

ESTIMATION AND PREDICTION BASED ON KIES REAL LIFETIMES TURBOCHARGER TYPE II CENSORED DATA

Ghassan K. Abufoudeh*[†], Raed R. Abu Awwad* , Samer A. Alokaily*, Maalee Almheidat** and Suad Alhihi***

*University of Petra, Jordan.

**The University of Jordan, Jordan.

*** Al-Balqa Applied University, Jordan.

ABSTRACT

This paper presents an investigation into estimation and prediction problems pertaining to life data derived from Kies distribution, utilizing type II censored data. The estimation of scale and shape parameters for Kies Distribution is accomplished through the utilization of both Maximum Likelihood and Bayesian methodologies. The Gibbs sampling technique is employed for generating Markov Chain Monte Carlo (MCMC) samples, and it has been utilized for the purpose of computing Bayes estimates and constructing symmetric credible intervals. The method also includes an approach to estimate the density of future ordered data points before deriving their associated predictions. The proposed estimation and predictor method's performance has been evaluated through simulation studies. In order to validate the outcomes of the simulation, empirical data on the operational lifespan of turbochargers is utilized and subsequently scrutinized.

KEYWORDS: Kies distribution, type II censored data, maximum likelihood estimation, Bayes estimation, Bayes prediction, Gibbs sampling, MCMC samples.

MSC: 94A17, 62F10, 62F15, 62N01.

RESUMEN

Este documento presenta una investigación sobre problemas de estimación y predicción relacionados con datos de vida derivados de la distribución de Kies, utilizando datos censurados de tipo II. La estimación de los parámetros de escala y forma para la distribución de Kies se lleva a cabo mediante la utilización de metodologías de Máxima Verosimilitud y Bayesiana. La técnica de muestreo de Gibbs se emplea para generar muestras de Monte Carlo por Cadenas de Markov (MCMC), y se ha utilizado con el propósito de calcular estimaciones de Bayes y construir intervalos creíbles simétricos. El método también incluye un enfoque para estimar la densidad de futuros puntos de datos ordenados antes de derivar sus predicciones asociadas. El rendimiento del método de estimación y predicción propuesto se ha evaluado a través de estudios de simulación. Para validar los resultados de la simulación, se utilizan datos empíricos sobre la vida operativa de los turbos y se someten a un escrutinio posterior.

[†]Corresponding author: Ghassan K. Abufoudeh, gabufoudeh@uop.edu.jo

PALABRAS CLAVE: Distribución de Kies, datos censurados de tipo II, estimación de máxima verosimilitud, estimación de Bayes, predicción de Bayes, muestreo de Gibbs, muestras MCMC.

1. INTRODUCTION

The Kies distribution was first suggested by Kies in 1958. It is based on the Weibull distribution. Kumar and Dharmaja (2014) looked into the Kies distribution, which has a growing, declining, and “bathtub” hazard rate function. They also showed that it is a good choice to the extended Weibull distribution. Kumer and Dharmaja (2013) looked into the reduced Kies distribution, which is a type of the Kies distribution with one parameter. They found that it has some unique characteristics that are similar to those of the Weibull distribution. Kumar and Dharmaja (2017) also showed and looked into an exponented version of the reduced Kies distribution with only two factors. Kumar and Dharmaja (2017) also came up with a new distribution called the modified Kies distribution, which is a broader version of the extended reduced Kies distribution. Kundu and Raqib (2012) wrote about Bayesian inference and forecast of the two parameters of the Weibull distribution for type II censored data. Kundu and Howlader (2010) wrote about Bayesian inference and forecast of the inverse Weibull distribution for type II censored data. Ghassan and el at. (2016) looked into how to use Rayleigh type II filtered data to figure out the remaining life of ball bearings. Pradhan and Kundu (2011) looked into how the Bayesian method can be used to measure and predict the two-parameter Gamma distribution. Bayes estimation with a squared error loss function was looked at. For the scale parameter, a Gamma prior was used, and for the shape parameter, a log-concave prior was used. They came up with a method called Gibbs sampling to get samples from the posterior density that could be used to make approximate Bayes estimates and build believable intervals. Al-Hussaini (1999) looked into the Bayesian prediction problem for a wide range of life distributions. Nesreen and el at (2021) examined Bayesian and classical inference for the Kies distribution parameters using recorded data. The probability density function (PDF) and cumulative distribution function (CDF) of the two parameters Kies distribution are:

$$f(x; \beta, \lambda) = \frac{\beta \lambda x^{\beta-1}}{(1-x)^{\beta+1}} e^{-\lambda \left(\frac{x}{1-x}\right)^\beta}, \quad (1.1)$$

and

$$F(x; \beta, \lambda) = 1 - e^{-\lambda \left(\frac{x}{1-x}\right)^\beta}, \quad (1.2)$$

where $0 < x < 1$, and $\lambda > 0$ and $\beta > 0$ are the scale and shape parameters, respectively.

2. TYPE II CENSORED DATA

Assume that a set of n objects are being monitored until the point of failure. The objects under consideration in reliability study experiments may include systems, components, or computer chips, while in clinical trials, patients may be subjected to specific drug or clinical conditions. In such cases, the lifetimes $\underline{X} = (X_1, X_2, \dots, X_n)$ of these objects or patients may be modeled using the Kies distribution, which is characterized by a probability density function (PDF).

It is possible to conclude an experiment at the r th failure, resulting in a type II censored sample at time $X_{r:n}$. The variable r is held constant, whereas the duration of the experiment, denoted as $X_{r:n}$, is subject to randomness. The probability function in this instance can be expressed as:

$$L(\beta, \lambda | \underline{X}) = \frac{n!}{(n-r)!} \prod_{i=1}^r f(x_i | \beta, \lambda) [1 - F(x_r; \beta, \lambda)]^{n-r}, x_1 < x_2 < \dots < x_r. \quad (2.1)$$

3. MAXIMUM LIKELIHOOD ESTIMATION

Using type II censored data, we calculate the maximum likelihood estimators (MLEs) for the Kies model parameters β and λ .

