

Research Paper

Reliability analysis of Kumaraswamy distribution under progressive first-failure censoring

SOMAYEH GHAFOURI^{*1}, MANOJ KUMAR RASTOGI²

¹DEPARTMENT OF MATHEMATICS, FACULTY OF SCIENCES,
ARAK UNIVERSITY, ARAK 38156-8-8349, IRAN

²DEPARTMENT OF STATISTICS, PATNA UNIVERSITY, PATNA, BIHAR, 800005, INDIA

Received: September 17, 2020/ Revised: August 18, 2021 / Accepted: September 08, 2021

Abstract: In this article, we consider the estimation of the parameters and reliability characteristics of Kumaraswamy distribution using progressive first failure censored samples. First, we derive the maximum likelihood estimates using an expectation-maximization algorithm and compute the observed information of the parameters that can be used for constructing asymptotic confidence intervals. We also compute the Bayes estimates of the parameters using Lindley approximation as well as the Metropolis-Hastings algorithm. Furthermore, we derive the highest posterior density credible intervals. Simulation studies are conducted to evaluate the performance of the point and interval estimators. Finally, two examples of real data sets are provided to illustrate the proposed procedures.

Keywords: Bayes estimation; EM algorithm; Kumaraswamy distribution; Maximum likelihood method; Progressive first-failure censoring.

Mathematics Subject Classification (2010): 30C80, 62F10, 62F15, 62N01, 65C05.

1 Introduction

Censoring is very useful for life testing experiments due to the cost and time restrictions on the collection of data. In the conventional Type I and Type II censoring schemes, there is not any flexibility on the removal of items at stages other than the terminal stage of the test. To overcome this difficulty, the progressive Type II censoring scheme was introduced in the literature. For more details on this, one can refer

^{*}Corresponding author: s_ghafouri@araku.ac.ir

to Cohen and Norgard (1977), Soliman (2005), Balakrishnan and Aggarwala (2000) and Balakrishnan and Cramer (2014). Progressive Type II censoring is a generalization of Type II censoring in which the removal of units is allowed in the middle of an experiment. Moreover, Johnson (1964) introduced a first-failure censoring in which the experimenter may be interested in testing units into several groups, each group as a collection of test units, and then run all the test units simultaneously until the first failure is observed in each group. The first-failure censoring does not allow groups to be omitted from the test at stages other than the final termination stage. However, this is much desirable in practical life. This censoring scheme is more applicable where the lifetime of a product is very high and examination facilities are sparse but the examination material is relatively cheap (see, Balasooriya (1995)). Some recent works on progressive first-failure censoring could be found in Ahmadi et al. (2013) and Heba et al. (2017).

A first-failure censoring scheme can be described as follows. Suppose that n independent groups with k units in each group are put on a life test. r_1 groups and the group in which the first failure (say $x_{1:m:n:k}^{r_1}$) occurs has been removed randomly, r_2 groups and the group in which the second failure (say $x_{2:m:n:k}^{r_2}$) occurs has been removed randomly. This procedure is continued until the m th failure where r_m ($m \leq n$) groups and the group in which the m th failure (say $x_{m:m:n:k}^{r_m}$) occurs has been randomly removed. So $(x_{1:m:n:k}^{r_1}, x_{2:m:n:k}^{r_2}, \dots, x_{m:m:n:k}^{r_m})$ are progressively first-failure observed sample and $n = m + r_1 + r_2 + \dots + r_m$ ($r_1, r_2, \dots, r_m$ are pre-determined).

If the failure times of $n \times k$ items originally in the test are from a continuous distribution function $F_x(x)$ and probability density function $f_X(x)$, then the joint probability density function of $x_{1:m:n:k}^{r_1}, x_{2:m:n:k}^{r_2}, \dots, x_{m:m:n:k}^{r_m}$ (For convenience, let us show them by $x_{1:m:n:k}, x_{2:m:n:k}, \dots, x_{m:m:n:k}$) is given by

$$f_{1,2,\dots,m}(x_{1:m:n:k}, x_{2:m:n:k}, \dots, x_{m:m:n:k}) = ck^m \prod_{i=1}^m f(x_{i:m:n:k}) (1 - F(x_{i:m:n:k}))^{k(r_i+1)-1}, \quad (1)$$

where $x_{i:m:n:k} > 0$, $i = 1, 2, \dots, m$ and

$$c = n(n - r_1 - 1)(n - r_1 - r_2 - 2) \dots (n - r_1 - r_2 \dots r_{m-1} - m + 1).$$

A random variable T is said to have a Kumaraswamy distribution with two positive shape parameters α and β , if its probability density function (pdf) and cumulative distribution function are given, respectively, by

$$\begin{aligned} f_T(t) &= \alpha \beta t^{\beta-1} (1 - t^\beta)^{\alpha-1}, \quad 0 < t < 1, \\ F_T(t) &= 1 - (1 - t^\beta)^\alpha, \quad 0 < t < 1. \end{aligned} \quad (2)$$

Then, the corresponding reliability function and hazard function of X are given, respectively, by

$$R(t) = (1 - t^\beta)^\alpha, \quad 0 < t < 1, \quad (3)$$

$$h(t) = \alpha \beta t^{\beta-1} (1 - t^\beta)^{-1}, \quad 0 < t < 1. \quad (4)$$

The Kumaraswamy distribution is suitable for outcomes with lower and upper bounds from natural phenomena and applications such as average height of individuals in a certain population, scores of a test, atmospheric temperatures, meteorological

inference, hydrological data, economic data for example unemployment data, etc. For details see Sindhu et al. (2013).

The distribution have many probabilistic similarities with the beta distribution. For instance, its density function is similar to the beta distribution, i.e., depending on the different values of its shape parameters, can be unimodal, increasing, decreasing, or constant. The choices of shape parameters can be transformed the Kumaraswamy distribution to various distributions, such as uniform, beta, exponential, and generalized beta distributions. One of the advantages of the Kumaraswamy distribution over the beta distribution is that it has a simple and closed form of distribution function, and so, corresponding quantiles can be easily obtained. Also, according to Gilchrist (1997), the models with closed-form distribution functions have advantage of being used in modeling of reliability data. For more details see Sultana et al. (2017) and Reyad and Ahmed (2016). The basic properties of this distribution such as properties of its skewness and kurtosis, the maximum likelihood estimators (MLEs) for its parameters along with summarized similarities and differences between the Beta and Kumaraswamy distributions were discussed by Jones (2009).

The interest of this paper is to provide classical and Bayesian inferences for the parameters of Kumaraswamy distribution using the first-failure censored samples. We first obtain the ML estimates of the parameters and their approximate confidence intervals. Also, by considering various symmetric and asymmetric loss functions, some expressions are provided as the Bayes estimates of the parameters and reliability characteristics of the model. Since these expressions can not simplified to nice closed forms, we employ Lindley method and M-H algorithm to compute the Bayes estimates. In addition, HPD credible intervals are derived.

The rest of the contents of this paper is provided as follows. In Section 2, the MLEs of unknown parameters as well as the Fisher information matrix are computed. In Section 3, we derive Bayesian estimates of unknown parameters in terms of different loss functions. A simulation study and two real data sets are analyzed in Section 4.

2 Maximum likelihood estimators

In this section, we derive the MLEs of unknown parameters α and β of the Kumaraswamy distribution under progressive first-failure censoring with joint probability density function in (1). Let $X_{1:m:n:k}, X_{2:m:n:k}, \dots, X_{m:m:n:k}$ denotes a progressively first-failure censored sample from the Kumaraswamy distribution with censoring scheme $(r_1, r_2, \dots, r_m)$. Using Balakrishnan and Aggarwala (2000), the likelihood function of α and β can be obtained as

$$L(\alpha, \beta \mid \mathbf{x}) = ck^m \alpha^m \beta^m \prod_{i=1}^m x_{i:m:n:k}^{\beta-1} \left(1 - x_{i:m:n:k}^{\beta}\right)^{\alpha k(r_i+1)-1},$$

where $\mathbf{x} = (x_{1:m:n:k}, x_{2:m:n:k}, \dots, x_{m:m:n:k})$ denotes the observed value of $\mathbf{X} = (X_{1:m:n:k}, X_{2:m:n:k}, \dots, X_{m:m:n:k})$. The log-likelihood function can be rewritten as (ignoring constant terms)

$$l(\alpha, \beta \mid \mathbf{x}) = m \log \alpha + m \log \beta + (\beta - 1) \sum_{i=1}^m \log x_{i:m:n:k}$$

$$+ \sum_{i=1}^m \{\alpha k(r_i + 1) - 1\} \log(1 - x_{i:m:n:k}^\beta),$$

and the corresponding likelihood equations for unknown parameters are obtained as

$$\frac{\partial l}{\partial \alpha} = \frac{m}{\alpha} + k \sum_{i=1}^m (r_i + 1) \log(1 - x_{i:m:n:k}^\beta) = 0, \quad (5)$$

$$\begin{aligned} \frac{\partial l}{\partial \beta} &= \frac{m}{\beta} + \sum_{i=1}^m \log x_{i:m:n:k} - \sum_{i=1}^m \{\alpha k(r_i + 1) - 1\} x_{i:m:n:k}^\beta \\ &\quad \times (1 - x_{i:m:n:k}^\beta)^{-1} \log x_{i:m:n:k} = 0. \end{aligned} \quad (6)$$

From (5), we find the MLE of α as

$$\hat{\alpha} = \left[\frac{-k \sum_{i=1}^m (r_i + 1) \log(1 - x_{i:m:n:k}^\beta)}{m} \right]^{-1}. \quad (7)$$

Substituting (7) into (6), the MLE of β can be obtained by solving the following equation:

$$\frac{m}{\hat{\beta}} + \sum_{i=1}^m \log x_{i:m:n:k} - \sum_{i=1}^m \{\hat{\alpha} k(r_i + 1) - 1\} x_{i:m:n:k}^{\hat{\beta}} (1 - x_{i:m:n:k}^{\hat{\beta}})^{-1} \log x_{i:m:n:k} = 0. \quad (8)$$

It is observed that the analytical solution of (8) is not in a closed form. So some numerical techniques for instance, the Newton- Raphson and Broydan method may be used. We applied the package “nleqslv” in (software) R to find the solution of unknown parameters α and β (see Ghitany et al., 2013). For more details, references and discussions of mentioned numerical techniques, see Huang and Wu (2012) and Balakrishnan et al. (2009). Now as a consequence, the MLEs of $R(t)$ and $h(t)$, respectively defined as $\hat{R}(t)$ and $\hat{h}(t)$, are obtained as

$$\hat{R}(t) = (1 - t^{\hat{\beta}})^{\hat{\alpha}} \quad \text{and} \quad \hat{h}(t) = \hat{\alpha} \hat{\beta} t^{\hat{\beta}-1} (1 - t^{\hat{\beta}})^{-1}, \quad 0 < t < 1.$$

In the next section, we use the expectation maximization (EM) algorithm for MLE.

2.1 EM algorithm

In this article, we use the EM algorithm to compute desired estimators. Dempster et al. (1977) introduced the EM algorithm to handle the incomplete or missing data problem. Derivation of MLEs for unknown parameters of Kumaraswamy distribution can be considered as a missing value problem in the sense of Ng. et al. (2002). Suppose $\mathbf{x} = (x_{1:m:n}, x_{2:m:n}, \dots, x_{m:m:n})$ is observed data and $\mathbf{Z} = (\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_m)$ is censored data. Note that $\mathbf{Z}_j$ is a $1 \times r_j$ vector $(Z_{j1}, Z_{j2}, \dots, Z_{jr_j}), j = 1, 2, \dots, m$. Observe that $(\mathbf{x}, \mathbf{Z})$ constitutes the complete data set, call it $\mathbf{W}$. The likelihood function using complete data may be taken in the form of

$$LW_c(\alpha, \beta | \mathbf{W}) \propto \prod_{i=1}^m f(x_{i:m:n:k}) \left(1 - F(x_{i:m:n:k})\right)^{k(R_i+1)-1}$$

$$\begin{aligned}
&= \prod_{i=1}^m f(x_{i:m:n:k}) \left(f(w_j | w_j > x_{i:m:n:k}) \right)^{k(R_i+1)-1}, \\
&= \prod_{i=1}^m [f(x_{i:m:n:k}) (f(Z_{ij}))^{k(R_i+1)-1}] \\
&= \prod_{i=1}^m \left[\alpha \beta x_{i:m:n:k}^{\beta-1} (1 - x_{i:m:n:k}^\beta)^{\alpha-1} \prod_{j=1}^{k(R_i+1)-1} \alpha \beta Z_{ij}^{\beta-1} (1 - Z_{ij}^\beta)^{\alpha-1} \right]
\end{aligned}$$

