24274
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 00353005	0. 00104338	0. 0358643	0. 0820196
$MSE\hat{\beta}_{J(G,q_2)}$	0. 0594643	0. 00168219	0. 0107436	0. 892197

Table 4: Case of the informative prior ($a = 5, b = 5, g = 2.25, d = 1.5$) for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 100, m = 50$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(sq)}$	0.737225	0.739426	0.521221	0.885239
$\hat{\beta}_{J(sq)}$	1.35356	1.27154	1.1394	1.77227
$MSE\hat{\alpha}_{J(sq)}$	0.00138646	0.00155892	0.0320107	0.0345245
$MSE\hat{\beta}_{J(sq)}$	0.00287351	0.000830179	0.0283527	0.223834
$\hat{\alpha}_{J(L,c_1)}$	0.708988	0.730624	0.517475	0.861584
$\hat{\beta}_{J(L,c_1)}$	1.28582	1.24427	1.11242	1.70859
$MSE\hat{\alpha}_{J(L,c_1)}$	0.0000902551	0.000941795	0.0333478	0.0263644
$MSE\hat{\beta}_{J(L,c_1)}$	0.000244843	0.00315656	0.0333302	0.167854
$\hat{\alpha}_{J(L,c_2)}$	0.766217	0.74759	0.524632	0.925361
$\hat{\beta}_{J(L,c_2)}$	1.43217	1.30408	1.16641	1.97776
$MSE\hat{\alpha}_{J(L,c_2)}$	0.00439583	0.00227057	0.030837	0.05085
$MSE\hat{\beta}_{J(L,c_2)}$	0.0175773	0.0000171227	0.0183237	0.459538
$\hat{\alpha}_{J(G,q_1)}$	0.691277	0.723281	0.51107	0.857539
$\hat{\beta}_{J(G,q_1)}$	1.28785	1.2438	1.1089	1.71836
$MSE\hat{\alpha}_{J(G,q_1)}$	0.0000968041	0.000546192	0.0357128	0.0250495
$MSE\hat{\beta}_{J(G,q_1)}$	0.000188515	0.00320982	0.0362909	0.176003
$\hat{\alpha}_{J(G,q_2)}$	0.763467	0.746957	0.525538	0.911456
$\hat{\beta}_{J(G,q_2)}$	1.39343	1.28807	1.15452	1.83182
$MSE\hat{\alpha}_{J(G,q_2)}$	0.00403757	0.00221058	0.0305812	0.0448225
$MSE\hat{\beta}_{J(G,q_2)}$	0.	0.	0.	0.

Table 5: Case of the non- informative prior for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 150, m = 75$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(MLE)}$	0.72112	0.713616	0.508697	0.954735
$\hat{\beta}_{J(MLE)}$	1.34085	1.30254	1.16488	1.95855
$MSE\hat{\alpha}_{J(MLE)}$	0.01442	0.00454603	0.0437057	0.0875172
$MSE\hat{\beta}_{J(MLE)}$	0.0380325	0.0140043	0.0286881	0.512945
$\hat{\alpha}_{J(sq)}$	0.715059	0.711814	0.507016	0.94822
$\hat{\beta}_{J(sq)}$	1.33469	1.30311	1.16689	1.95443
$MSE\hat{\alpha}_{J(sq)}$	0.000227018	0.000139639	0.0372431	0.0616152
$MSE\hat{\beta}_{J(sq)}$	0.00120538	0.0000102245	0.0177357	0.428302
$\hat{\alpha}_{J(L,c_1)}$	0.696944	0.706215	0.504541	0.92207
$\hat{\beta}_{J(L,c_1)}$	1.28834	1.2824	1.14797	1.84278
$MSE\hat{\alpha}_{J(L,c_1)}$	0.0000108459	0.0000387813	0.0382046	0.0493224
$MSE\hat{\beta}_{J(L,c_1)}$	0.000145529	0.000310845	0.0231199	0.294739
$\hat{\alpha}_{J(L,c_2)}$	0.733825	0.717473	0.509515	0.975349
$\hat{\beta}_{J(L,c_2)}$	1.38265	1.32374	1.18548	2.06815
$MSE\hat{\alpha}_{J(L,c_2)}$	0.00114759	0.000305647	0.0362848	0.0758395
$MSE\hat{\beta}_{J(L,c_2)}$	0.00685836	0.000569508	0.0131836	0.590367
$\hat{\alpha}_{J(G,q_1)}$	0.684075	0.701615	0.500722	0.913288
$\hat{\beta}_{J(G,q_1)}$	1.38188	1.2965	1.15931	2.03035
$MSE\hat{\alpha}_{J(G,q_1)}$	0.000258179	3.08229×10^{-6}	0.0397131	0.0455055
$MSE\hat{\beta}_{J(G,q_1)}$	0.00675703	0.0000143405	0.0198018	0.533851
$\hat{\alpha}_{J(G,q_2)}$	0.732456	0.717105	0.510288	0.967457
$\hat{\beta}_{J(G,q_2)}$	1.44646	1.32865	1.19176	2.14379
$MSE\hat{\alpha}_{J(G,q_2)}$	0.0010565	0.00029288	0.0359914	0.071548
$MSE\hat{\beta}_{J(G,q_2)}$	0.0215836	0.000823355	0.0118313	0.712986