Consider $X_{1:n} < X_{2:n} < \dots < X_{r:n}$ as a type II censored sample of size r ($1 < r < n$). Using (1), (2), and (3), the likelihood function is given as

$$L(\beta, \lambda | \text{data}) = \frac{n!}{(n-r)!} \beta^r \lambda^r \prod_{i=1}^r \frac{x_i^{\beta-1}}{(1-x_i)^{\beta+1}} e^{-\lambda \sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta} e^{-\lambda(n-r)\left(\frac{x_r}{1-x_r}\right)^\beta}.$$

The differentiation of the natural logarithm of the likelihood function with respect to both β and λ when set to zero allows us to obtain

$$\frac{r}{\beta} + \sum_{i=1}^r \ln x_i - \sum_{i=1}^r \ln(1-x_i) - \frac{r \left[\sum_{i=1}^r \ln \left(\frac{x_i}{1-x_i} \right) \cdot \left(\frac{x_i}{1-x_i} \right)^\beta + (n-r) \ln \left(\frac{x_r}{1-x_r} \right) \cdot \left(\frac{x_r}{1-x_r} \right)^\beta \right]}{\sum_{i=1}^r \left(\frac{x_i}{1-x_i} \right)^\beta + (n-r) \left(\frac{x_r}{1-x_r} \right)^\beta} = 0, \quad (3.1)$$

and

$$\lambda = \frac{r}{\sum_{i=1}^r \left(\frac{x_i}{1-x_i} \right)^\beta + (n-r) \left(\frac{x_r}{1-x_r} \right)^\beta}. \quad (3.2)$$

The MLE of β , denoted as $\hat{\beta}$, is a numerical solution of Eq. (4). After calculating the MLE of β , Eq.(5) provides the MLE of λ , denoted as $\hat{\lambda}$. Balakrishnan and Kateri (2008) provide additional information on the existence and uniqueness of these MLEs.

4. BAYES ESTIMATE AND CREDIBLE INTERVALS

The Bayesian inference and life testing plan requires specific assumptions about prior distributions to calculate Bayes estimates and credible intervals for β as well as λ using type II censored data. The natural choice for the joint prior distribution of β and λ appears as follows:

$$g(\beta, \lambda) \propto \lambda^{a_1-1} e^{-b_1 \lambda} \beta^{a_2-1} e^{-b_2 \beta},$$

with the hyperparameters $a_1, a_2, b_1, b_2 > 0$.

The joint posterior distribution for parameters β and λ can be derived using the joint prior while working with Type II censoring data where $X_{1:n} < X_{2:n} < \dots < X_{r:n}$, ($1 \leq r \leq n$).

$$\begin{aligned} \pi(\beta, \lambda | \text{data}) &\propto \beta^{r+a_2-1} \prod_{i=1}^r \frac{x_i^{\beta-1}}{(1-x_i)^{\beta+1}} e^{-b_2 \beta} \lambda^{r+a_1-1} \\ &\times e^{-\lambda \left[\sum_{i=1}^r \left(\frac{x_i}{1-x_i} \right)^\beta + (n-r) \left(\frac{x_r}{1-x_r} \right)^\beta + b_1 \right]}. \end{aligned}$$

For parameters estimations, authors in the literature employ a variety of error loss functions to address the challenge of estimating parameters. The most two common error loss functions are:

First: square error loss function, defined as:

$$L_S(\theta, \hat{\theta}) = (\theta - \hat{\theta})^2$$

Second: LINEX loss function, defined as:

$$L_L(\theta, \hat{\theta}) = \left(\frac{\hat{\theta}}{\theta}\right)^{a^*} - a^* \ln\left(\frac{\hat{\theta}}{\theta}\right) - 1, a^* \neq 0,$$

where $\hat{\theta}$ is the estimate of θ .

4.1. Case I: Shape Parameter Known

From $\pi(\beta, \lambda | data)$, the posterior density of λ given β and data is:

$$\pi_1(\lambda | \beta, data) \propto \lambda^{r+a_1-1} e^{-\lambda \left[\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1 \right]}$$

That is the conditional density of λ given β and data is:

$$Gamma\left(r + a_1 - 1, \sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1\right) \quad (4.1)$$

Therefore, the Bayes estimate of λ with respect to square error loss function the posterior mean, namely,

$$\hat{\lambda}_{Bayes,S} = \frac{r + a_1 - 1}{\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1}$$

and the Bayes estimate of λ with respect to LINEX loss function is:

$$\begin{aligned} \hat{\lambda}_{Bayes,L} &= [E_{\pi_1}(\lambda^{-a^*} | \beta, data)]^{-\frac{1}{a^*}} \\ &= \left[\frac{\Gamma(r+a_1-a^*)}{\Gamma(r+a_1)} \left(\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1 \right)^{a^*} \right]^{-\frac{1}{a^*}} \end{aligned}$$

To obtain the $(1 - \alpha)100\%$ Bayesian credible interval (L, U) for λ such that

$$\int_L^U \pi_1(\lambda | \beta, data) d\lambda = 1 - \alpha$$

Therefore,

$$\Pr(L < \lambda < \infty) = 1 - \frac{\alpha}{2} \quad \text{and} \quad \Pr(U < \lambda < \infty) = \frac{\alpha}{2}$$

Using $\pi_1(\lambda | \beta, data)$ and $\Pr(L < \lambda < \infty) = 1 - \frac{\alpha}{2}$ we obtain:

$$\Gamma\left(r + a_1, \left(\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1\right) L\right) = \left(1 - \frac{\alpha}{2}\right) \Gamma(r + a_1)$$

Similarly for U we obtain:

$$\Gamma\left(r + a_1, \left(\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1\right) U\right) = \frac{\alpha}{2} \Gamma(r + a_1)$$

The solution of above two equations utilizing proper numerical approaches enables the derivation of Bayesian CIs for λ .

4.2. Case II: Shape Parameter Unknown

The conditional density of β given λ and data is:

$$\begin{aligned} \pi_2(\beta|\lambda, data) &= \int_0^\infty \pi(\beta, \lambda|data) d\lambda \\ &= \beta^{r+a_2-1} e^{-b_2\beta} \prod_{i=1}^r \frac{x_i^{\beta-1}}{(1-x_i)^{\beta+1}} \cdot \frac{\Gamma(r+a_1)}{\left(\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1\right)^{r+a_1}} \end{aligned}$$

Therefore,

$$\pi_2(\beta|\lambda, data) \propto \beta^{r+a_2-1} e^{-b_2\beta} \prod_{i=1}^r \frac{x_i^{\beta-1}}{(1-x_i)^{\beta+1}} \cdot \frac{1}{\left(\sum_{i=1}^r \left(\frac{x_i}{1-x_i}\right)^\beta + (n-r) \left(\frac{x_r}{1-x_r}\right)^\beta + b_1\right)^{r+a_1}} \quad (4.2)$$

There exists no explicit solution for $\pi_2(\beta|\lambda, data)$ so we differentiate Eq.(8) twice with respect to β after applying natural logarithms to the both sides of $\pi_2(\beta|\lambda, data)$ to obtain $\frac{\partial^2}{\partial \beta^2} \ln \pi_2(\beta|\lambda, data) < 0$. Since $\pi_2(\beta|\lambda, data)$ is Log-Concave we apply Deveroye's (1984) method to generate samples of β from the density function $\pi_2(\beta|\lambda, data)$ before using them to estimate BEs of β and λ under both square error and LINEX loss functions.