The corresponding log-likelihood function of $\mathbf{W}$ is

$$\begin{aligned}
L_c(\alpha, \beta | \mathbf{W}) &= nk \log \alpha + nk \log \beta \\
&+ (\beta - 1) \left(\sum_{i=1}^m \log x_{i:m:n:k} + \sum_{i=1}^m \sum_{j=1}^{k(r_i+1)-1} \log Z_{ij} \right) \\
&+ (\alpha - 1) \left(\sum_{i=1}^m \log(1 - x_{i:m:n:k}^\beta) + \sum_{i=1}^m \sum_{j=1}^{k(r_i+1)-1} \log(1 - Z_{ij}^\beta) \right).
\end{aligned}$$

The E-step of the EM algorithm involves computation of pseudo log-likelihood function. The involved expectations are computed as

$$\begin{aligned}
E(\log Z_{ij} | Z_{ij} > c) &= \frac{\alpha}{\beta(1 - F_X(c; \alpha, \beta))} \int_0^{1-c^\beta} u^{\alpha-1} \log(1 - u) du \\
&= A(c, \alpha, \beta), \quad \text{say,} \tag{9}
\end{aligned}$$

$$E(\log(1 - Z_{ij}^\beta) | Z_{ij} > c) = \log(1 - c^\beta) - \frac{1}{\alpha} = B(c, \alpha, \beta), \quad \text{say.} \tag{10}$$

The pseudo log-likelihood function denoted by $L_s(\alpha, \beta)$ and observe that

$$\begin{aligned}
L_s(\alpha, \beta) &= nk \log \alpha + nk \log \beta + (\beta - 1) \sum_{i=1}^m \log x_{i:m:n:k} \\
&+ (\alpha - 1) \sum_{i=1}^m \log(1 - x_{i:m:n:k}^\beta) \\
&+ (\alpha - 1) \sum_{i=1}^m \sum_{j=1}^{k(r_i+1)-1} E(\log(1 - Z_{ij}^\beta) | Z_{ij} > x_{i:m:n:k}) \\
&+ (\beta - 1) \sum_{i=1}^m \sum_{j=1}^{k(r_i+1)-1} E(\log Z_{ij} | Z_{ij} > x_{i:m:n:k}). \tag{11}
\end{aligned}$$

In the M-step of the EM algorithm, the function $L_s(\alpha, \beta)$ is maximized with respect to α and β . Now, if the estimates of (α, β) at the j th stage are $(\alpha^{(j)}, \beta^{(j)})$, then $(\alpha^{(j+1)}, \beta^{(j+1)})$ can be derived by maximizing the equation

$$G(\alpha, \beta) = nk \log \alpha + nk \log \beta + (\beta - 1) \sum_{i=1}^m \log x_{i:m:n:k}$$

$$+(\alpha - 1) \sum_{i=1}^m \log(1 - x_{i:m:n:k}^\beta) + (\beta - 1)\tilde{A} + (\alpha - 1)\tilde{B}, \quad (12)$$

where

$$\begin{aligned} \tilde{A} &= \sum_{i=1}^m (k(r_i + 1) - 1) A(x_{i:m:n:k}, \alpha^{(j)}, \beta^{(j)}), \\ \tilde{B} &= \sum_{i=1}^m (k(r_i + 1) - 1) B(x_{i:m:n:k}, \alpha^{(j)}, \beta^{(j)}). \end{aligned}$$

The maximization of (12) can be easily obtained as follows. Using the fixed point method (see, Pradhan and Kundu (2009)), we observe that the updated estimate of β , namely $\beta^{(j+1)}$, can be obtained by solving the equation

$$h(\beta) = \beta, \quad (13)$$

where

$$\begin{aligned} h(\beta) &= \left[-\frac{1}{nk} \sum_{i=1}^m \log x_{i:m:n:k} - \frac{\tilde{A}}{nk} \right. \\ &\quad \left. + \frac{(\hat{\alpha}(\beta) - 1)}{nk} \sum_{i=1}^m x_{i:m:n:k}^\beta (1 - x_{i:m:n:k}^\beta)^{-1} \log x_{i:m:n:k} \right]^{-1}, \end{aligned}$$

and $\hat{\alpha}(\beta) = nk / (\sum_{i=1}^m \log(1 - e^{-\beta x_{i:m:n:k}^{-2}}) + \tilde{B})$. Subsequently, the updated estimate of α can be computed using the expression $\alpha^{(j+1)} = \hat{\alpha}(\beta^{(j+1)})$.

Let $\boldsymbol{\theta} = (\alpha, \beta)$. Next, we derive the observed Fisher information matrix $I_{\mathbf{X}}(\boldsymbol{\theta})$ using the missing value principle from Louis (1982). It can be used for construction of the asymptotic confidence intervals of the unknown parameters. Suppose that $I_{\mathbf{W}}(\boldsymbol{\theta})$ denotes the complete information and $I_{\mathbf{W}|\mathbf{X}}(\boldsymbol{\theta})$ represents the missing information. Let $I_{\mathbf{W}|\mathbf{X}}^{i:m:n:k}(\boldsymbol{\theta})$ is the representative of the Fisher information for single observation being censored at the time of the i th failure $x_{i:m:n:k}$. Then, we have

$$I_{\mathbf{X}}(\boldsymbol{\theta}) = I_{\mathbf{W}}(\boldsymbol{\theta}) - I_{\mathbf{W}|\mathbf{X}}(\boldsymbol{\theta}), \quad (14)$$

where

$$\begin{aligned} I_{\mathbf{W}}(\boldsymbol{\theta}) &= -E \left[\frac{\partial^2 L_c(\boldsymbol{\theta}|\mathbf{W})}{\partial \boldsymbol{\theta}^2} \right], \\ I_{\mathbf{W}|\mathbf{X}}(\boldsymbol{\theta}) &= \sum_{i=1}^m (k(r_i + 1) - 1) I_{\mathbf{W}|\mathbf{X}}^{i:m:n:k}(\boldsymbol{\theta}), \\ I_{\mathbf{W}|\mathbf{X}}^{i:m:n:k}(\boldsymbol{\theta}) &= -E_{Z_i|X_{i:m:n:k}} \left[\frac{\partial^2 \log f_{Z_i}(z_i | x_{i:m:n:k}, \boldsymbol{\theta})}{\partial \boldsymbol{\theta}^2} \right]. \end{aligned}$$

Elements of both matrices are presented below. We denote (i, j) th, $i, j = 1, 2$, element of $I_{\mathbf{W}}(\boldsymbol{\theta})$ by $a_{ij}(\alpha, \beta)$, where

$$a_{11}(\alpha, \beta) = \frac{nk}{\alpha^2}, \quad a_{22}(\alpha, \beta) = \frac{nk}{\beta^2} + \frac{nk\alpha(\alpha - 1)}{\beta^2} \int_0^1 u^{\alpha-3} (1 - u) (\log(1 - u))^2 du,$$

$$a_{12}(\alpha, \beta) = a_{21}(\alpha, \beta) = \frac{nk\alpha}{\beta} \int_0^1 u^{\alpha-2} (1-u) \log(1-u) du.$$

Proceeding similarly, we can get

$$I_{\mathbf{W}|\mathbf{X}}^{i:m:n:k}(\boldsymbol{\theta}) = \begin{pmatrix} b_{11}(x_{i:m:n:k}; \alpha, \beta) & b_{12}(x_{i:m:n:k}; \alpha, \beta) \\ b_{21}(x_{i:m:n:k}; \alpha, \beta) & b_{22}(x_{i:m:n:k}; \alpha, \beta) \end{pmatrix},$$

where

$$\begin{aligned} b_{11}(x_{i:m:n:k}; \alpha, \beta) &= \frac{1}{\alpha^2}, \\ b_{12}(x_{i:m:n:k}; \alpha, \beta) &= b_{21}(x_{i:m:n:k}; \alpha, \beta) = h_1(x_{i:m:n:k}; \alpha, \beta) - \frac{x_{i:m:n:k}^\beta \log x_{i:m:n:k}}{1 - x_{i:m:n:k}^\beta}, \\ b_{22}(x_{i:m:n:k}; \alpha, \beta) &= \frac{1}{\beta^2} + (\alpha - 1)h_2(x_{i:m:n:k}; \alpha, \beta) - \frac{\alpha x_{i:m:n:k}^\beta (\log x_{i:m:n:k})^2}{(1 - x_{i:m:n:k}^\beta)^2}, \end{aligned}$$

with

$$\begin{aligned} h_1(c; \alpha, \beta) &= \frac{\alpha}{\beta(1 - F_X(c; \alpha, \beta))} \int_0^{1-c^\beta} (1-u) u^{\alpha-2} \log(1-u) du, \\ h_2(c; \alpha, \beta) &= \frac{\alpha}{\beta^2(1 - F_X(c; \alpha, \beta))} \int_0^{1-c^\beta} (1-u) u^{\alpha-3} (\log(1-u))^2 du. \end{aligned}$$

Furthermore, the asymptotic variance-covariance matrix of $\hat{\boldsymbol{\theta}} = (\hat{\alpha}, \hat{\beta})$, can be obtained by inverting $I(\hat{\boldsymbol{\theta}})$.

3 Bayesian estimation

This section deals with finding Bayes estimates of the parameters α , β , $R(t)$ and $h(t)$. For this aim, we consider different loss functions which are defined as

$$\begin{aligned} \text{squared error loss : } L_1(v, \eta) &= (\eta - v)^2, \\ \text{linex loss : } L_2(v, \eta) &= e^{p(\eta-v)} - p(\eta-v) - 1, \quad p \neq 0, \\ \text{entropy loss : } L_3(v, \eta) &\propto \left(\frac{\eta}{v}\right)^w - w \log\left(\frac{\eta}{v}\right) - 1, \quad w \neq 0. \end{aligned}$$

In each case, η represents an estimate for the unknown parameter v . It is well known that the Bayes estimate of v under the squared error loss function is the posterior mean. Furthermore, the Bayes estimates of v with respect to the loss functions L_2 and L_3 can be obtained from their posterior distributions as

$$\begin{aligned} \text{linex loss : } \eta_{B_2} &= -\frac{1}{p} \log \left\{ E_v \left(e^{-pv} | \mathbf{x} \right) \right\}, \\ \text{entropy loss : } \eta_{B_3} &= (E_v(v^{-w} | \mathbf{x}))^{\frac{-1}{w}}, \end{aligned}$$

provided that required expectations exist. Note that the Bayes estimate η_{B_1} for v under the loss function L_1 is given by η_{B_3} with $w = -1$.

According to Tripathi and Rastogi (2016), the loss function L_1 is naturally symmetric and equally penalizes both over and under estimation. However, in many reliability researches over estimation is more serious than under estimation. Therefore, inference based only on a symmetric loss function may not be practically effective. For the loss function L_2 , underestimation is heavily penalized when p is negative and overestimation is considered to be more serious for positive p . In addition, Tripathi and Rastogi (2016) remarks that in the literature, the asymmetric loss function L_3 first introduced by Calabria and Pulcini (1994). The constant w is known as the shape parameter of this loss function. When $w < 0$, the under estimation is treated to be more serious than the over estimation and the opposite is true for $w > 0$. It is worthwhile to mention that $w = 1$ yields Bayes estimates under a quadratic loss function $(\eta - v)^2/v$. For more details about more applications of the loss function L_3 see Calabria and Pulcini (1996).

To conduct a Bayesian analysis, some prior distributions for the parameters are required. Here, we use independent priors for α and β from gamma $G(a_1, b_1)$ and $G(a_2, b_2)$ distributions, respectively. So, the prior distribution of (α, β) can be written as

$$\pi(\alpha, \beta) \propto \alpha^{a_1-1} e^{-b_1\alpha} \beta^{a_2-1} e^{-b_2\beta}, \quad \alpha, \beta, a_1, b_1, a_2, b_2 > 0. \quad (15)$$

The gamma distribution can accommodate a variety of shapes depending upon parameter values. Thus the family of gamma distributions is highly flexible in nature and can be considered as suitable priors for α and β . It also includes a non-informative prior distribution (Jeffreys prior). One may also refer to Kundu and Pradhan (2009) for a further discussion on this aspect.

After some calculations, the joint posterior distribution of (α, β) is given by

$$\begin{aligned} \pi(\alpha, \beta | \mathbf{x}) &\propto \alpha^{m+a_1-1} \beta^{m+a_2-1} e^{-b_1\alpha} e^{-\beta(b_2-\sum_{i=1}^m \log x_{i:m:n:k})} \\ &\times \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1}. \end{aligned} \quad (16)$$

Thus, the Bayes estimates of α with respect to the loss functions L_2 and L_3 are given, respectively, by

$$\begin{aligned} \tilde{\alpha}_{B_2} &= \frac{-1}{p} \log \left\{ \int_0^\infty \int_0^\infty \alpha^{m+a_1-1} \beta^{m+a_2-1} e^{-\alpha(p+b_1)} e^{-\beta(b_2-\sum_{i=1}^m \log x_{i:m:n:k})} \right. \\ &\quad \times \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1} d\alpha d\beta \Big\}, \\ \tilde{\alpha}_{B_3} &= \left\{ \int_0^\infty \int_0^\infty \alpha^{m+a_1-w-1} \beta^{m+a_2-1} e^{-b_1\alpha} e^{-\beta(b_2-\sum_{i=1}^m \log x_{i:m:n:k})} \right. \\ &\quad \times \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1} d\alpha d\beta \Big\}^{\frac{-1}{w}}. \end{aligned}$$

The Bayes estimate of α with respect to the loss functions L_1 , say $\tilde{\alpha}_{B_1}$, can be obtained from $\tilde{\alpha}_{B_3}$ with $w = -1$. In a similar manner, we can get the Bayes estimates of β with respect to the loss functions L_1 , L_2 and L_3 .