Table 6: Case of the informative prior ($a = 5, b = 5, g = 2.25, d = 1.5$) for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 150, m = 75$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(sq)}$	0. 721325	0. 722029	0. 51694	0. 902605
$\hat{\beta}_{J(sq)}$	1. 32762	1. 28272	1. 14507	1. 82928
$MSE\hat{\alpha}_{J(sq)}$	0. 000454966	0. 000486429	0. 0335132	0. 0411032
$MSE\hat{\beta}_{J(sq)}$	0. 000764375	0. 000303652	0. 0240175	0. 280342
$\hat{\alpha}_{J(L,c_1)}$	0. 702823	0. 71636	0. 514491	0. 883198
$\hat{\beta}_{J(L,c_1)}$	1. 28235	1. 26376	1. 12794	1. 76476
$MSE\hat{\alpha}_{J(L,c_1)}$	$9. 65932 \times 10^{-6}$	0. 000268601	0. 0344149	0. 0336426
$MSE\hat{\beta}_{J(L,c_1)}$	0. 00032072	0. 00132675	0. 0296689	0. 216368
$\hat{\alpha}_{J(L,c_2)}$	0. 739803	0. 727419	0. 519267	0. 929467
$\hat{\beta}_{J(L,c_2)}$	1. 37635	1. 30404	1. 16442	1. 96373
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 00158744	0. 000753381	0. 0326684	0. 0526694
$MSE\hat{\beta}_{J(L,c_2)}$	0. 00585736	0. 0000164744	0. 0183957	0. 440584
$\hat{\alpha}_{J(G,q_1)}$	0. 689226	0. 711429	0. 510397	0. 878167
$\hat{\beta}_{J(G,q_1)}$	1. 28263	1. 26339	1. 12583	1. 77766
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 000121119	0. 000131665	0. 0359498	0. 0318255
$MSE\hat{\beta}_{J(G,q_1)}$	0. 000310608	0. 00135394	0. 0304055	0. 228516
$\hat{\alpha}_{J(G,q_2)}$	0. 738481	0. 727079	0. 519959	0. 920811
$\hat{\beta}_{J(G,q_2)}$	1. 35273	1. 29357	1. 15609	1. 87121
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 00148355	0. 000734851	0. 0324191	0. 048783
$MSE\hat{\beta}_{J(G,q_2)}$	0.	0.	0.	0.