We propose using Gibbs sampling to generate samples of parameters $\{(\beta_j, \lambda_j); j = 1, 2, \dots, M\}$ which we will use to obtain BEs as well as to form CIs.

Algorithm (1):

1. Generate β from the Log-Concave density function $\pi_2(\beta|\lambda, data)$.
2. Produce a λ value after sampling from the conditional distribution $\pi_1(\lambda|\beta, data)$ at every instance of β .
3. Repeat step 1 and step 2, M times to obtain Markov Chain Monte Carlo (MCMC) samples $\{(\beta_j, \lambda_j); j = 1, 2, \dots, M\}$.
4. Obtain the BEs of β and λ under square error loss function as:

$$\begin{aligned} \hat{\beta}_{Bayes, S} &= \frac{1}{M} \sum_{i=1}^M \beta_i, \\ \hat{\lambda}_{Bayes, S} &= \frac{1}{M} \sum_{i=1}^M \lambda_i \end{aligned}$$

and

$$\begin{aligned}\hat{Var}(\beta|data) &= \frac{1}{M} \sum_{i=1}^M \left(\beta_i - \hat{\beta}_{Bayes,S} \right)^2, \\ \hat{Var}(\lambda|data) &= \frac{1}{M} \sum_{i=1}^M \left(\lambda_i - \hat{\lambda}_{Bayes,S} \right)^2\end{aligned}$$

5. Obtain the BEs of β and λ under LINEX loss function as:

$$\begin{aligned}\hat{\beta}_{Bayes,L} &= \left[\frac{1}{M} \sum_{i=1}^M \frac{1}{\beta_i^{a^*}} \right]^{-\frac{1}{a^*}}, \\ \hat{\lambda}_{Bayes,L} &= \left[\frac{1}{M} \sum_{i=1}^M \frac{1}{\lambda_i^{a^*}} \right]^{-\frac{1}{a^*}}\end{aligned}$$

and

$$\begin{aligned}\hat{Var}(\beta|data) &= \frac{1}{M} \sum_{i=1}^M \left(\beta_i - \hat{\beta}_{Bayes,L} \right)^2, \\ \hat{Var}(\lambda|data) &= \frac{1}{M} \sum_{i=1}^M \left(\lambda_i - \hat{\lambda}_{Bayes,L} \right)^2\end{aligned}$$

6. To compute the credible interval, we order:

$$\lambda_1, \lambda_2, \dots, \lambda_M \text{ as } \lambda_{(1)} < \lambda_{(2)} < \dots < \lambda_{(M)}$$

and

$$\beta_1, \beta_2, \dots, \beta_M \text{ as } \beta_{(1)} < \beta_{(2)} < \dots < \beta_{(M)}$$

and then, the $(1 - \alpha)100\%$ symmetric credible intervals (CIs) of β and λ respectively are:

$$\left[\beta_{[M\frac{\alpha}{2}]}, \beta_{[M(1-\frac{\alpha}{2})]} \right] \text{ and } \left[\lambda_{[M\frac{\alpha}{2}]}, \lambda_{[M(1-\frac{\alpha}{2})]} \right]$$

5. BAYES PREDICTION AND PREDICTIVE INTERVALS

An essential Bayes analytical component makes use of prediction methods for future sample observations based on current “informative sample” data collection. Our main assessment revolves around determining the posterior predictive densities for future observations based on the present data collection. We produce future experimental observations through predictions made by analyzing results obtained from an informative experiment.

The observed data sequence $x_{1:n} < x_{2:n} < \dots < x_{r:n}$ serves as a known informative sample. We aim to predict values from $X_{s:n}$ where $r < s \leq n$. Posterior predictive density of $X_{s:n}$ given observed data $\underline{X} = (x_{1:n}, x_{2:n}, \dots, x_{r:n})$ appears as follows:

$$\pi_{X_s}(x|data) = \int_0^\infty \int_0^\infty h_{X_s|data}(x|\beta, \lambda) \pi(\beta, \lambda|data) d\beta d\lambda, x_s > x_r,$$

where $h_{X_s|data}(x|\beta, \lambda)$ is the conditional density of X_s given β, λ and data $\underline{X}$, see for example Chen, Shao and Ibrahim (2000).

The Markovian property of conditional order statistics according to David and Nagaraja (2003) allows us to determine that the conditional PDF of $X_{s:n}$ given $\underline{X}$ reduces to the conditional PDF of $X_{s:n}$ given $X_{r:n}$ such that $r+1 \leq s \leq n$.

$$h_{X_s|data}(x|\beta, \lambda) = h_{X_s|X_r}(x|\beta, \lambda) = \frac{h_{r,s:n}(x_r, x_s)}{h_{r:n}(x_r)} \quad (5.1)$$

The joint PDF of these order statistics consists of $h_{r,s:n}(x_r, x_s)$ components from n sampled values drawn from $G(\cdot)$. The conditional density of $X_{s:n}$ under the condition that $X_{r:n}$ occurs simply matches the marginal density of the $(s-r)$ th order statistic derived from a sample of $(n-r)$ elements drawn from a left-truncated version of $G(\cdot)$ starting from x_r .

Using Eq.(1), Eq.(2) and Eq.(9), we get:

$$h_{X_s|data}(x|\beta, \lambda) = c_2 \left[1 - e^{-\lambda \left[\left(\frac{x}{1-x} \right)^\beta - \left(\frac{x_r}{1-x_r} \right)^\beta \right]} \right]^{s-r-1} \times e^{-\lambda(n-s+1) \left[\left(\frac{x}{1-x} \right)^\beta - \left(\frac{x_r}{1-x_r} \right)^\beta \right]} \frac{\beta \lambda x^{\beta-1}}{(1-x)^{\beta+1}}$$

where $c_2 = \frac{(n-r)!}{(s-r-1)!(n-s)!}$.