Next, the Bayes estimates of $R(t)$ with respect to the loss functions L_2 and L_3 can be written, respectively, as

$$\begin{aligned}\tilde{R}_{B_2}(t) &= \frac{-1}{p} \log \left\{ \int_0^\infty \int_0^\infty \alpha^{m+a_1-1} \beta^{m+a_2-1} e^{-\alpha b_1} e^{-\beta(b_2 - \sum_{i=1}^m \log x_{i:m:n:k})} \right. \\ &\quad \times \left. \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1} e^{-p(1-t^\beta)^\alpha} d\alpha d\beta \right\}, \\ \tilde{R}_{B_3}(t) &= \left\{ \int_0^\infty \int_0^\infty \alpha^{m+a_1-1} \beta^{m+a_2-1} e^{-b_1\alpha} e^{-\beta(b_2 - \sum_{i=1}^m \log x_{i:m:n:k})} \right. \\ &\quad \times \left. \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1} ((1-t^\beta)^\alpha)^{-w} d\alpha d\beta \right\}^{\frac{-1}{w}}.\end{aligned}$$

Also, the estimate $\tilde{R}_{B_1}(t)$ can be obtained from $\tilde{R}_{B_3}(t)$ with $w = -1$. In a similar manner, we can get the Bayes estimates of $h(t)$ with respect to the loss functions L_1 , L_2 and L_3 .

It is clear that all above Bayes estimators are involved in double integrals which do not have simple closed forms. Therefore, in next sections, we employ two popular approximation procedures to calculate the approximate Bayes estimates of the parameters.

3.1 Lindley's method

Lindley (1980) proposed his procedure to approximate Bayes estimators of the unknown parameters and it is observed that the approximation works quite well. This has been used by several authors to obtain the approximate Bayes estimators. For details see Lindley (1980) or Press and Tanur (2001).

Recall that the Bayes estimates of the previous section are computed as expectation of some parametric functions $g(\alpha, \beta)$ where their expectations are taken with respect to the posterior distribution $\pi(\alpha, \beta | \mathbf{x})$. Consider the ratio of two integrals as

$$I(X) = \frac{\int_{\theta_1, \theta_2} g(\theta_1, \theta_2) e^{l(\theta_1, \theta_2) + \rho(\theta_1, \theta_2)} d(\theta_1, \theta_2)}{\int_{\theta_1, \theta_2} e^{l(\theta_1, \theta_2) + \rho(\theta_1, \theta_2)} d(\theta_1, \theta_2)}, \quad (17)$$

where $g(\theta_1, \theta_2)$ is a function of θ_1 and θ_2 only, $l(\theta_1, \theta_2)$, is a log-likelihood and $\rho(\theta_1, \theta_2) = \log \pi(\theta_1, \theta_2)$. According to Lindley (1980), $I(X)$ can be approximated by

$$\tilde{I} = g(\theta) + 0.5[U + l_{30}^* V_{12} + l_{03}^* V_{21} + l_{21}^* C_{12} + l_{12}^* C_{21}] + \rho_1 U_{12} + \rho_2 U_{21}, \quad (18)$$

where

$$U = \sum_{i=1}^2 \sum_{j=1}^2 \omega_{ij} \tau_{ij}, \quad l_{ij}^* = \frac{\partial^{i+j} l}{\partial \theta_1^i \partial \theta_2^j}, \quad i, j = 0, 1, 2, 3, \quad i+j = 3,$$

$$\rho_i = \frac{\partial \rho}{\partial \theta_i}, \quad \omega_i = \frac{\partial g}{\partial \theta_i}, \quad \omega_{ij} = \frac{\partial^2 g}{\partial \theta_i \partial \theta_j}, \quad \rho = \log \pi(\theta_1, \theta_2), \quad U_{ij} = \omega_i \tau_{ii} + \omega_j \tau_{ji}$$

$$V_{ij} = \tau_{ii}(\omega_i \tau_{ii} + \omega_j \tau_{ij}), \quad C_{ij} = 3 \omega_i \tau_{ii} \tau_{ij} + \omega_j (\tau_{ii} \tau_{jj} + 2 \tau_{ij}^2),$$

where $\tau_{i,j}$ denotes the (i, j) th elements of the matrix $[-\frac{\partial^2 l}{\partial \theta_i \partial \theta_j}]^{-1}$. We apply the above method for our estimation problem by assuming $\theta = (\alpha, \beta)$ instead of (θ_1, θ_2) .

In our case, the required quantities of the expression (18) are derived as

$$l_{12}^* = -k \sum_{i=1}^m (r_i + 1) x_{i:m:n:k}^{\hat{\beta}} \log(x_{i:m:n:k}) (1 - x_{i:m:n:k}^{\hat{\beta}})^{-1}, \quad l_{21}^* = 0, \quad l_{30}^* = \frac{2m}{\hat{\beta}^3},$$

$$l_{03}^* = \frac{2m}{\hat{\beta}^3} - \sum_{i=1}^m (\hat{\alpha} k (r_i + 1) - 1) x_{i:m:n:k}^{\hat{\beta}} (\log(x_{i:m:n:k}))^3$$

$$\times (1 + x_{i:m:n:k}^{\hat{\beta}}) (1 - x_{i:m:n:k}^{\hat{\beta}})^{-3},$$

$$\tau_{11} = \frac{P}{N}, \quad \tau_{22} = \frac{M}{N}, \quad \tau_{12} = \tau_{21} = -\frac{T}{N}, \quad M = \frac{m}{\hat{\alpha}^2},$$

$$P = \frac{m}{\hat{\beta}^2} + \sum_{i=1}^m (\hat{\alpha} k (r_i + 1) - 1) x_{i:m:n:k}^{\hat{\beta}} (\log(x_{i:m:n:k}))^2 (1 - x_{i:m:n:k}^{\hat{\beta}})^{-2}$$

$$T = k \sum_{i=1}^m (r_i + 1) x_{i:m:n:k}^{\hat{\beta}} \log(x_{i:m:n:k}) (1 - x_{i:m:n:k}^{\hat{\beta}})^{-1}, \quad N = MP - T^2,$$

$$\rho_1 = \frac{(a_1 - 1)}{\hat{\alpha}} - b_1, \quad \rho_2 = \frac{(a_2 - 1)}{\hat{\beta}} - b_2.$$

By substituting the MLEs $\hat{\alpha}$ and $\hat{\beta}$ of the parameters α and β into expressions of (18), the approximate Bayes estimates can be obtained.

In the case that we estimate α under the loss function L_2 , it can be seen that

$$g(\alpha, \beta) = e^{-p\alpha}, \quad \omega_1 = -p e^{-p\alpha}, \quad \omega_{11} = p^2 e^{-p\alpha}, \quad \omega_2 = \omega_{22} = \omega_{12} = \omega_{21} = 0.$$

By utilizing these expressions in (18), the Bayes estimate of α can be written as

$$\tilde{\alpha}_{B_2} = \frac{-1}{p} \log \left\{ e^{-p\hat{\alpha}} + \frac{1}{N} (P\rho_1 - T\rho_2) + \frac{0.5}{N^2} (P^2 l_{30}^* - T M l_{03}^* + (MP + 2T^2) l_{12}^*) \right\}.$$

To compute the Bayes estimate of α under the loss function L_3 , suppose $g(\alpha, \beta) = \alpha^{-w}$. Then,

$$\omega_1 = -w \alpha^{-w-1}, \quad \omega_{11} = w(w+1) \alpha^{-w-2}, \quad \omega_2 = \omega_{22} = \omega_{12} = \omega_{21} = 0.$$

Also, we have

$$\tilde{\alpha}_{B_3} = \left\{ \hat{\alpha}^{-w} + \frac{1}{N} (M\rho_2 - T\rho_1) + \frac{0.5}{N^2} (M^2 l_{03}^* - T P l_{30}^* - 3 T M l_{12}^*) \right\}^{-w}.$$

Similarly, we can derive the Bayes estimates of β , $R(t)$ and $h(t)$ with respect to the loss functions L_1 , L_2 and L_3 .

3.2 Gibbs sampling

In this section, we adopt Gibbs sampling method to extract random samples from the conditional densities of the parameters and use them for calculating the Bayes estimates. For more details about Gibbs sampler, see Geman and Geman (1984), Gilks et al. (1995), Gilks and Wild (1992) and Smith and Roberts (1993). From (16), the conditional posterior densities of α and β can be extracted, respectively, as

$$\pi_1^*(\alpha \mid \beta, \mathbf{x}) \propto G(m + a_1, b_1) \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1} \quad (19)$$

$$\pi_2^*(\beta \mid \alpha, \mathbf{x}) \propto G(m + a_2, b_2 - \sum_{i=1}^m \log x_{i:m:n:k}) \prod_{i=1}^m (1 - x_{i:m:n:k}^\beta)^{\alpha k(1+r_i)-1}. \quad (20)$$

Note that the conditional densities in (19) and (20) are not in the form of well known distributions. By trial and drawing plots, we can conclude that they are similar to the normal distribution. Therefore, in the following algorithm, we employ the well-known Metropolis-Hastings (M-H) with the normal proposal distribution to generate samples from these distributions.

- i. Let initial values of the parameters be $(\alpha^{(0)}, \beta^{(0)})$ and set $l = 1$.
- ii. Considering the proposal distribution $q(\alpha) \equiv I(\alpha > 0)N(\alpha^{(l-1)}, 1)$ for the M-H method, generate $\alpha^{(l)}$ from $\pi_1^*(\alpha \mid \beta^{(l-1)}, \mathbf{x})$.
- iii. Generate $\beta^{(l)}$ from $\pi_2^*(\beta \mid \alpha^{(l)}, \mathbf{x})$ using M-H method with the proposal distribution $q(\beta) \equiv I(\beta > 0)N(\beta^{(l-1)}, 1)$.
- iv. Compute $R^{(l)}(t)$ and $h^{(l)}(t)$ from (3) and (4) and set $l = l + 1$.
- v. Repeat Steps ii-iv M times and obtain $\alpha^{(l)}$, $\beta^{(l)}$, $R^{(l)}(t)$ and $h^{(l)}(t)$ for $l = 1, \dots, M$.

Using the generated random samples from the above Gibbs technique, the Bayes estimates of the parameters α , β , $R(t)$ and $h(t)$ under the squared error loss function become $\tilde{\alpha}_{B1} = \sum_{l=1}^M \alpha^{(l)}/M$, $\tilde{\beta}_{B1} = \sum_{l=1}^M \beta^{(l)}/M$, $\tilde{R}_{B1}(t) = \sum_{l=1}^M R^{(l)}(t)/M$ and $\tilde{h}_{B1}(t) = \sum_{l=1}^M h^{(l)}(t)/M$, respectively.

Next, under the loss function L_2 , the Bayes estimates of the parameters are given by

$$\tilde{\alpha}_{B2} = -\frac{1}{p} \log \left(\frac{1}{M} \sum_{l=1}^M e^{-p\alpha^{(l)}} \right), \quad \tilde{\beta}_{B2} = -\frac{1}{p} \log \left(\frac{1}{M} \sum_{l=1}^M e^{-p\beta^{(l)}} \right), \quad (21)$$

$$\tilde{R}_{B2}(t) = -\frac{1}{p} \log \left(\frac{1}{M} \sum_{l=1}^M e^{-pR^{(l)}(t)} \right), \quad \tilde{h}_{B2}(t) = -\frac{1}{p} \log \left(\frac{1}{M} \sum_{l=1}^M e^{-ph^{(l)}(t)} \right). \quad (22)$$

Finally, Bayes estimates of the parameters under the loss function L_3 can be computed as

$$\tilde{\alpha}_{B3} = \left(\frac{1}{M} \sum_{l=1}^M \left(\frac{1}{\alpha^{(l)}} \right)^w \right)^{-1/w}, \quad \tilde{\beta}_{B3} = \left(\frac{1}{M} \sum_{l=1}^M \left(\frac{1}{\beta^{(l)}} \right)^w \right)^{-1/w}, \quad (23)$$

$$\tilde{R}_{B_3}(t) = \left(\frac{1}{M} \sum_{l=1}^M \left(\frac{1}{R^{(l)}(t)} \right)^w \right)^{-1/w}, \quad \tilde{h}_{B_3}(t) = \left(\frac{1}{M} \sum_{l=1}^M \left(\frac{1}{h^{(l)}(t)} \right)^w \right)^{-1/w}. \quad (24)$$

Therefore, we can obtain the Bayesian credible and HPD intervals for both α and β using samples generated from the above algorithm. It should be noticed that the $100(1 - \xi)\%$ HPD interval for the unknown parameter $\theta = (\alpha, \beta)$ can be easily constructed using the MH samples. The idea developed by Chen and shao (1999). First, arrange the sample $\theta_1, \theta_2, \dots, \theta_Q$ in increasing order. Then, the $100(1 - \tau)\%$ credible interval for θ is obtained as $(\theta_1, \theta_{\lfloor (1-\tau)Q+1 \rfloor}) \dots (\theta_{\lfloor Q\tau \rfloor}, \theta_Q)$ in which $\lfloor z \rfloor$ denotes the greatest integer less than or equal to z . Among all such credible intervals, the shortest one is the HPD interval.