Table 7: Case of the non- informative prior for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 200, m = 100$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(MLE)}$	0. 718539	0. 711885	0. 502877	0. 945524
$\hat{\beta}_{J(MLE)}$	1. 33538	1. 30183	1. 15957	1. 9347
$MSE\hat{\alpha}_{J(MLE)}$	0. 0108399	0. 00307613	0. 04418	0. 0767062
$MSE\hat{\beta}_{J(MLE)}$	0. 0289834	0. 0114091	0. 0268066	0. 455624
$\hat{\alpha}_{J(sq)}$	0. 713978	0. 710514	0. 501567	0. 940508
$\hat{\beta}_{J(sq)}$	1. 33072	1. 30226	1. 16111	1. 93109
$MSE\hat{\alpha}_{J(sq)}$	0. 000195487	0. 000110579	0. 0393759	0. 0578444
$MSE\hat{\beta}_{J(sq)}$	0. 000944584	$5. 34086 \times 10^{-6}$	0. 0193462	0. 398286
$\hat{\alpha}_{J(L,c_1)}$	0. 700283	0. 706346	0. 499741	0. 921028
$\hat{\beta}_{J(L,c_1)}$	1. 29553	1. 28667	1. 14766	1. 84695
$MSE\hat{\alpha}_{J(L,c_1)}$	$7. 40628 \times 10^{-7}$	0. 0000403393	0. 0401042	0. 0488566
$MSE\hat{\beta}_{J(L,c_1)}$	0. 0000242805	0. 000178157	0. 0233552	0. 29921
$\hat{\alpha}_{J(L,c_2)}$	0. 728045	0. 714717	0. 503408	0. 960558
$\hat{\beta}_{J(L,c_2)}$	1. 36685	1. 3178	1. 17442	2. 01674
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 000787991	0. 000216744	0. 0386491	0. 0678982
$MSE\hat{\beta}_{J(L,c_2)}$	0. 00448083	0. 000319554	0. 0158608	0. 513853
$\hat{\alpha}_{J(G,q_1)}$	0. 690044	0. 702855	0. 496836	0. 913975
$\hat{\beta}_{J(G,q_1)}$	1. 36549	1. 29731	1. 15584	1. 98539
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 000101263	$8. 3781 \times 10^{-6}$	0. 0412767	0. 0457913
$MSE\hat{\beta}_{J(G,q_1)}$	0. 00430885	$8. 11502 \times 10^{-6}$	0. 0208654	0. 469892
$\hat{\alpha}_{J(G,q_2)}$	0. 727056	0. 714452	0. 504005	0. 954799
$\hat{\beta}_{J(G,q_2)}$	1. 41562	1. 32146	1. 17905	2. 0726
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 00073331	0. 000209004	0. 0384148	0. 0649276
$MSE\hat{\beta}_{J(G,q_2)}$	0. 0134273	0. 000461414	0. 01482	0. 597297

Table 8: Case of the informative prior ($a = 5, b = 5, g = 2.25, d = 1.5$) for the exact value ($\alpha_J = 0.7, \beta_J = 1.3$) and loss function L1, L2 (for values ($c_1 = 3, c_2 = -3$)), L3 (for values ($q_1 = 3, q_2 = -3$)), $n = 200, m = 100$.

	IKUM	EIW	IKUM TO EIW	EIW TO IKUM
$\hat{\alpha}_{J(sq)}$	0. 718995	0. 718148	0. 508938	0. 908431
$\hat{\beta}_{J(sq)}$	1. 32594	1. 28701	1. 14557	1. 84288
$MSE\hat{\alpha}_{J(sq)}$	0. 000360889	0. 000329774	0. 0365071	0. 0434624
$MSE\hat{\beta}_{J(sq)}$	0. 00067376	0. 000170948	0. 0239058	0. 294786
$\hat{\alpha}_{J(L,c_1)}$	0. 705066	0. 71394	0. 507125	0. 892527
$\hat{\beta}_{J(L,c_1)}$	1. 29131	1. 27242	1. 13337	1. 78542
$MSE\hat{\alpha}_{J(L,c_1)}$	0. 000026361	0. 000194689	0. 0372022	0. 037099
$MSE\hat{\beta}_{J(L,c_1)}$	0. 0000794169	0. 000766842	0. 0273954	0. 235804
$\hat{\alpha}_{J(L,c_2)}$	0. 732885	0. 722201	0. 510683	0. 928398
$\hat{\beta}_{J(L,c_2)}$	1. 36247	1. 30294	1. 15934	1. 9397
$MSE\hat{\alpha}_{J(L,c_2)}$	0. 00108284	0. 000493492	0. 0358445	0. 0521704
$MSE\hat{\beta}_{J(L,c_2)}$	0. 0039152	8.69927×10^{-6}	0. 0198225	0. 409229
$\hat{\alpha}_{J(G,q_1)}$	0. 694365	0. 710262	0. 504062	0. 887639
$\hat{\beta}_{J(G,q_1)}$	1. 29139	1. 27211	1. 13759	1. 79815
$MSE\hat{\alpha}_{J(G,q_1)}$	0. 0000340376	0. 000105757	0. 0383923	0. 0352435
$MSE\hat{\beta}_{J(G,q_1)}$	0. 0000780421	0. 000784167	0. 0721376	0. 248307
$\hat{\alpha}_{J(G,q_2)}$	0. 731923	0. 721951	0. 511233	0. 922242
$\hat{\beta}_{J(G,q_2)}$	1. 34469	1. 29513	1. 15346	1. 87439
$MSE\hat{\alpha}_{J(G,q_2)}$	0. 0010203	0. 000482466	0. 035637	0. 0493997
$MSE\hat{\beta}_{J(G,q_2)}$	0.	0.	0.	0.