By using the binomial expansion, we have:

$$h_{X_s|data}(x|\beta, \lambda) = c_2 \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i e^{-\lambda(n-s+i+1) \left[\left(\frac{x}{1-x} \right)^\beta - \left(\frac{x_r}{1-x_r} \right)^\beta \right]} \frac{\beta \lambda x^{\beta-1}}{(1-x)^{\beta+1}}$$

So, the posterior predictive density of $X_{s:n}$ at any point $x > x_r$ is:

$$\pi_{X_s}(x|data) = c_2 \int_0^\infty \int_0^\infty \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i e^{-\lambda(n-s+i+1) \left[\left(\frac{x}{1-x} \right)^\beta - \left(\frac{x_r}{1-x_r} \right)^\beta \right]} \frac{\beta \lambda x^{\beta-1}}{(1-x)^{\beta+1}} \times \pi(\beta, \lambda|data) d\beta d\lambda$$

Posterior predictive density function exists in the above format yet it fails to achieve attractability status and therefore prevents the derivation of an explicit predictive Bayes estimate. The Bayes predictor (BP) determines the distribution of $X_{s:n}$ when using an SEL function according to:

$$\begin{aligned} X_{s:n}^{BP} &= E_{\pi_{X_s}}(x|data) \\ &= c_2 \int_{x_r}^1 x \int_0^\infty \int_0^\infty \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i \\ &\quad \times e^{-\lambda(n-s+i+1) \left[\left(\frac{x}{1-x} \right)^\beta - \left(\frac{x_r}{1-x_r} \right)^\beta \right]} \frac{x^{\beta-1}}{(1-x)^{\beta+1}} \pi(\lambda, \mu|data) d\beta d\lambda dx. \end{aligned} \quad (5.2)$$

The Gibbs sampling method produces MC samples $\{(\beta_j, \lambda_j); j = 1, 2, \dots, M\}$ to determine the simulation Bayes predictor

$$\begin{aligned} \hat{X}_{s:n}^{BP} &= \frac{c_2}{M} \sum_{j=1}^M \beta_j \lambda_j \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i e^{\lambda_j(n-s+i+1) \left(\frac{x_r}{1-x_r} \right)^\beta} \\ &\quad \times \int_{x_r}^1 \frac{x \cdot x^{\beta_j}}{(1-x)^{\beta_j+1}} e^{-\lambda_j(n-s+i+1) \left(\frac{x}{1-x} \right)^\beta} dx \end{aligned}$$

By applying transformation $z = \left(\frac{x}{1-x}\right)^{\beta_j}$ and incomplete gamma function definition as $\Gamma(a, c) = \int_c^\infty x^{a-1} e^{-x} dx$, $a > 0, c > 0$ the Bayes predictor computes $X_{s:n}$ as follows:

$$\begin{aligned} \hat{X}_{s:n}^{BP} &= \frac{c_2}{M} \sum_{j=1}^M \lambda_j \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i e^{\lambda_j(n-s+i+1)\left(\frac{x_r}{1-x_r}\right)^{\beta_j}} \\ &\times \sum_{k=0}^{\infty} (-1)^k \frac{\Gamma\left(\frac{\beta_j+k+1}{\beta_j}, \lambda_j(n-s+i+1)\left(\frac{x_r}{1-x_r}\right)^{\beta_j}\right)}{[\lambda_j(n-s+i+1)]^{\frac{\beta_j+k+1}{\beta_j}}} \end{aligned} \quad (5.3)$$

Algorithm (2):

1. Generate MCMC sampling $\{(\beta_j, \lambda_j); j = 1, 2, \dots, M\}$.
2. Compute $\hat{X}_{s:n}^{BP}$ based on β_j and λ_j obtained in step 1.
3. Repeat step 1 and step 2, M times to obtain $\{\hat{X}_1, \hat{X}_2, \dots, \hat{X}_M\}$.
4. Compute mean square error (MSE) as follows:

$$MSE(\hat{X}_{s:n}) = \frac{1}{M} \sum_{i=1}^M \left(\hat{X}_i - \text{Average} \hat{X} \right)^2, \text{ where } \text{Average} \hat{X} = \frac{1}{M} \sum_{i=1}^M \hat{X}_i$$

Research seeks to develop dual prediction intervals for order statistics X_s . To achieve this we require the predictive survival function of X_s defined as:

$$\begin{aligned} S_{X_s|data}(x|\beta, \lambda) &= Pr(X > x) = \int_x^1 h_{X_s|data}(z|\beta, \lambda) dz \\ &= c_2 \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} (-1)^i \int_x^1 e^{-\lambda(n-s+i+1)\left[\left(\frac{z}{1-z}\right)^\beta - \left(\frac{x_r}{1-x_r}\right)^\beta\right]} \frac{z^{\beta-1}}{(1-z)^{\beta+1}} dz. \end{aligned}$$

Using the transformation $z = \left(\frac{x}{1-x}\right)^\beta$ we get:

$$S_{X_s|data}(x|\beta, \lambda) = c_2 \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} \frac{(-1)^i}{(n-s+i+1)} e^{-\lambda(n-s+i+1)\left[\left(\frac{x}{1-x}\right)^\beta - \left(\frac{x_r}{1-x_r}\right)^\beta\right]}$$

The predictive survival function for X_s follows the SEL function.

$$\begin{aligned} S_{X_s|data}^P(x|\beta, \lambda) &= \int_0^\infty \int_0^\infty c_2 \sum_{i=0}^{s-r-1} \binom{s-r-1}{i} \frac{(-1)^i}{(n-s+i+1)} e^{-\lambda(n-s+i+1)\left[\left(\frac{x}{1-x}\right)^\beta - \left(\frac{x_r}{1-x_r}\right)^\beta\right]} \\ &\pi(\beta, \lambda|data) d\beta d\lambda. \end{aligned}$$

The collected MCMC samples $\{(\beta_j, \lambda_j); j = 1, 2, \dots, M\}$ from the Gibbs sampling lead to the following simulation estimator of the predictive survival function

$$\hat{S}_{X_s|data}^P(x|\beta, \lambda) = c_2 \sum_{j=1}^M \left[\sum_{i=0}^{s-r-1} \binom{s-r-1}{i} \frac{(-1)^i}{(n-s+i+1)} e^{-\lambda_j(n-s+i+1)\left[\left(\frac{x}{1-x}\right)^{\beta_j} - \left(\frac{x_r}{1-x_r}\right)^{\beta_j}\right]} \right].$$

Using a suitable numerical technique solve the following non-linear equations to find the $(1-\alpha)100\%$ predictive interval of X_s for both the lower bound (L) and the upper bound (U).

$$\hat{S}_{X_s|data}^P(L) = 1 - \frac{\alpha}{2} \quad \text{and} \quad \hat{S}_{X_s|data}^P(U) = \frac{\alpha}{2} \quad (5.4)$$

6. SIMULATIONS AND DATA ANALYSIS

6.1. Simulations

In this section, we conduct a series of numerical experiments to evaluate the performance of the proposed methods under different sampling schemes and priors using type II censored data. The data is generated from Kies distribution by assuming $\beta = 2$, $\lambda = 1$. When computing Bayesian Estimates (BEs) using the square error loss function and LINEX loss function, two types of priors are considered for both variables β and λ . The two priors considered are Prior 0 and Prior 1.