4 Numerical comparisons

In this section, we analyze the performance of different estimators which discussed in the previous sections by a Monte Carlo simulation study. Samples are generated from the Kumaraswamy distribution under progressive first failure censoring. The true value of parameters can be taken an arbitrarily choice such as $(\alpha, \beta) = (1.5, 0.5)$. We carry out the simulation procedures for different combinations of (n, m, k, r) and consider independent gamma priors under loss functions L_1, L_2 and L_3 for Bayes estimates. The Hyperparameters are given by $(a_1, b_1, a_2, b_2) = (3, 2, 1, 2)$.

Here, it is noticeable that some censoring schemes with repetitive elements have been represented by sort notations such as $(0^{*2}) = (0, 0)$ and $(2^{*3}) = (2, 2, 2)$. In each case, we obtain MLEs and Bayes estimators of unknown parameters. The Lindley and M-H algorithm methods are used for deriving Bayes estimates. Also, the 95% asymptotic confidence interval using MLEs and 95% HPD credible interval for unknown parameters can be derived at sufficient nominal level. All cases in the simulation study have been done 5000 times using R software and basically, we use *nleqslv* package for solving non linear equations. The conclusions are given for two different group sizes $k = 2, 4$ and two different number of groups $n = 30, 50$. The results of the Monte Carlo simulation are tabulated and reported below and some important conclusions have been drawn.

In Tables 1 to 4, the average value and MSE of unknown parameters of α and β are reported using the MLE (the EM algorithm), the Lindley and M-H algorithm methods. Whenever the sample size n increases, the MSE of estimates decreases. Moreover, as m increases, the bias and MSE decrease. Also, when the group size k increases, the MSE of α increases while the MSE of β decreases. It is observed that among all Bayes estimate methods, the M-H algorithm performs better and has quite good performance in comparison to the MLE method. It is also observed that as the failure proportion m/n increases, the point estimates even become better. In addition, in most of cases, when the group size k increases, the information reduces as well as the bias and MSE increase.

In Tables 5 to 8, the average value as well as the MSE for all estimates of the reliability function $R(t)$ are provided for $t = 0.1$ and $t = 0.5$, respectively. It may be noted that the MLE competes quite well with different Bayes estimates. However, the Bayes estimates have superior performance in comparison to MLEs. The selection

of $p = 1$ under the loss function L_2 creates the best average estimates of $R(t)$ when $t = 0.1$. Table 4 illustrates that the selection of $w = 1$ under the loss function L_3 provides almost good conditions for estimation of $R(t)$ when $t = 0.5$. These results hold good under various censoring schemes.

The average value and MSE for all estimates of $h(t)$ are reported in Tables 9 and 10. Based on the estimation of $h(t)$ with $t = 0.25$, it can be seen that the Bayes estimates under the loss function L_2 with $p = -0.5$ and $p = 0.5$ have almost well performance in comparison to the other estimates. In this case, the Bayes estimation of $h(t)$ with $t = 0.25$ under the loss function L_2 , gain the best performance under almost all the censoring schemes. Also, it is easy to see that the MLEs gain well competence in comparison to the corresponding Bayes estimates for both cases.

We can conclude the similar behavior for various other censoring schemes as well. Also, by the effective increase in sample sizes, the absolute bias and MSE of all proposed estimates tend to decrease.

In Tables 11 and 12, the average length (AL) of asymptotic confidence interval (ACI) and HPD are presented along with their coverage probability (CP). It turns out that the length of ACI as well as HPD interval become narrower as the sample size n increases. The HPD interval in terms of the AL, performs better than the ACI. Also, as the group size k increases, the AL of ACI increases for the parameter α and decreases for the β under both ACI and HPD intervals. In some cases, the CP of unknown parameters attain their prescribed nominal confidence level.

5 Data Analysis

In this section, for illustrative purposes, we have analyzed two real data sets which have been recently considered by Sharaf El-Deen et al. (2014). They fitted these real data sets to the Kumaraswamy distribution and found that the Kumaraswamy distribution fits both real data sets reasonably good. Also, they obtained useful inferences for the prescribed model.

Data set 1. The first application is the annual water level behind the High Dam during the flood time from 1968 to 2010 obtained from Abdel-Latif and Yacoub (2011). The highest water level of the Dam is 182 meter (m) above the mean sea level. The data points are listed below as follows.

0.83021, 0.84521, 0.87736, 0.89280, 0.86923, 0.88851, 0.90989, 0.94741, 0.94340, 0.94791, 0.95076, 0.94104, 0.94027, 0.93604, 0.91137, 0.89890, 0.85917, 0.86390, 0.84972, 0.83351, 0.90335, 0.89983, 0.89285, 0.90098, 0.92005, 0.93209, 0.94692, 0.94923, 0.96417, 0.96016, 0.96016, 0.96587, 0.96620, 0.96538, 0.96230, 0.94538, 0.93181, 0.92664, 0.95285, 0.96044, 0.95219, 0.93291.

The corresponding Kolmogorov-Smirnov p-value based on the above data is 0.6728. For more details on goodness of fit discussions, one may refer to Sharaf El-Deen et al. (2014).

Table 1: Average and estimated MSE (in the second line) values of all estimates of α for different choices of n, m, k, r .

(n, m)	r	k	$\hat{\alpha}$		p			w		
					-0.5	0.5	1.0	-0.5	0.5	1.0
(30,20)	(10, 0*19)	2	1.785	Lindley	$\tilde{\alpha}_{B_1}$	$\tilde{\alpha}_{B_2}$		$\tilde{\alpha}_{B_3}$		
					1.436	1.684	1.430	1.422	1.418	1.397
		0.319	Lindley		0.089	0.048	0.044	0.073	0.055	0.062
				M-H	1.491	1.510	1.472	1.453	1.478	1.451
					0.168	0.173	0.164	0.162	0.169	0.170
									0.172	0.172
	(0*19, 10)	4	1.911	Lindley	1.443	1.855	1.352	1.381	1.323	1.340
					0.152	0.050	0.049	0.067	0.060	0.067
		0.343	Lindley		1.529	1.548	1.509	1.490	1.515	1.488
				M-H	0.183	0.188	0.178	0.175	0.183	0.184
									0.186	0.186
(30,25)	(5, 0*24)	2	2.009	Lindley	1.482	1.964	1.277	1.345	1.239	1.298
					0.188	0.067	0.056	0.065	0.124	0.147
		0.420	Lindley		1.534	1.555	1.513	1.492	1.520	1.491
				M-H	0.212	0.220	0.196	0.201	0.212	0.212
									0.213	0.213
	(0*5, 1*10, 0*5)	4	2.101	Lindley	1.604	1.200	1.314	1.263	1.079	1.190
					0.237	0.085	0.073	0.084	0.082	0.087
		0.465	Lindley		1.491	1.511	1.470	1.450	1.476	1.446
				M-H	0.224	0.231	0.219	0.215	0.225	0.229
									0.231	0.231
(30,25)	(5, 0*24)	2	1.768	Lindley	1.444	1.767	1.379	1.398	1.345	1.357
					0.121	0.056	0.042	0.050	0.050	0.057
		0.305	Lindley		1.527	1.548	1.507	1.488	1.514	1.486
				M-H	0.184	0.190	0.179	0.175	0.183	0.184
									0.186	0.186
	(0*24, 5)	4	1.897	Lindley	1.441	1.874	1.263	1.317	1.212	1.255
					0.175	0.093	0.090	0.092	0.093	0.106
		0.265	Lindley		1.491	1.510	1.472	1.453	1.477	1.450
				M-H	0.196	0.200	0.194	0.192	0.198	0.202
									0.205	0.205
(30,25)	(5, 0*24)	2	1.768	Lindley	1.495	1.662	1.483	1.464	1.483	1.451
					0.070	0.086	0.078	0.082	0.067	0.071
		0.221	Lindley		1.511	1.529	1.494	1.476	1.499	1.475
				M-H	0.152	0.157	0.148	0.145	0.151	0.152
									0.153	0.153
	(0*24, 5)	4	1.830	Lindley	1.447	1.732	1.419	1.413	1.425	1.404
					0.110	0.074	0.042	0.068	0.063	0.077
		0.277	Lindley		1.478	1.497	1.459	1.440	1.465	1.438
				M-H	0.160	0.163	0.157	0.155	0.160	0.163
									0.165	0.165
(30,25)	(5, 0*24)	2	1.791	Lindley	1.443	1.694	1.427	1.416	1.432	1.409
					0.100	0.068	0.042	0.067	0.090	0.092
		0.206	Lindley		1.517	1.536	1.499	1.481	1.505	1.479
				M-H	0.177	0.182	0.173	0.170	0.177	0.178
									0.179	0.179
	(0*10, 1*5, 0*10)	4	1.903	Lindley	1.487	1.773	1.342	1.369	1.323	1.337
					0.138	0.065	0.054	0.063	0.063	0.066
		0.289	Lindley		1.487	1.506	1.467	1.448	1.473	1.445
				M-H	0.181	0.184	0.178	0.176	0.182	0.186
									0.189	0.189
(30,25)	(5, 0*24)	2	1.773	Lindley	1.469	1.668	1.461	1.444	1.458	1.428
					0.074	0.061	0.052	0.058	0.053	0.066
		0.240	Lindley		1.514	1.533	1.495	1.477	1.501	1.476
				M-H	0.162	0.167	0.158	0.156	0.162	0.163
									0.164	0.164
	(0*10, 1*5, 0*10)	4	1.825	Lindley	1.434	1.768	1.384	1.388	1.376	1.367
					0.133	0.056	0.044	0.049	0.092	0.098
		0.232	Lindley		1.532	1.551	1.513	1.494	1.519	1.492
				M-H	0.178	0.184	0.174	0.171	0.178	0.179
									0.180	0.180

Table 2: Average and estimated MSE (in the second line) values of all estimates of α for different choices of n, m, k, r .

						p			w		
						-0.5	0.5	1.0	-0.5	0.5	1.0
(n, m)	r	k	$\hat{\alpha}$		$\tilde{\alpha}_{B_1}$	$\tilde{\alpha}_{B_2}$			$\tilde{\alpha}_{B_3}$		
(30,30)	(0*30)	2	1.710 0.186	Lindley M-H	1.539	1.630	1.503	1.473	1.515	1.471	1.453
					0.073	0.086	0.068	0.079	0.069	0.077	0.084
					1.541	1.558	1.523	1.506	1.529	1.506	1.494
					0.159	0.165	0.154	0.150	0.156	0.157	0.157
		4	1.791 0.279	Lindley M-H	1.466	1.695	1.459	1.443	1.465	1.432	1.421
					0.085	0.079	0.067	0.072	0.072	0.079	0.095
					1.509	1.526	1.492	1.475	1.497	1.473	1.461
					0.161	0.166	0.157	0.155	0.161	0.162	0.163
(50,30)	(20,0*29)	2	1.707 0.266	Lindley M-H	1.556	1.636	1.520	1.489	1.532	1.488	1.470
					0.078	0.095	0.074	0.083	0.075	0.082	0.088
					1.528	1.546	1.511	1.494	1.517	1.494	1.483
					0.151	0.156	0.146	0.143	0.147	0.149	0.149
		4	1.779 0.241	Lindley M-H	1.503	1.687	1.485	1.463	1.494	1.457	1.440
					0.080	0.086	0.052	0.074	0.073	0.082	0.094
					1.531	1.550	1.513	1.495	1.519	1.494	1.482
					0.168	0.173	0.163	0.159	0.166	0.167	0.168
(0*29,20)	(0*29,20)	2	1.842 0.293	Lindley M-H	1.476	1.761	1.421	1.416	1.429	1.410	1.403
					0.103	0.073	0.059	0.062	0.061	0.065	0.069
					1.509	1.528	1.489	1.470	1.495	1.469	1.455
					0.163	0.169	0.159	0.156	0.163	0.164	0.165
		4	1.929 0.260	Lindley M-H	1.521	1.949	1.322	1.345	1.323	1.342	1.347
					0.161	0.099	0.090	0.093	0.098	0.116	0.131
					1.491	1.510	1.473	1.455	1.478	1.452	1.439
					0.189	0.193	0.186	0.183	0.190	0.193	0.195
(50,40)	(10,0*39)	2	1.6434 0.246	Lindley M-H	1.561	1.602	1.524	1.493	1.538	1.496	1.478
					0.090	0.091	0.080	0.082	0.086	0.086	0.088
					1.535	1.551	1.519	1.504	1.525	1.504	1.494
					0.124	0.128	0.120	0.117	0.113	0.122	0.121
		4	1.683 0.254	Lindley M-H	1.547	1.631	1.506	1.473	1.523	1.475	1.455
					0.079	0.094	0.073	0.081	0.084	0.089	0.096
					1.516	1.533	1.499	1.482	1.504	1.481	1.470
					0.124	0.128	0.100	0.118	0.123	0.123	0.124
(0*39,10)	(0*39,10)	2	1.692 0.264	Lindley M-H	1.546	1.640	1.511	1.481	1.523	1.479	1.460
					0.074	0.092	0.071	0.081	0.071	0.078	0.085
					1.536	1.552	1.519	1.503	1.525	1.503	1.492
					0.137	0.142	0.133	0.129	0.136	0.135	0.135
		4	1.786 0.246	Lindley M-H	1.500	1.716	1.486	1.464	1.489	1.451	1.435
					0.084	0.060	0.046	0.052	0.053	0.057	0.062
					1.514	1.531	1.497	1.481	1.503	1.480	1.468
					0.150	0.154	0.146	0.143	0.149	0.149	0.150
(50,50)	(0*50)	2	1.645 0.206	Lindley M-H	1.587	1.625	1.552	1.523	1.566	1.528	1.511
					0.092	0.092	0.080	0.084	0.076	0.080	0.080
					1.539	1.554	1.525	1.511	1.530	1.511	1.501
					0.118	0.122	0.114	0.111	0.117	0.115	0.115
		4	1.624 0.183	Lindley M-H	1.549	1.595	1.509	1.475	1.524	1.478	1.458
					0.082	0.079	0.075	0.076	0.079	0.079	0.082
					1.547	1.562	1.532	1.517	1.537	1.517	1.507
					0.135	0.139	0.131	0.129	0.130	0.133	0.133