6 Conclusion

From the simulation studies we noted that:

In general, the Bayesian estimators have mean square error less than that of the MLE. Increasing the sample size leads to decrease mean square error and increase the accuracy of estimators.

To study of the effect of the distribution model on the robustness of Bayes estimation, two different distributions are chosen. (IKUM) distribution and (EIW) distribution each having two shape parameters.

Values of the parameters are chosen such that the two distributions are close in shape. Data are generated from each distribution and used to estimate the parameters of the two distributions, for the same prior distribution and loss function.

The results are shown for $((n, m) = (50, 25), (100, 50), (150, 75) \text{ and } (200, 100))$ in tables (1, 3, 5, 7) for non-informative prior and tables (2, 4, 6, 8) for informative prior.

The results in each table present the estimates when sampling from one distribution and use the data to estimate the parameters of the two distributions, (when sampling from (IKUM) distribution we get estimates in column 3 and column 5).

Similarly for the (EIW) distribution estimates are in columns 2 and 4) with the mean square error in each case.

The mean square error of parameter α in case generate data from the distribution (IKUM) to the distribution (EIW) is the best.

While the mean square error of parameter β in case generate data from the distribution (EIW) to the distribution (IKUM) is the best.

The effect of the prior information on the robustness of Bayes estimation can be seen in this study by comparing the results in tables (1, 3, 5 and 7) with the corresponding values in tables (2, 4, 6 and 8).

In case of non-informative prior and informative gamma prior with different values of shape-parameters, with the loss function and distribution are deterministic.

To study the effect of the element of loss on the robustness of Bayes estimates, the results of using different loss functions (square error, LINEX, general entropy loss functions) are also shown in all tables when fixing the distribution model and the prior distribution.

The effect of increasing the number of items put on test (n) and the number of observed failures (m) can be shown by comparing the results in all tables.

The simulation also stresses the importance of robust as shown in the case studied.

Conflicts of Interests

The authors declare that they have no conflicts of interests

References

- [1] D. Rios Insua, J. O. Bergerand F. Ruggeri, "Robust Bayesian Analysis", Lecture Notes in Statistics, Springer-Verlag New York, vol 152, (2000).
 - [2] Y. Vander Heyden, A. Nijhuis, J. Smeyers-Verbeke, B. G. M. Vandeginste and D. L. Massart, "Guidance for Robustness/Ruggedness Tests in Method Validation", *J Pharm Biomed Anal.*, **24**(5-6), 723-53, (2001 Mar).
 - [3] A . Aljuaid, "Estimating the Parameters of an Exponentiated Inverted Weibull Distribution under Type-II Censoring", *Applied Mathematical Sciences*, Vol. **7**, no. 35, 1721 - 1736, (2013).
 - [4] A. M. Abd AL-Fattah, A. A. EL-Helbawy and G. R. AL-Dayian, "Inverted Kumaraswamy Distribution: Properties And Estimation", *Pak. J. Statist.*, Vol. **33**(1), 37-61, (2017).
 - [5] M. S. Khan, G. R. Pasha and A. H. Pasha, "Theoretical analysis of inverse weibull distribution", *WSEAS Transactions on Mathematics*, **7**, 2, (2008).
 - [6] E. Cramer and G. Iliopoulos. "Adaptive progressive Type-II censoring. " *Test*, **19**, no. 2 , 342-358, (2010).
 - [7] P. Kumaraswamy, "A generalized probability density function for double-bounded random processes", *Journal of Hydrology*, **46** (1-2), 79-88, (1980).
 - [8] D. V. Lindley, " Approximate Bayesian Methods", *Trabajos de Estadística y de Investigación Operativa*, Vol **31**, No:1, pp. 223-245, (1980).
-