Prior 0 is a non-informative prior with $a_1 = b_1 = a_2 = b_2 = 0$. Prior 1, on the other hand, is an informative prior with $a_1 = 2, b_1 = 1, a_2 = 1, b_2 = 1$. The Maximum Likelihood Estimates (MLEs) and the Bayesian Estimates (BEs) based on square error loss function and LINEX loss function are calculated for each sampling scheme. Addition, 1000 Monte Carlo (MC) samples are used to produce 95% credible intervals. This simulation study presents the average Bayes estimates, mean squared errors (MSEs), and coverage percentages lengths that are calculated from 1000 replications. Algorithm (1) is used to calculate the numerical results of the Bayesian estimates for both parameters, as well as their corresponding mean squared errors (MSEs). Table 1 and table 2 show that the maximum likelihood estimators (MLEs) and the Bayesian estimates with respect to square error loss function and LINEX loss function and show improvement as the available information (as r becomes larger). The aforementioned observation holds for both Prior 0 and Prior 1. Addition—it is worth noting that the Bayesian estimators perform commendably compared to the maximum likelihood estimators of parameter and (as shown in Table 1) for both priors. when comparing the Bayes estimates (BEs) obtained under Prior 0 and Prior 1 with respect two loss functions, it is evident that the BEs obtained with Prior 1 (an informative prior) perform better than those obtained with Prior 0 (a non-informative prior), as measured by the mean squared errors (MSEs). Furthermore, it is noteworthy that all BEs associated with the squared error loss function and LINEX loss function for both β and λ fall within their respective credible intervals (CLs). Table 3 shows that the average length of credible intervals decreases as r increases while holding n constant for both β and λ .

In order to calculate the predictors, we only consider the prior 1 and the square error loss function. Using type II censored data, we have derive point predictors by algorithm (2) and 95% prediction intervals (PIs) by using Eq.(10) for the absent order statistics $X_{s:n}$, $r + 1 \leq s \leq n$. We then simulate the Bayesian predictors for the absent order statistics $X_{s:n}$, $r + 1 \leq s \leq n$, based on MC samples $\{(\beta_i, \lambda_i), i = 1, 2, \dots, M\}$, where $M = 1000$. Table 4 shows that the expected values for the absent order statistics $X_{s:n}$ are very close to each other and remain within their respective prediction intervals across all schemes.

6.2. Data analysis

We analyze the real lifetimes data of size $n = 40$ from Xu et al.(2003), which represents the time to failure (10^3 h) of turbocharger of one type of engine. The data are: 1.6, 3.5, 4.8, 5.4, 6.0, 6.5, 7.0, 7.3, 7.7, 8.0, 8.4, 2.0, 3.9, 5.0, 5.6, 6.1, 6.5, 7.1, 7.3, 7.8, 8.1, 8.4, 2.6, 4.5, 5.1, 5.8, 6.3, 6.7, 7.3, 7.7, 7.9, 8.3, 8.5, 3.0, 4.6, 5.3, 6.0, 8.7, 8.8, 9.0.

Table 1: MLEs, Bayes estimates and credible interval based on type II censored data with respect to square error loss function, when Prior 0 and Prior 1 are used.

Scheme	MLE		Bayes estimates (Prior 0)		Bayes estimates (Prior 1)	
	β	λ	β	λ	β	λ
Scheme 1:	2.4602	1.0996	2.2075	1.2742	2.2314	1.1330
$n = 25, r = 10$	(1.2358)	(1.1578)	(0.8540)	(0.9069)	(0.3656)	(0.4029)
95% CIs			(1.0126; 3.6495)	(0.2717; 3.0619)	(1.2711; 3.2543)	(0.3389; 2.0300)
Scheme 2:	2.2967	1.0030	2.0380	1.0568	2.0578	0.9516
$n = 25, r = 15$	(0.4792)	(0.1676)	(0.1880)	(0.1485)	(0.1631)	(0.0926)
95% CIs			(1.1679; 3.0267)	(0.3871; 1.9601)	(1.2805; 2.9153)	(0.4190; 1.7777)
Scheme 3:	2.1224	1.0448	2.0480	1.0113	2.0254	1.0260
$n = 25, r = 20$	(0.1494)	(0.1287)	(0.1156)	(0.1264)	(0.1019)	(0.0723)
95% CIs			(1.3425; 2.8188)	(0.4625; 1.6429)	(1.3467; 2.7270)	(0.4699; 1.5992)
Scheme 4:	2.2774	1.1608	2.1933	1.0929	2.1406	0.9982
$n = 40, r = 20$	(0.4064)	(0.1928)	(0.3134)	(0.1802)	(0.2432)	(0.1559)
95% CIs			(1.2916; 2.9254)	(0.4748; 1.8264)	(1.3696; 2.8464)	(0.4579; 1.7499)
Scheme 5:	2.1679	1.0640	2.1089	1.0222	2.0652	1.0279
$n = 40, r = 25$	(0.1897)	(0.1424)	(0.1500)	(0.0988)	(0.1143)	(0.0873)
95% CIs			(1.3536; 2.7422)	(0.5083; 1.6856)	(1.3864; 2.7194)	(0.5049; 1.5942)
Scheme 6:	2.1101	1.0543	2.0808	1.0152	2.0863	1.0080
$n = 40, r = 30$	(0.1423)	(0.0356)	(0.1209)	(0.0585)	(0.0971)	(0.0730)
95% CIs			(1.4144; 2.6389)	(0.5350; 1.5032)	(1.4402; 2.5744)	(0.5509; 1.4937)

Table 2: Bayes estimates and credible interval based on type II censored data with respect to LINEX loss function, when Prior 0 and Prior 1 are used.