Table 3: Average and estimated MSE (in the second line) values of all estimates of β for different choices of n, m, k, r .

						p			w		
						-0.5	0.5	1.0	-0.5	0.5	1.0
(n, m)	r	k	$\hat{\beta}$		$\tilde{\beta}_{B_1}$	$\tilde{\beta}_{B_2}$			$\tilde{\beta}_{B_3}$		
(30,20)	(10,0 ^{*19})	2	0.533954 0.014	Lindley	0.494	0.495	0.492	0.490	0.490	0.484	0.480
				M-H	0.005	0.005	0.005	0.005	0.005	0.005	0.006
					0.495	0.497	0.494	0.492	0.492	0.486	0.482
					0.007	0.007	0.007	0.007	0.007	0.008	0.008
		4	0.534 0.012	Lindley	0.484	0.485	0.484	0.482	0.483	0.481	0.480
				M-H	0.013	0.014	0.010	0.012	0.009	0.011	0.013
					0.501	0.502	0.500	0.499	0.499	0.495	0.493
					0.005	0.005	0.005	0.005	0.005	0.005	0.005
	(0 ^{*19} ,10)	2	0.551 0.018	Lindley	0.469	0.466	0.472	0.472	0.473	0.471	0.473
				M-H	0.017	0.024	0.020	0.021	0.015	0.017	0.018
					0.500	0.501	0.498	0.497	0.497	0.491	0.488
					0.007	0.007	0.007	0.007	0.007	0.007	0.008
		4	0.546 0.025	Lindley	0.444	0.436	0.449	0.453	0.452	0.464	0.464
				M-H	0.023	0.024	0.017	0.018	0.016	0.018	0.021
					0.489	0.490	0.488	0.487	0.487	0.483	0.481
					0.005	0.005	0.005	0.005	0.005	0.005	0.005
	(0 ^{*5} ,1 ^{*10} ,0 ^{*5})	2	0.538 0.014	Lindley	0.482	0.482	0.482	0.479	0.481	0.476	0.475
				M-H	0.007	0.008	0.008	0.008	0.008	0.008	0.009
					0.500	0.502	0.499	0.498	0.497	0.492	0.489
					0.007	0.007	0.007	0.007	0.007	0.007	0.007
		4	0.526 0.010	Lindley	0.466	0.465	0.467	0.468	0.467	0.466	0.466
				M-H	0.007	0.007	0.006	0.006	0.007	0.007	0.007
					0.491	0.492	0.490	0.489	0.489	0.485	0.484
					0.005	0.005	0.005	0.005	0.005	0.005	0.005
(30,25)	(5,0 ^{*24})	2	0.535 0.013	Lindley	0.502	0.504	0.500	0.498	0.498	0.492	0.489
				M-H	0.005	0.005	0.005	0.005	0.005	0.005	0.005
					0.500	0.501	0.499	0.497	0.497	0.491	0.489
					0.008	0.008	0.008	0.007	0.008	0.008	0.008
		4	0.529 0.020	Lindley	0.495	0.496	0.495	0.490	0.493	0.488	0.487
				M-H	0.011	0.01	0.010	0.012	0.010	0.011	0.014
					0.495	0.495	0.494	0.493	0.493	0.489	0.487
					0.005	0.005	0.005	0.005	0.005	0.005	0.005
	(0 ^{*24} ,5)	2	0.531941 0.023	Lindley	0.492	0.493	0.492	0.489	0.490	0.484	0.481
				M-H	0.009	0.012	0.008	0.009	0.011	0.012	0.011
					0.497	0.498	0.496	0.495	0.494	0.489	0.486
					0.007	0.007	0.007	0.007	0.007	0.008	0.008
		4	0.535 0.021	Lindley	0.485	0.484	0.487	0.482	0.486	0.482	0.483
				M-H	0.012	0.010	0.018	0.019	0.010	0.012	0.014
					0.491	0.492	0.490	0.490	0.490	0.486	0.484
					0.005	0.005	0.005	0.005	0.005	0.005	0.005
	(0 ^{*10} ,1 ^{*5} ,0 ^{*10})	2	0.531 0.011	Lindley	0.495	0.497	0.493	0.492	0.492	0.486	0.483
				M-H	0.004	0.004	0.004	0.004	0.004	0.004	0.005
					0.501	0.502	0.500	0.498	0.498	0.493	0.490
					0.007	0.007	0.007	0.007	0.007	0.007	0.007
		0.527 0.009	Lindley	0.488	0.488	0.488	0.487	0.487	0.487	0.484	0.484
			M-H	0.007	0.006	0.006	0.006	0.006	0.007	0.011	0.014
					0.502	0.503	0.501	0.500	0.500	0.497	0.495
					0.004	0.004	0.004	0.004	0.004	0.004	0.004

Table 7: Average and MSE values of all estimates of $R(t)$ for different choices of n, m, k, r and $t = 0.5$.

						p			w		
						-0.5	0.5	1.0	-0.5	0.5	1.0
(n, m)	r	k	$\hat{R}(t)$		$\tilde{R}_{B_1}(t)$	$\tilde{R}_{B_2}(t)$			$\tilde{R}_{B_3}(t)$		
(30,20)	(10, 0* ¹⁹)	2	0.149 0.004	Lindley	0.179	0.180	0.178	0.177	0.174	0.162	0.154
					0.002	0.002	0.002	0.002	0.002	0.002	
				M-H	0.177	0.177	0.176	0.176	0.175	0.173	
					0.004	0.004	0.004	0.004	0.004	0.004	
		4	0.145 0.005	Lindley	0.186	0.188	0.185	0.184	0.181	0.183	
					0.002	0.003	0.002	0.002	0.002	0.002	
				M-H	0.179	0.180	0.179	0.178	0.177	0.174	
					0.004	0.004	0.004	0.004	0.004	0.004	
	(0* ¹⁹ , 10)	2	0.138 0.005	Lindley	0.191	0.192	0.191	0.190	0.189	0.220	
					0.002	0.002	0.002	0.002	0.002	0.022	
				M-H	0.176	0.177	0.176	0.176	0.175	0.172	
					0.004	0.004	0.004	0.004	0.004	0.003	
		4	0.142 0.007	Lindley	0.207	0.207	0.206	0.206	0.210	0.195	
					0.003	0.003	0.003	0.003	0.003	0.003	
				M-H	0.184	0.185	0.184	0.183	0.182	0.177	
					0.006	0.006	0.006	0.006	0.006	0.005	
	(0 ⁵ , 1 ¹⁰ , 0 ⁵)	2	0.145 0.005	Lindley	0.184	0.185	0.183	0.182	0.179	0.172	
					0.002	0.003	0.002	0.002	0.002	0.002	
				M-H	0.175	0.175	0.175	0.174	0.174	0.171	
					0.004	0.004	0.004	0.004	0.004	0.004	
		4	0.141 0.006	Lindley	0.194	0.196	0.194	0.193	0.195	0.218	
					0.002	0.003	0.002	0.002	0.002	0.002	
				M-H	0.177	0.177	0.176	0.176	0.174	0.171	
					0.005	0.005	0.005	0.005	0.005	0.004	
(30,25)	(5, 0* ²⁴)	2	0.149 0.004	Lindley	0.174	0.174	0.173	0.172	0.169	0.157	
					0.002	0.002	0.002	0.002	0.002	0.002	
				M-H	0.173	0.173	0.173	0.173	0.172	0.170	
					0.003	0.003	0.003	0.003	0.003	0.003	
		4	0.144 0.004	Lindley	0.178	0.180	0.177	0.176	0.173	0.162	
					0.002	0.002	0.001	0.001	0.001	0.001	
				M-H	0.174	0.175	0.174	0.174	0.173	0.169	
					0.004	0.004	0.004	0.004	0.004	0.003	
	(0* ²⁴ , 5)	2	0.146 0.004	Lindley	0.176	0.177	0.175	0.175	0.171	0.160	
					0.002	0.002	0.001	0.001	0.001	0.001	
				M-H	0.175	0.175	0.175	0.175	0.174	0.171	
					0.003	0.003	0.003	0.003	0.003	0.003	
		4	0.140 0.005	Lindley	0.184	0.186	0.184	0.183	0.180	0.191	
					0.002	0.002	0.001	0.001	0.001	0.002	
				M-H	0.180	0.180	0.179	0.179	0.178	0.174	
					0.005	0.005	0.004	0.004	0.004	0.004	
	(0 ¹⁰ , 1 ⁵ , 0 ¹⁰)	2	0.146 0.004	Lindley	0.174	0.175	0.173	0.172	0.169	0.157	
					0.002	0.002	0.002	0.002	0.001	0.001	
				M-H	0.172	0.172	0.172	0.172	0.171	0.168	
					0.003	0.003	0.003	0.003	0.003	0.003	
		4	0.145 0.005	Lindley	0.182	0.183	0.181	0.181	0.177	0.170	
					0.002	0.002	0.002	0.002	0.001	0.001	
				M-H	0.178	0.178	0.177	0.177	0.176	0.172	
					0.004	0.004	0.004	0.004	0.004	0.004	

Table 9: Average and MSE values of all estimates of $h(t)$ for different choices of n, m, k, r and $t = 0.25$.

						p			w		
						-0.5	0.5	1.0	-0.5	0.5	1.0
(n, m)	r	k	$\tilde{R}(t)$		$\tilde{R}_{B_1}(t)$	$\tilde{R}_{B_2}(t)$			$\tilde{R}_{B_3}(t)$		
(30,20)	(10, 0*19)	2	3.775 0.727	Lindley	3.366	3.470	3.260	3.180	3.332	3.270	3.243
					0.224	0.235	0.273	0.342	0.238	0.277	0.302
				M-H	2.933	2.970	2.896	2.858	2.920	2.892	2.878
					0.477	0.480	0.477	0.481	0.481	0.490	0.496
		4	3.890 0.603	Lindley	3.278	3.361	3.183	3.115	3.250	3.190	3.168
					0.220	0.222	0.281	0.313	0.210	0.261	0.283
				M-H	2.998	3.041	2.954	2.910	2.982	2.950	2.933
					0.520	0.527	0.519	0.522	0.524	0.535	0.542
(0*19, 10)	2	3.941 0.694	Lindley	3.167	3.338	3.115	3.078	3.163	3.132	3.125	
				0.278	0.275	0.322	0.462	0.228	0.252	0.257	
			M-H	2.981	3.022	2.939	2.897	2.966	2.935	2.919	
				0.565	0.572	0.564	0.568	0.570	0.580	0.587	
	4	4.042 0.810	Lindley	2.814	3.510	2.935	2.961	2.891	2.926	2.930	
				0.395	0.302	0.418	0.442	0.322	0.383	0.466	
			M-H	2.921	2.971	2.871	2.822	2.903	2.864	2.844	
				0.618	0.619	0.625	0.640	0.627	0.650	0.663	
(0*5, 1*10, 0*5)	2	3.895 0.719	Lindley	3.291	3.314	3.224	3.157	3.269	3.223	3.203	
				0.293	0.168	0.256	0.343	0.203	0.250	0.283	
			M-H	2.950	2.990	2.910	2.871	2.936	2.905	2.890	
				0.518	0.519	0.522	0.529	0.524	0.537	0.545	
	4	4.038 0.814	Lindley	3.081	3.433	3.102	3.100	3.094	3.087	3.080	
				0.385	0.213	0.291	0.350	0.181	0.265	0.297	
			M-H	2.885	2.932	2.837	2.790	2.867	2.831	2.812	
				0.570	0.569	0.577	0.590	0.578	0.599	0.611	
(30,25)	(5, 0*24)	2	3.746 0.682	Lindley	3.421	3.519	3.313	3.232	3.387	3.328	3.300
					0.227	0.228	0.249	0.294	0.234	0.258	0.273
				M-H	2.979	3.014	2.944	2.908	2.967	2.941	2.928
					0.420	0.425	0.418	0.420	0.422	0.428	0.432
	4	3.838 0.641	Lindley	3.378	3.489	3.246	3.165	3.342	3.270	3.246	
				0.274	0.274	0.270	0.375	0.181	0.233	0.251	
			M-H	2.978	3.022	2.933	2.889	2.962	2.930	2.914	
				0.444	0.454	0.439	0.442	0.446	0.453	0.458	
(0*24, 5)	2	3.717 0.664	Lindley	3.351	3.446	3.248	3.176	3.318	3.259	3.232	
				0.221	0.233	0.263	0.340	0.232	0.271	0.294	
			M-H	2.979	3.016	2.941	2.904	2.965	2.938	2.924	
				0.467	0.472	0.466	0.469	0.470	0.478	0.482	
	4	3.917 0.594	Lindley	3.251	3.366	3.173	3.125	3.231	3.176	3.165	
				0.295	0.226	0.253	0.301	0.169	0.227	0.231	
			M-H	2.958	3.005	2.911	2.865	2.941	2.907	2.890	
				0.486	0.496	0.482	0.485	0.489	0.499	0.506	
(0*10, 1*5, 0*10)	2	3.805 0.624	Lindley	3.442	3.550	3.329	3.248	3.407	3.342	3.317	
				0.216	0.233	0.237	0.285	0.222	0.243	0.261	
			M-H	3.005	3.043	2.967	2.929	2.992	2.965	2.951	
				0.446	0.456	0.440	0.438	0.447	0.451	0.454	
	4	3.869 0.689	Lindley	3.336	3.426	3.236	3.164	3.304	3.244	3.217	
				0.180	0.163	0.246	0.285	0.185	0.231	0.264	
			M-H	2.998	3.043	2.952	2.907	2.982	2.948	2.932	
				0.491	0.501	0.487	0.489	0.494	0.503	0.509	

Table 10: Average and MSE values of all estimates of $h(t)$ for different choices of n, m, k, r and $t = 0.25$.

						p			w		
						-0.5	0.5	1.0	-0.5	0.5	1.0
(n, m)	r	k	$\hat{R}(t)$		$\tilde{R}_{B_1}(t)$	$\tilde{R}_{B_2}(t)$			$\tilde{R}_{B_3}(t)$		
(30,30)	(0*30)	2	3.687 0.702	Lindley	3.439	3.515	3.334	3.260	3.407	3.347	3.320
					0.222	0.209	0.229	0.265	0.225	0.236	0.246
				M-H	3.017	3.050	2.983	2.949	3.005	2.982	2.969
					0.421	0.430	0.414	0.412	0.421	0.424	0.427
		4	3.855 0.698	Lindley	3.449	3.577	3.335	3.261	3.412	3.344	3.322
					0.194	0.237	0.235	0.310	0.200	0.241	0.253
				M-H	2.976	3.014	2.936	2.897	2.962	2.933	2.919
					0.453	0.459	0.450	0.452	0.455	0.462	0.467
(50,30)	(20,0*29)	2	3.707 0.675	Lindley	3.475	3.559	3.369	3.286	3.441	3.382	3.356
					0.215	0.182	0.216	0.241	0.215	0.224	0.233
				M-H	3.023	3.057	2.989	2.955	3.011	2.988	2.975
					0.425	0.435	0.418	0.415	0.425	0.428	0.430
		4	3.756 0.945	Lindley	3.459	3.584	3.325	3.228	3.420	3.342	3.308
					0.212	0.189	0.233	0.293	0.218	0.241	0.260
				M-H	2.999	3.042	2.955	2.911	2.984	2.952	2.937
					0.451	0.462	0.444	0.443	0.452	0.457	0.461
	(0*29, 20)	2	3.830 0.678	Lindley	3.386	3.503	3.270	3.197	3.351	3.287	3.259
					0.167	0.161	0.221	0.250	0.170	0.203	0.227
				M-H	0.973	3.014	2.930	2.889	2.958	2.927	2.911
					0.468	0.473	0.468	0.472	0.472	0.482	0.488
		4	3.917 0.605	Lindley	3.209	3.478	3.133	3.084	3.192	3.142	3.118
					0.266	0.223	0.250	0.287	0.203	0.206	0.241
				M-H	2.948	2.996	2.900	2.852	2.931	2.895	2.877
					0.538	0.542	0.538	0.544	0.543	0.558	0.567
(50,40)	(10,0*39)	2	3.668 0.521	Lindley	3.481	3.545	3.410	3.345	3.458	3.420	3.398
					0.194	0.194	0.202	0.219	0.197	0.209	0.213
				M-H	3.021	3.053	2.989	2.957	3.010	2.988	2.977
					0.322	0.330	0.317	0.315	0.322	0.323	0.324
		4	3.745 0.699	Lindley	3.506	3.581	3.389	3.297	3.474	3.406	3.375
					0.236	0.170	0.229	0.254	0.237	0.241	0.248
				M-H	2.998	3.038	2.957	2.917	2.984	2.955	2.940
					0.358	0.369	0.353	0.353	0.359	0.363	0.366
	(0*39, 10)	2	3.701 0.681	Lindley	3.480	3.561	3.384	3.302	3.449	3.394	3.368
					0.214	0.182	0.218	0.240	0.213	0.223	0.230
				M-H	2.996	3.031	2.961	2.926	2.984	2.959	2.947
					0.349	0.356	0.346	0.347	0.351	0.355	0.357
		4	3.760 0.686	Lindley	3.451	3.590	3.306	3.217	3.410	3.328	3.302
					0.198	0.257	0.256	0.342	0.200	0.250	0.247
				M-H	2.999	3.041	2.957	2.916	2.985	2.955	2.940
					0.438	0.450	0.431	0.430	0.439	0.443	0.446
(50,50)	(0*50)	2	3.615 0.422	Lindley	3.480	3.530	3.417	3.359	3.461	3.426	3.406
					0.201	0.199	0.201	0.208	0.203	0.208	0.209
				M-H	3.031	3.060	3.001	2.972	3.021	3.001	2.990
					0.303	0.312	0.297	0.293	0.303	0.303	0.303
		4	3.677 0.506	Lindley	3.499	3.566	3.401	3.322	3.471	3.416	3.388
					0.210	0.184	0.204	0.225	0.211	0.215	0.218
				M-H	3.055	3.091	3.019	2.983	3.043	3.018	3.005
					0.376	0.386	0.370	0.367	0.376	0.378	0.379

Table 11: The 95% confidence and HPD intervals of α and β for different choices of n, m, k, r .

(n, m)	r	k	α		β	
			Asy. CI AL/CP	HPD AL/CP	Asy. CI AL/CP	HPD AL/CP
(30,20)	(10, 0^{*19})	2	(0.380, 1.949)	(1.025, 1.964)	(0.232, 0.585)	(0.351, 0.653)
			1.568/0.983	0.939/0.729	0.353/0.865	0.3021/0.889
		4	(0.240, 1.905)	(1.060, 2.00)	(0.247, 0.573)	(0.381, 0.629)
			1.664/0.999	0.949/0.717	0.325/0.886	0.247/0.247
	$(0^{*19}, 10)$	2	(0.206, 2.048)	(1.053, 2.023)	(0.217, 0.609)	(0.364, 0.642)
			1.84201/0.999	0.970/0.698	0.392/0.958	0.278/0.878
		4	(0.034, 2.071)	(1.008, 1.971)	(0.235, 0.612)	(0.377, 0.603)
			2.036/0.994	0.963/0.668	0.377/0.984	0.226/0.865
	$(0^{*5}, 1^{*10}, 0^{*5})$	2	(0.289, 2.011)	(1.057, 2.011)	(0.231, 0.593)	(0.364, 0.646)
			1.722/0.991	0.954/0.729	0.362/0.904	0.281/0.894
		4	(0.142, 1.990)	(1.029, 1.965)	(0.248, 0.586)	(0.380, 0.608)
			1.848/0.999	0.936/0.681	0.338/0.954	0.228/0.872
(30,25)	(5, 0^{*24})	2	(0.464, 1.920)	(1.069, 1.975)	(0.244, 0.577)	(0.362, 0.649)
			1.456/0.977	0.905/0.732	0.333/0.863	0.287/0.887
		4	(0.320, 1.849)	(1.019, 1.945)	(0.259, 0.566)	(0.379, 0.616)
			1.529/0.999	0.925/0.724	0.307/0.891	0.236/0.892
	$(0^{*24}, 5)$	2	(0.385, 1.963)	(1.062, 1.985)	(0.239, 0.594)	(0.363, 0.639)
			1.578/0.983	0.923/0.715	0.354/0.919	0.275/0.861
		4	(0.228, 1.925)	(1.023, 1.965)	(0.252, 0.580)	(0.379, 0.606)
			1.696/0.997	0.941/0.714	0.327/0.952	0.226/0.866
	$(0^{*10}, 1^{*5}, 0^{*10})$	2	(0.426, 1.951)	(1.052, 1.983)	(0.246, 0.585)	(0.365, 0.645)
			1.524/0.977	0.930/0.774	0.338/0.914	0.279/0.892
		4	(0.280, 1.877)	(1.072, 2.009)	(0.260, 0.571)	(0.391, 0.618)
			1.597/0.997	0.936/0.721	0.311/0.931	0.226/0.886
(30,30)	(0^{*30})	2	(0.534, 1.930)	(1.0961, 1.998)	(0.260, 0.585)	(0.371, 0.645)
			1.395/0.972	0.902/0.755	0.324/0.901	0.274/0.897
		4	(0.382, 1.810)	(1.067, 1.958)	(0.267, 0.558)	(0.387, 0.609)
			1.428/0.994	0.891/0.721	0.290/0.874	0.221/0.876
	(50,30)	2	(0.522, 1.763)	(1.093, 1.981)	(0.260, 0.538)	(0.375, 0.632)
			1.241/0.971	0.888/0.742	0.278/0.729	0.257/0.881
		4	(0.398, 1.731)	(1.078, 1.987)	(0.279, 0.541)	(0.394, 0.609)
			1.333/0.989	0.909/0.717	0.261/0.786	0.214/0.866
	$(0^{*29}, 20)$	2	(0.339, 1.855)	(1.0485, 1.987)	(0.252, 0.572)	(0.378, 0.613)
			1.515/0.993	0.938/0.736	0.319/0.922	0.234/0.872
		4	(0.185, 1.902)	(1.035, 1.949)	(0.268, 0.577)	(0.395, 0.586)
			1.7168/0.997	0.914/0.69	0.308/0.969	0.190/0.839
	$(0^{*10}, 2^{*10}, 0^{*10})$	2	(0.431, 1.819)	(1.068, 1.961)	(0.267, 0.556)	(0.383, 0.616)
			1.388/0.987	0.893/0.761	0.288/0.855	0.232/0.889
		4	(0.301, 1.813)	(1.061, 1.973)	(0.284, 0.556)	(0.399, 0.589)
			1.512/0.996	0.912/0.738	0.271/0.924	0.190/0.867
(50,40)	(10, 0^{*39})	2	(0.613, 1.746)	(1.118, 1.977)	(0.277, 0.539)	(0.384, 0.623)
			1.132/0.944	0.858/0.769	0.261/0.75	0.239/0.869
		4	(0.480, 1.673)	(1.083, 1.964)	(0.289, 0.530)	(0.399, 0.597)
			1.192/0.975	0.881/0.772	0.240/0.78	0.197/0.878

Table 12: The 95% confidence and HPD intervals of α and β for different choices of n, m, k, r .