Scheme	a^*	Bayes estimates (Prior 0)		Bayes estimates (Prior 1)	
		β	λ	β	λ
Scheme 1: $n = 25, r = 10$	0.1	2.1931 (0.4732)	1.0607 (0.2979)	2.1092 (0.2504)	1.0593 (0.1355)
		95% CIs (1.2231; 3.8731)	(0.4673; 2.5242)	(1.2058; 3.2341)	(0.5003; 1.8947)
	1.0	2.1922 (0.4733)	1.0035 (0.2661)	2.1087 (0.2504)	1.0082 (0.1239)
		95% CIs (1.2215; 3.8719)	(0.4429; 2.3644)	(1.2053; 3.2333)	(0.4769; 1.8132)
	3.0	2.1900 (0.4736)	0.8760 (0.2014)	2.1074 (0.2505)	0.8894 (0.1032)
		95% CIs (1.2173; 3.8690)	(0.3918; 1.9967)	(1.2044; 3.2317)	(0.4250; 1.6538)
Scheme 2: $n = 25, r = 15$	0.1	2.0977 (0.3253)	1.0076 (0.1593)	2.0841 (0.1804)	0.9736 (0.0578)
		95% CIs (1.3883; 3.6002)	(0.5911; 2.1491)	(1.3460; 2.9860)	(0.6491; 1.6128)
	1.0	2.0973 (0.3253)	0.9740 (0.1486)	2.0838 (0.1804)	0.9437 (0.0544)
		95% CIs (1.3885; 3.5997)	(0.5725; 2.0711)	(1.3457; 2.9857)	(0.6296; 1.5600)
	3.0	2.0965 (0.3254)	0.8982 (0.1268)	2.0832 (0.1505)	0.8769 (0.0474)
		95% CIs (1.3878; 3.5985)	(0.5242; 1.8719)	(1.3451; 2.9851)	(0.5874; 1.4594)
Scheme 3: $n = 25, r = 20$	0.1	2.0821 (0.1603)	0.9914 (0.0498)	1.9860 (0.1305)	0.9531 (0.0455)
		95% CIs (1.4206; 3.001)	(0.5929; 1.4919)	(1.3342; 2.7638)	(0.6255; 1.4171)
	1.0	2.0919 (0.1603)	0.9675 (0.0474)	1.9858 (0.1305)	0.9313 (0.0435)
		95% CIs (1.4204; 3.0009)	(0.5782; 1.4607)	(1.3341; 2.7636)	(0.6120; 1.3764)
	3.0	2.0815 (0.1603)	0.9148 (0.0423)	1.9854 (0.1305)	0.8827 (0.0392)
		95% CIs (1.4198; 3.0003)	(0.5464; 1.3949)	(1.3335; 2.7634)	(0.5809; 1.2964)
Scheme 4: $n = 40, r = 20$	0.1	2.0560 (0.2039)	1.0633 (0.1604)	2.0803 (0.1394)	0.9961 (0.0696)
		95% CIs (1.2822; 3.0161)	(0.6558; 1.9932)	(1.4947; 2.9992)	(0.6326; 1.6825)
	1.0	2.0558 (0.2039)	1.0375 (0.1532)	2.0802 (0.1394)	0.9728 (0.0663)
		95% CIs (1.2819; 3.0180))	(0.6398; 1.9471)	(1.4945; 2.9991)	(0.6159; 1.6389)
	3.0	2.0556 (0.2040)	0.9799 (0.13875)	2.0799 (0.1394)	0.9219 (0.0594)
		95% CIs (1.2812; 3.0179)	(0.5959; 1.8443)	(1.4942; 2.9989)	(0.5815; 1.5201)
Scheme 5: $n = 40, r = 25$	0.1	2.0628 (0.1358)	1.0277 (0.0557)	2.0576 (0.1111)	1.0104 (0.0473)
		95% CIs (1.3797; 2.8294)	(0.6695; 1.6729)	(1.4985; 2.7977)	(0.6766; 1.3765)
	1.0	2.0627 (0.1359)	1.0082 (0.0538)	2.0576 (0.1111)	0.9921 (0.0456)
		95% CIs (1.3796; 2.8294)	(0.6586; 1.6362)	(1.4984; 2.7977)	(0.6632; 1.3502)
	3.0	2.0578 (0.1111)	0.9514 (0.0420)	2.0575 (0.1111)	0.9514 (0.0420)
		95% CIs (1.4983; 2.7976)	(0.6332; 1.2868)	(1.4983; 2.7976)	(0.6332; 1.2867)
Scheme 6: $n = 40, r = 30$	0.1	2.0406 (0.1017)	0.9693 (0.0394)	2.0264 (0.0822)	0.9902 (0.0333)
		95% CIs (1.4879; 2.6313)	(0.6311; 1.4775)	(1.4923; 2.5491)	(0.6666; 1.3279)
	1.0	2.0405 (0.1016)	0.9544 (0.0382)	2.0263 (0.0822)	0.9753 (0.0324)
		95% CIs (1.4878; 2.6312)	(0.6211; 1.4574)	(1.4923; 2.5490)	(0.6552; 1.2996)
	3.0	2.0404 (0.1017)	0.9214 (0.0355)	2.0262 (0.0822)	0.9422 (0.0304)
		95% CIs (1.4876; 2.6311)	(0.5945; 1.4133)	(1.4922; 2.5488)	(0.6316; 1.2500)

Table 3: Average credible intervals lengths based on type II censored data and the coverage percentages, when Prior 0 and Prior 1 are used.

Scheme	Bayes (Prior0)		Bayes (Prior1)	
	β	λ	β	λ
Scheme 1: $n = 25; r = 10$	2.4573 (0.87)	2.4885 (0.90)	1.8097 (0.91)	1.5425 (0.85)
Scheme 2: $n = 25; r = 15$	1.9831 (0.97)	1.6910 (0.94)	1.7587 (0.90)	1.4729 (0.97)
Scheme 3: $n = 25; r = 20$	1.5587 (0.96)	1.1519 (0.94)	1.3563 (0.94)	1.1815 (0.94)
Scheme 4: $n = 40; r = 20$	1.5772 (0.92)	1.3625 (0.91)	1.3803 (0.88)	1.3293 (0.91)
Scheme 5: $n = 40; r = 25$	1.4091 (0.93)	1.1395 (0.95)	1.2768 (0.95)	1.2928 (0.92)
Scheme 6: $n = 40; r = 30$	1.3330 (0.91)	1.0893 (0.95)	1.1528 (0.93)	0.9358 (0.90)

Table 4: Point predictors and PIs for the missing order statistics $X_{s:n}, r+1 \leq s \leq n$.