(n, m)	r	k	α		β	
			Asy. CI AL/CP	HPD AL/CP	Asy. CI AL/CP	HPD AL/CP
$(0^{*39}, 10)$	2	2	(0.530, 1.764)	(1.107, 1.983)	(0.272, 0.550)	(0.387, 0.617)
			1.23375/0.969	0.876/0.751	0.277/0.844	0.229/0.887
	4	2	(0.391, 1.738)	(1.083, 1.963)	(0.287, 0.548)	(0.401, 0.587)
			1.347/0.989	0.880/0.722	0.261/0.904	0.186/0.853
$(0^{*15}, 1^{*10}, 0^{*15})$	2	2	(0.575, 1.763)	(1.097, 1.962)	(0.279, 0.541)	(0.385, 0.616)
			1.188/0.957	0.865/0.759	0.262/0.796	0.230/0.896
	4	2	(0.442, 1.708)	(1.083, 1.947)	(0.293, 0.536)	(0.404, 0.590)
			1.265/0.994	0.863/0.758	0.243/0.835	0.186/0.865
$(50, 50)$	2	2	(0.685, 1.748)	(1.128, 1.957)	(0.294, 0.544)	(0.390, 0.617)
			1.063/0.925	0.828/0.771	0.249/0.817	0.226/0.886
	4	2	(0.542, 1.641)	(1.129, 1.974)	(0.299, 0.523)	(0.407, 0.591)
			1.099/0.955	0.844/0.743	0.224/0.759	0.184/0.879

Grouped real data set: In this case k is 2

{0.83021, 0.83351}, {0.84521, 0.84972}, {0.85917, 0.86390}, {0.86923, 0.87736},
 {0.88851, 0.89280}, {0.89285, 0.89890}, {0.89983, 0.90098}, {0.90335, 0.90989},
 {0.91137, 0.92005}, {0.92664, 0.93181}, {0.93208, 0.93291}, {0.93604, 0.94027},
 {0.94104, 0.94340}, {0.94538, 0.94692}, {0.94741, 0.94791}, {0.94923, 0.95076},
 {0.95219, 0.95285}, {0.96016, 0.96016}, {0.96043, 0.96230}, {0.96417, 0.96538},
 {0.96587, 0.96620}.

Grouped real data set: In this case k is 3

{0.83021, 0.83351, 0.84521}, {0.84972, 0.85917, 0.86390},
 {0.86923, 0.87736, 0.88851}, {0.89280, 0.89285, 0.89890},
 {0.89983, 0.90098, 0.90335}, {0.90989, 0.91137, 0.92005},
 {0.92664, 0.93181, 0.93208}, {0.93291, 0.93604, 0.94027},
 {0.94104, 0.94340, 0.94538}, {0.94692, 0.94741, 0.94791},
 {0.94923, 0.95076, 0.95219}, {0.95285, 0.96016, 0.96016},
 {0.96043, 0.96230, 0.96417}, {0.96538, 0.96587, 0.96620}.

We have generated progressive first failure censored samples from the above data with different n, k, m and censoring schemes. The generated samples are presented in Table 13. The maximum likelihood estimates of unknown parameters say $\hat{\alpha}$, $\hat{\beta}$, and reliability characteristics are obtained using the EM algorithm under arbitrarily chosen censoring schemes. We also compute different Bayes estimates using Lindley's approximation and M-H algorithm. These Bayes estimates are obtained under noninformative prior distribution as we take $(a_1, b_1, a_2, b_2) = (0, 0, 0, 0)$. The maximum likelihood estimates and all the Bayes estimates under noninformative prior are tabulated in Table 14. It can be seen that all the estimates are close to each other. The results of reliability function for $t = 0.95$ are given in Table 15. In Table 16, the hazard rate function is calculated for $t = 0.85$. Also, in Table 17, the 95% confidence intervals are reported

under classical as well as Bayesian approach. We can conclude that among all the intervals in terms of their length, the HPD performs better.

Table 13: Progressive first failure censored samples for different n, k, m, r .

(n, k, m)	r	Progressive first failure censored sample
(21, 2, 16)	$(5, 0^{*15})$	(0.83021, 0.84521, 0.85917, 0.89285, 0.89983, 0.90335, 0.91137, 0.93209, 0.94104, 0.94741, 0.94923, 0.95219, 0.96016, 0.960439, 0.96417, 0.96587)
	$(0^{*15}, 5)$	(0.83021, 0.84521, 0.85917, 0.86923, 0.88851, 0.89285, 0.89983, 0.90335, 0.91137, 0.92664, 0.93209, 0.93604, 0.94104, 0.94538, 0.94741, 0.94923)
	$(0^{*7}, 2, 3, 0^{*7})$	(0.83021, 0.84521, 0.85917, 0.86923, 0.88851, 0.89285, 0.89983, 0.90335, 0.91137, 0.92664, 0.93209, 0.94104, 0.94923, 0.95219, 0.96417, 0.96587)
(14, 3, 10)	$(4, 0^{*9})$	(0.83021, 0.84972, 0.86923, 0.89983, 0.92664, 0.93291, 0.94104, 0.94923, 0.95285, 0.96538)
	$(0^{*9}, 4)$	(0.83021, 0.84972, 0.86923, 0.8928, 0.89983, 0.90989, 0.92664, 0.93291, 0.94104, 0.94692)
	$(0^{*4}, 2, 2, 0^{*4})$	(0.83021, 0.84972, 0.86923, 0.8928, 0.89983, 0.90989, 0.93291, 0.94692, 0.94923, 0.96538)

Table 14: Point estimates of α and β based on generated progressive first failure censored samples from data set 1.

k	r	$\hat{\alpha}$		$\tilde{\alpha}_{B_1}$		$\tilde{\alpha}_{B_2}$			$\tilde{\alpha}_{B_3}$		
		$\hat{\beta}$		$\tilde{\beta}_{B_1}$		$\tilde{\beta}_{B_2}$			$\tilde{\beta}_{B_3}$		
						$p=-0.5$	$p=0.5$	$p=1.0$	$w=-0.5$	$w=0.5$	$w=1.0$
2	(5, 0*15)	2.436	Lindley	2.522	2.805	2.216	1.997	2.391	2.156	2.065	
			M-H	1.469	1.470	1.469	1.469	1.469	1.469	1.469	
		27.283	Lindley	26.961	27.783	23.566	24.177	26.581	25.864	25.548	
			M-H	22.950	22.950	22.950	22.949	22.950	22.950	22.950	
	(0*15, 5)	2.097	Lindley	2.193	2.448	1.913	1.704	2.055	1.813	1.724	
			M-H	1.444	1.444	1.444	1.443	1.444	1.443	1.443	
		24.756	Lindley	24.263	28.212	21.133	21.712	23.877	23.163	22.847	
			M-H	23.096	23.096	23.096	23.096	23.096	23.096	23.096	
	(0*7, 2, 3, 0*7)	2.810	Lindley	2.884	3.295	2.445	2.182	2.716	2.423	2.315	
			M-H	1.511	1.511	1.511	1.511	1.511	1.511	1.511	
		26.844	Lindley	26.256	27.344	23.322	23.864	25.927	25.313	25.038	
			M-H	23.026	23.026	23.026	23.026	23.026	23.026	23.026	
3	(4, 0*9)	1.780	Lindley	1.884	2.127	1.615	1.410	1.729	1.467	1.380	
			M-H	1.592	1.593	1.591	1.590	1.592	1.590	1.590	
		25.827	Lindley	25.367	29.911	21.627	22.452	24.841	23.874	23.458	
			M-H	22.979	22.979	22.978	22.978	22.979	22.979	22.978	
	(0*9, 4)	1.224	Lindley	1.308	1.449	1.155	1.021	1.182	0.975	0.911	
			M-H	1.503	1.503	1.502	1.502	1.502	1.502	1.502	
		23.125	Lindley	22.362	24.175	18.884	19.738	21.780	20.756	20.338	
			M-H	22.922	22.992	22.991	22.991	22.992	22.992	22.992	
	(0*4, 2, 2, 0*4)	1.838	Lindley	1.921	2.215	1.604	1.381	1.742	1.454	1.365	
			M-H	1.596	1.596	1.596	1.596	1.596	1.596	1.596	
		26.184	Lindley	25.266	26.684	21.962	22.813	24.767	23.874	23.497	
			M-H	23.039	23.041	23.039	23.039	23.039	23.039	23.039	

Table 15: Point estimates of $R(t)$ ($t = 0.95$) based on generated progressive first failure censored samples from data set 1.

k	r	$\hat{R}(t)$		$\tilde{R}_{B_1}(t)$	p			w		
					-0.5	0.5	1.0	-0.5	0.5	1.0
2	$(5, 0^{*15})$	0.501	Lindley	0.512	$\tilde{R}_{B_2}(t)$			$\tilde{R}_{B_3}(t)$		
			M-H	0.500	0.514	0.510	0.508	0.508	0.501	0.497
	$(0^{*15}, 5)$	0.500	Lindley	0.511	0.501	0.500	0.500	0.500	0.500	0.500
			M-H	0.517	0.513	0.509	0.507	0.507	0.499	0.495
	$(0^{*7}, 2, 3, 0^{*7})$	0.441	Lindley	0.457	0.517	0.517	0.516	0.516	0.516	0.516
			M-H	0.435	0.459	0.454	0.452	0.452	0.442	0.437
3	$(4, 0^{*9})$	0.576	Lindley	0.592	0.435	0.435	0.435	0.435	0.435	0.435
			M-H	0.588	0.594	0.589	0.587	0.587	0.578	0.574
	$(0^{*9}, 4)$	0.6400	Lindley	0.651	0.588	0.588	0.588	0.588	0.588	0.588
			M-H	0.624	0.653	0.649	0.646	0.647	0.640	0.637
	$(0^{*4}, 2, 2, 0^{*4})$	0.573	Lindley	0.591	0.624	0.6241	0.624	0.624	0.623	0.623
			M-H	0.576	0.594	0.588	0.585	0.586	0.576	0.572
					0.576	0.576	0.576	0.576	0.576	0.575

Table 16: Point estimates of $h(t)$ ($t = 0.85$) based on generated progressive first failure censored samples from data set 1.

k	r	$\hat{h}(t)$		$\tilde{h}_{B_1}(t)$	p			w		
					-0.5	0.5	1.0	-0.5	0.5	1.0
2	$(5, 0^{*15})$	0.939	Lindley	1.009	$\tilde{h}_{B_2}(t)$			$\tilde{h}_{B_3}(t)$		
			M-H	0.946	1.061	0.953	0.898	0.949	0.838	0.796
	$(0^{*15}, 5)$	1.104	Lindley	1.150	0.946	0.946	0.946	0.946	0.946	0.946
			M-H	1.057	1.202	1.096	1.044	1.101	1.010	0.973
	$(0^{*7}, 2, 3, 0^{*7})$	1.145	Lindley	1.206	1.058	1.057	1.056	1.057	1.056	1.055
			M-H	1.165	1.262	1.146	1.089	1.154	1.056	1.015
3	$(4, 0^{*9})$	0.825	Lindley	0.875	1.165	1.165	1.164	1.165	1.164	1.164
			M-H	0.798	0.923	0.824	0.775	0.814	0.706	0.668
	$(0^{*9}, 4)$	0.795	Lindley	0.824	0.798	0.797	0.797	0.797	0.796	0.796
			M-H	0.844	0.860	0.788	0.753	0.779	0.697	0.666
	$(0^{*4}, 2, 2, 0^{*4})$	0.816	Lindley	0.868	0.844	0.844	0.843	0.844	0.843	0.843
			M-H	0.808	0.908	0.826	0.784	0.816	0.722	0.685
					0.808	0.807	0.807	0.807	0.806	0.806

Table 17: 95% Interval estimates of α and β based on generated progressive first failure censored samples from data set 1.

k	r	α		β	
		Asy. CI	HPD	Asy. CI	HPD
2	$(5, 0^{*15})$	(0.223, 4.650)	(1.410, 1.548)	(14.575, 39.990)	(22.866, 22.984)
	$(0^{*15}, 5)$	(0, 4.203)	(1.415, 1.477)	(12.515, 36.997)	(23.039, 23.146)
	$(0^{*7}, 2, 3, 0^{*7})$	(0.110, 5.510)	(1.453, 1.549)	(15.054, 38.634)	(22.990, 23.052)
3	$(4, 0^{*9})$	(0, 3.838)	(1.511, 1.678)	(11.227, 40.426)	(22.923, 23.017)
	$(0^{*9}, 4)$	(0, 2.764)	(1.469, 1.558)	(8.505, 37.746)	(22.945, 23.063)
	$(0^{*4}, 2, 2, 0^{*4})$	(0, 4.094)	(1.561, 1.632)	(11.731, 40.638)	(22.976, 23.097)

Data set 2. The second application is the annual capacity of Naser Lake during the flood time from 1968 to 2010 obtained from Abdel-Latif and Yacoub (2011). Naser Lake is located in the lower Nile River Basin at the border between Egypt and Sudan.

The total capacity of the lake is 162.3×10^9 Cubic meter (m^3) at its highest water level. The lake was created from the late 1960 to the 1968 together with the construction of the Aswan High Dam upstream of the old Aswan Dam, about 5 km south of the city of Aswan. The data are as follows

0.24627, 0.28607, 0.37806, 0.43019, 0.35268, 0.41096, 0.49328, 0.66771, 0.64725, 0.67060, 0.68576, 0.63536, 0.63148, 0.61090, 0.49926, 0.44941, 0.31706, 0.33086, 0.29125, 0.25058, 0.46691, 0.45298, 0.42667, 0.45736, 0.53641, 0.59180, 0.66543, 0.67757, 0.76278, 0.73937, 0.73937, 0.77270, 0.77461, 0.76980, 0.75187, 0.65730, 0.59051, 0.56598, 0.69747, 0.74097, 0.69377, 0.59581.

The corresponding Kolmogorov-Smirnov p-value using Data set 2 is 0.6758. For more details about goodness of fit studies like Data Set 1, we again refer to Sharaf El-Deen et al. (2014).

Grouped real data set: In this case k is 2

{0.24627,0.25058},{0.28607,0.29125},{0.31706,0.33086},{0.35268,0.37806},
{0.41096,0.42667},{0.43019,0.44941},{0.45298,0.45736},{0.46691,0.49328},
{0.49926,0.53641},{0.56598,0.59051},{0.59180,0.59581},{0.61090,0.63148},
{0.63536,0.64725},{0.65730,0.66543},{0.66771,0.67060},{0.67757,0.68576},
{0.69377,0.69747},{0.73937,0.73937},{0.74097,0.75187},{0.76278,0.76980},
{0.77270,0.77461}.

Grouped real data set: In this case k is 3

{0.24627,0.25058,0.28607},{0.29125,0.31706,0.33086},
{0.35268,0.37806,0.41096},{0.42667,0.43019,0.44941},
{0.45298,0.45736,0.46691},{0.49328,0.49926,0.53641},
{0.56598,0.59051,0.59180},{0.59581,0.61090,0.63148},
{0.63536,0.64725,0.65730},{0.66543,0.66771,0.67060},
{0.67757,0.68576,0.69377},{0.69747,0.73937,0.73937},
{0.74097,0.75187,0.76278},{0.76980,0.77270,0.77461}.

Progressive first failure censored samples are generated from the above data with different n, k, m and censoring schemes. The generated samples are presented in Table 18.

Table 19 contains the point estimates of unknown parameters α and β under classical and Bayesian approach. In Table 20, the reliability function is calculated for $t = 0.95$. The results of hazard rate function for $t = 0.85$ are given in Table 21. Table 22 contains the 95% confidence and HPD intervals of unknown parameters. All the Bayes estimates and HPD intervals are evaluated under noninformative prior distribution. We see that the HPD with respect to its length performs better than the ACI. We found out that results are so consistent with simulation studies.

Table 20: Point estimates of $R(t)$ ($t = 0.95$) based on generated progressive first failure censored samples from data set 2.

					p			w		
					-0.5	0.5	1.0	-0.5	0.5	1.0
k	r	$\hat{R}(t)$		$\tilde{R}_{B_1}(t)$	$\tilde{R}_{B_2}(t)$			$\tilde{R}_{B_3}(t)$		
2	$(5, 0^{*15})$	0.171	Lindley	0.196	0.198	0.195	0.193	0.185	0.163	0.153
			M-H	0.177	0.177	0.177	0.177	0.177	0.177	
	$(0^{*15}, 5)$	0.255	Lindley	0.280	0.282	0.277	0.274	0.268	0.246	0.236
			M-H	0.254	0.254	0.254	0.254	0.254	0.254	
	$(0^{*7}, 2, 3, 0^{*7})$	0.225	Lindley	0.251	0.254	0.249	0.246	0.240	0.217	0.207
			M-H	0.219	0.219	0.219	0.219	0.219	0.219	
3	$(4, 0^{*9})$	0.315	Lindley	0.350	0.354	0.346	0.341	0.336	0.308	0.295
			M-H	0.305	0.305	0.305	0.305	0.305	0.305	
	$(0^{*9}, 4)$	0.429	Lindley	0.455	0.460	0.450	0.445	0.443	0.419	0.408
			M-H	0.441	0.441	0.440	0.440	0.440	0.439	0.438
	$(0^{*4}, 2, 2, 0^{*4})$	0.334	Lindley	0.371	0.376	0.366	0.360	0.355	0.322	0.307
			M-H	0.336	0.336	0.336	0.336	0.336	0.335	0.335

Table 21: Point estimates of $h(t)$ ($t = 0.85$) based on generated progressive first failure censored samples from data set 2.

					p			w		
					-0.5	0.5	1.0	-0.5	0.5	1.0
k	r	$\hat{h}(t)$		$\tilde{h}_{B_1}(t)$	$\tilde{h}_{B_2}(t)$			$\tilde{h}_{B_3}(t)$		
2	$(5, 0^{*15})$	0.358	Lindley	0.388	0.396	0.380	0.371	0.365	0.322	0.305
			M-H	0.352	0.352	0.352	0.352	0.352	0.352	
	$(0^{*15}, 5)$	0.283	Lindley	0.313	0.318	0.307	0.302	0.293	0.255	0.240
			M-H	0.280	0.280	0.280	0.280	0.280	0.279	
	$(0^{*7}, 2, 3, 0^{*7})$	0.251	Lindley	0.284	0.289	0.279	0.274	0.264	0.225	0.210
			M-H	0.257	0.257	0.256	0.256	0.254	0.253	
3	$(4, 0^{*9})$	0.188	Lindley	0.223	0.227	0.218	0.214	0.199	0.156	0.142
			M-H	0.197	0.197	0.196	0.196	0.194	0.193	
	$(0^{*9}, 4)$	0.213	Lindley	0.238	0.242	0.234	0.230	0.219	0.183	0.171
			M-H	0.202	0.202	0.202	0.202	0.202	0.202	
	$(0^{*4}, 2, 2, 0^{*4})$	0.180	Lindley	0.214	0.217	0.210	0.206	0.193	0.153	0.140
			M-H	0.189	0.189	0.189	0.189	0.189	0.189	

Table 22: The 95% interval estimates of α and β based on generated progressive first failure censored samples from data set 2.

k	r	α		β	
		Asy. CI	HPD	Asy. CI	HPD
2	$(5, 0^{*15})$	(0.291, 5.212)	(0.920, 1.079)	(1.888, 4.810)	(2.937, 3.073)
	$(0^{*15}, 5)$	(0, 4.212)	(0.992, 1.074)	(1.677, 4.955)	(3.002, 3.098)
	$(0^{*7}, 2, 3, 0^{*7})$	(0.068, 4.918)	(1.0551, 1.221)	(15.054, 38.634)	(2.889, 2.952)
3	$(4, 0^{*9})$	(0, 4.186)	(1.011, 1.160)	(1.644, 5.590)	(2.956, 3.045)
	$(0^{*9}, 4)$	(0, 2.712)	(0.916, 1.050)	(1.122, 5.005)	(3.034, 3.118)
	$(0^{*4}, 2, 2, 0^{*4})$	(0, 4.229)	(0.943, 1.025)	(1.525, 5.684)	(2.958, 3.035)

6 Conclusions

In this paper, we have discussed classical and Bayesian inferential procedures for the progressive first failure censored data from the Kumaraswamy distribution. We have

provided MLEs and Bayes estimates as well as the corresponding asymptotic CIs and HPD credible intervals. A simulation study has been carried out to compare the performance of different methods. The simulation results showed that as the sample size n , m and the group size k increase, the MSE decreases. Also, we can conclude that the M-H algorithm performs better than other Bayesian estimates as well as the MLE. In addition, with increasing the failure proportion m/n , the estimates improve. Based on interval estimation, the HPD interval performs better than the ACI in terms of AL. We investigated the relation between the group size k and AL under ACI and HPD. Although the consequences have been derived under the progressive first failure censoring schemes, similar methods can be extended for other censoring schemes. We believe that more work can be done in these contexts.

Acknowledgment

The authors would like to thank the referees for their careful reading and their useful comments which led to this considerably improved version.

References

- Abdel-Latif, M. and Yacoub, M. (2011). Effect of change of discharges at dongola station due to sedimentation on the water losses from Nasser lake. *Nile Basin Water Science & Engineering Journal*, **4**, 86–98.
- Ahmadi, M.V., Doostparast, M. and Ahmadi, J. (2013). Estimating the lifetime performance index with Weibull distribution based on progressively first-failure censoring scheme. *Journal of Computational and Applied Mathematics*, **239**, 93–102.
- Balakrishnan, N. and Aggarwala, R. (2000). *Progressive Censoring: Theory, Methods and Applications*. Boston: Birkhauser.
- Balakrishnan, N. and Cramer, E. (2014). *The Art of Progressive Censoring: Applications to Reliability and Quality*. New York: Birkhauser.
- Balakrishnan, N., Zhang, L. and Xie, Q. (2009) Inference for a simple step-stress model with type-I censoring and lognormally distributed lifetimes. *Communications in Statistics -Theory and Methods*, **38**, 1690–1709.
- Balasooriya, U. (1995). Failure-Censored reliability sampling plans for the exponential distribution. *Journal of Statistical Computation and Simulation*, **52**, 337–349.
- Calabria, H. and Pulcini, G. (1994). An engineering approach to Bayes estimation for the Weibull distribution. *Microelectronics Reliability*, **34**, 789–802.
- Calabria, H. and Pulcini, G. (1996). Point estimation under asymmetric loss functions for left truncated exponential samples. *Communications in Statistics -Theory and Methods*, **25**, 585–600.
- Chen, M.H., and Shao, Q. M. (1999). Monte Carlo estimation of Bayesian credible and HPD intervals. *Journal of Computational and Graphical Statistics*, **8**, 69–92.

- Cohen, A.C. and Norgaard, N.J. (1977). Progressively censored sampling in the three-parameter gamma distribution. *Technometrics*, **19**, 333–340.
- Dempster, A.P., Laird, N.M. and Rubin, D.B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society, Series B (Methodological)*, **39**, 1–38.
- Geman, S. and Geman, D. (1984). Stochastic relaxation, gibbs distributions and the Bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **6**, 721–741.
- Ghitany, M.E., Al-Mutairi, D.K., Balakrishnan, N. and Al-Enezi, L.J. (2013). Power Lindley distribution and associated inference. *Computational Statistics & Data Analysis*, **64**, 20–33.
- Gilchrist, W. (1997). Modelling with quantile distribution functions. *Journal of Applied Statistics*, **24**, 113–122.
- Gilks, W.R., Richardson, S., Spiegelhalter, D.J. (1995). Markov Chain Monte Carlo in Practice. Chapman & Hall: London.
- Gilks, W. and Wild, P. (1992). Adaptive rejection sampling for Gibbs sampling. *Applied Statistics*, **41**, 337–348.
- Heba, S.M., Ateya, S.F. and Al-Hussaini, E.K. (2017). Estimation based on progressive first-failure censoring from exponentiated exponential distribution. *Journal of Applied Statistics*, **44**, 1479–1494.
- Huang, S.R. and Wu, S.J. (2012). Bayesian estimation and prediction for Weibull model with progressive censoring. *Journal of Statistical Computation and Simulation*, **82**, 1607–1620.
- Johnson, L.G. (1964). *Theory and Technique of Variation Research*. Amsterdam: Elsevier.
- Jones, M.C. (2009). Kumaraswamy's distribution: A beta-type distribution with some tractability advantages. *Statistical Methodology*, **6**, 70–81.
- Kundu, D. and Pradhan, B. (2009). Bayesian inference and life testing plans for generalized exponential distribution. *Science in China Series A: Mathematics*, **52**, 1373–1388.
- Lindley, D.V. (1980). Approximate Bayesian methods. *Trabajos de Estadística*, **31**, 223–237.
- Louis, T.A. (1982). Finding the observed information matrix when using the EM algorithm. *Journal of the Royal Statistical Society, Series B (Methodological)*, **44**, 226–233.
- Ng, H.K.T., Chan, P.S. and Balakrishnan, N. (2002). Estimation of parameters from progressively censored data using EM algorithm. *Computational Statistics & Data Analysis*, **39**, 371–386.

- Pradhan, B. and Kundu, D. (2009). On progressively censored generalized exponential distribution. *Test*, **18**, 497–515.
- Press, S. J. and Tanur, J. M. (2001). The subjectivity of scientists and the Bayesian approach. Wiley: New York.
- Reyad, H.M., and Ahmed, S.O. (2016). Bayesian and E-Bayesian estimation for the Kumaraswamy distribution based on type-II censoring. *International Journal of Advanced Mathematical Sciences*, **4**, 10–17.
- Sharaf El-Deen, M.M., AL-Dayian, G.R. and EL-Helbawy, A.A. (2014). Statistical inference for Kumaraswamy distribution based on generalized order statistics with applications. *British Journal of Mathematics & Computer Science*, **4**, 1710–1743.
- Sindhu, T.N., Feroze, N., and Aslam, M. (2013). Bayesian analysis of the Kumaraswamy distribution under failure censoring sampling scheme. *International Journal of Advanced Science and Technology*, **51**, 39–58.
- Smith, A.F.M. and Roberts, G.O. (1993). Bayesian computation via the gibbs sampling and related Markov Chain Monte Carlo methods (with discussion). *Journal of the Royal Statistical Society, Series B*, **55**, 3–24.
- Soliman, A.A. (2005). Estimation of parameters of life from progressively censored data using Burr-XII model. *IEEE Transactions on Reliability*, **54**, 34–42.
- Sultana, F., Tripathi, Y.M., Rastogi, M.K. and Wu, S.J. (2017). Parameter estimation for the Kumaraswamy distribution based on hybrid censoring. *American Journal of Mathematical and Management Sciences*, DOI:10.1080/01966324.2017.1396943.
- Tripathi, Y.M. and Rastogi, M.K. (2016). Estimation using hybrid censored data from a generalized inverted exponential distribution. *Communications in Statistics -Theory and Methods*, **45**, 4858–4873.