Scheme 1: $n = 25, r = 10$	Predicted Values	MSE	95% PI
$X_{11:n}$	0.4203	0.0248	(0.4040; 0.4601)
$X_{12:n}$	0.4361	0.0368	(0.4081; 0.4838)
$X_{13:n}$	0.4512	0.0426	(0.4152; 0.5032)
$X_{14:n}$	0.4655	0.0482	(0.4239; 0.5200)
$X_{15:n}$	0.4798	0.0542	(0.4338; 0.5362)
$X_{16:n}$	0.4933	0.0604	(0.4442; 0.5508)
$X_{17:n}$	0.5068	0.0672	(0.4552; 0.5652)
$X_{18:n}$	0.5199	0.0743	(0.4668; 0.5797)
$X_{19:n}$	0.5344	0.0829	(0.4787; 0.5941)
$X_{20:n}$	0.5483	0.0919	(0.4913; 0.6086)
$X_{21:n}$	0.5637	0.1027	(0.5049; 0.6246)
$X_{22:n}$	0.5796	0.1147	(0.5192; 0.6414)
$X_{23:n}$	0.5974	0.1293	(0.5351; 0.6606)
$X_{24:n}$	0.6196	0.1473	(0.5540; 0.6856)
$X_{25:n}$	0.6524	0.1403	(0.5802; 0.7253)

Table 5: Point predictors and PIs for the missing order statistics $X_{s:n}$, $r + 1 \leq s \leq n$.

Scheme 2: $n = 25, r = 15$	Predicted Values	MSE	95% PI
$X_{16:n}$	0.4958	0.0404	(0.4818; 0.5306)
$X_{17:n}$	0.5099	0.0657	(0.4853; 0.5531)
$X_{18:n}$	0.5242	0.0757	(0.4916; 0.5726)
$X_{19:n}$	0.5388	0.0849	(0.4998; 0.5909)
$X_{20:n}$	0.5538	0.0949	(0.5096; 0.6091)
$X_{21:n}$	0.5694	0.1060	(0.5207; 0.6265)
$X_{22:n}$	0.5867	0.1193	(0.5335; 0.6463)
$X_{23:n}$	0.6058	0.1349	(0.5483; 0.6682)
$X_{24:n}$	0.6288	0.1517	(0.5662; 0.6949)
$X_{25:n}$	0.6623	0.1283	(0.5908; 0.7364)

Table 6: Point predictors and PIs for the missing order statistics $X_{s:n}$, $r + 1 \leq s \leq n$.

Scheme 3: $n = 25, r = 20$	Predicted Values	MSE	95% PI
$X_{21:n}$	0.5657	0.0426	(0.5501; 0.6038)
$X_{22:n}$	0.5831	0.0955	(0.5544; 0.6310)
$X_{23:n}$	0.6028	0.1169	(0.5630; 0.6580)
$X_{24:n}$	0.6263	0.1271	(0.5760; 0.6885)
$X_{25:n}$	0.6602	0.0789	(0.5961; 0.7332)

Before we analyzing the data, we divide each data value by 10. The well-known Kolomogrov-Smirnov (K-S) goodness of fit test is used to test whether the Kies distribution adequately fits this data set or not, we compuetd the MLE estimatores of β and λ and they are: 1.2740 and 0.2705 respectively, then the corresponding (K-S) distance become 0.0765 and the corresponding P-value is 0.9449. Therefore, we can not reject the hypothesis that the data comes from Kies distribution, so we can use it to analyze the real life turbocharger data set.

First, under the scheme: $n = 40; r = 25$, we compute the MLEs of β and λ and they are 1.3335 and 0.2691, respectively. Second, we compute the Bayes estimates and the 95% cerdible intervals of β and λ with respect to square error loss function and LINEX loss function under prior 0 and the results are represented in Table 5.

We now consider the prediction of the missing order statistics under the scheme: $n = 40; r = 25$. The predicted values and the 95% PI of the order statistics are shown in Table 6 when Prior 1 is used. It is observed that all predicted values with respect to square error loss function are all ordered and fall in their corresponding PI.

7. CONCLUSION

In this paper, classical and Bayesian estimation are proposed for the two parameter Kies distribution based on type II censored data. Gibbs sampling technique is used to generate samples for computing

Table 7: Point predictors and PIs for the missing order statistics $X_{s:n}$, $r + 1 \leq s \leq n$ for Scheme 4.

Scheme 4: $n = 40, r = 20$	Predicted Values	MSE	95% PI
$X_{21:n}$	0.4550	0.0356	(0.4461; 0.4781)
$X_{22:n}$	0.4638	0.0465	(0.4482; 0.4919)
$X_{23:n}$	0.4726	0.0506	(0.4519; 0.5056)
$X_{24:n}$	0.4812	0.0543	(0.4565; 0.5170)
$X_{25:n}$	0.4897	0.0582	(0.4617; 0.5274)
$X_{26:n}$	0.4983	0.0623	(0.4675; 0.5379)
$X_{27:n}$	0.5066	0.0666	(0.4735; 0.5477)
$X_{28:n}$	0.5151	0.0711	(0.4799; 0.5574)
$X_{29:n}$	0.5231	0.0757	(0.4865; 0.5665)
$X_{30:n}$	0.5321	0.0810	(0.4935; 0.5764)
$X_{31:n}$	0.5409	0.0865	(0.5007; 0.5863)
$X_{32:n}$	0.5485	0.0919	(0.5083; 0.5961)
$X_{33:n}$	0.5462	0.0972	(0.5162; 0.6055)
$X_{34:n}$	0.4873	0.1037	(0.5248; 0.6165)
$X_{35:n}$	0.5780	0.1131	(0.5336; 0.6276)
$X_{36:n}$	0.5904	0.1230	(0.5433; 0.6400)
$X_{37:n}$	0.6029	0.1338	(0.5541; 0.6536)
$X_{38:n}$	0.6170	0.1467	(0.5660; 0.6697)
$X_{39:n}$	0.6428	0.1826	(0.5809; 0.6913)
$X_{40:n}$	0.6623	0.1593	(0.6012; 0.7264)

Table 8: Point predictors and PIs for the missing order statistics $X_{s:n}$, $r + 1 \leq s \leq n$.

Scheme 5: $n = 40, r = 25$	Predicted Values	MSE	95% PI
$X_{26:n}$	0.4973	0.0471	(0.4886; 0.5197)
$X_{27:n}$	0.5056	0.0650	(0.4906; 0.5335)
$X_{28:n}$	0.5146	0.0708	(0.4943; 0.5466)
$X_{29:n}$	0.5231	0.0757	(0.4989; 0.5578)
$X_{30:n}$	0.5323	0.0812	(0.5044; 0.5695)
$X_{31:n}$	0.5398	0.0859	(0.5097; 0.5787)
$X_{32:n}$	0.5500	0.0925	(0.5169; 0.5906)
$X_{33:n}$	0.5595	0.0991	(0.5237; 0.6019)
$X_{34:n}$	0.5699	0.1067	(0.5316; 0.6143)
$X_{35:n}$	0.5789	0.1135	(0.5385; 0.6260)
$X_{36:n}$	0.5909	0.1234	(0.5479; 0.6393)
$X_{37:n}$	0.6050	0.1354	(0.5595; 0.6544)
$X_{38:n}$	0.6207	0.1498	(0.5720; 0.6732)
$X_{39:n}$	0.6421	0.1690	(0.5895; 0.6971)
$X_{40:n}$	0.6666	0.1520	(0.6061; 0.7309)

Table 9: Point predictors and PIs for the missing order statistics $X_{s:n}, r+1 \leq s \leq n$.

Scheme 6: $n = 40, r = 30$	Predicted Values	MSE	95% PI
$X_{31:n}$	0.5444	0.0584	(0.5358; 0.5667)
$X_{32:n}$	0.5535	0.0908	(0.5380; 0.5821)
$X_{33:n}$	0.5630	0.1005	(0.5418; 0.5960)
$X_{34:n}$	0.5727	0.1083	(0.5470; 0.6091)
$X_{35:n}$	0.5831	0.1165	(0.5532; 0.6225)
$X_{36:n}$	0.5945	0.1258	(0.5605; 0.6367)
$X_{37:n}$	0.6071	0.1366	(0.5691; 0.6524)
$X_{38:n}$	0.6216	0.1492	(0.5794; 0.6704)
$X_{39:n}$	0.6399	0.1622	(0.5923; 0.6933)
$X_{40:n}$	0.6672	0.1316	(0.6104; 0.7301)

Table 10: MLEs, Bayes estimates and credible interval based on type II censored turbocharger real life data.

MLE		Bayes estimates		
β λ		with respect to squar error loss function		
		β	λ	
Scheme: $n = 40, r = 25$	1.3335	0.2691	1.2735	0.2715
95% CIs			(1.2342; 1.2916)	(0.1709; 0.3934)
with respect to LINEX loss function				
		a^*	β	λ
		0.1	1.4339	0.2608
95%CIs			(1.4323; 1.4377)	(0.2541; 0.2681)
		1.0	1.4338	0.2544
95%CIs			(1.4322; 1.4377)	(0.2472; 0.2613)
		3.0	1.4336	0.2400
95%CIs			(1.4319; 1.4375)	(0.2299; 0.2478)

Table 11: Point predictors and 95% PIs for the missing real life turbocharger data.

Scheme: $n = 40, r = 25$	Predicted Values	95% PI
$X_{26:n}$	7.4008	(7.3026; 7.6576)
$X_{27:n}$	7.4996	(7.3259; 7.8222)
$X_{28:n}$	7.6054	(7.3688; 7.9748)
$X_{29:n}$	7.7060	(7.4217; 8.1217)
$X_{30:n}$	7.7932	(7.4810; 8.2009)
$X_{31:n}$	7.8978	(7.5558; 8.3140)
$X_{32:n}$	7.9967	(7.6152; 8.4675)
$X_{33:n}$	8.0928	(7.7010; 8.5349)
$X_{34:n}$	8.1826	(7.7763; 8.6306)
$X_{35:n}$	8.2867	(7.8632; 8.7375)
$X_{36:n}$	8.3973	(7.9657; 8.8434)
$X_{37:n}$	8.4773	(8.0358; 8.9062)
$X_{38:n}$	8.6057	(8.1516; 9.0385)
$X_{39:n}$	8.7466	(8.2663; 9.1879)
$X_{40:n}$	8.9473	(8.4665; 9.3686)

the Bayes estimates and to construct credible intervals. In addition, the posterior predictive density of a future observation based on the current data is estimated to predict the missing order data. The behavior of the proposed methods for different sampling schemes and priors are listed. In terms of the MSE, the scale and shape parameters estimated from the Bayesian method are better than the ones from MLE. Moreover, the Bayes estimate are significant under prior one. Furthermore, when the observed data are increased and the sample size is fixed, the average length of the credible intervals is decreased. Finally, the Bayesian predicted values are found to increase and to lie within the corresponding predictive intervals.

Conflicts of Interest: The authors declare that they have no conflict of interest.

RECEIVED: JULY, 2023.

REVISED: FEBRUARY, 2024.

REFERENCES

- [1] AL-HUSSAINI, E.K. (1999). Predicting Observable from a General Class of Distributions, **Journal of Statistical Planning and Inference** 79, 79-81.
- [2] BALAKRISHNAN, N. and KATERI, M. (2008). On the maximum likelihood estimation of parameters of Weibull distribution based on complete and censored data, **Statistics & Probability Letters** 78(17), 2971-2975.

- [3] CHEN, M.H., SHAO, Q.M. and ABRAHIM, J.G. (2000). **Monte Carlo Methods in Bayesian Computation**, Springer-Verlag, New York.
- [4] DAVID, H.A. and NAGARAJA, H.N. (2003). **Order Statistics**, 3rd Edition, Wiley, New York.
- [5] DEVROYE, L. (1984). A Simple Algorithm for Generating Random Variates with a Log-Concave Density, **Computing** 33, 247-257.
- [6] GHASSAN and el at. (2016). Statistical Inference of the Rest Lifetimes of Ball Bearings Residual Lifetimes Based on Rayleigh Type II Censored Data, **Journal of Mathematics and Statistics** 12(3):182-191.
- [7] KIES, J.A. (1958). The strength of glass performance, **Naval Research Lab Report No.5093**, Washington, D.C.
- [8] KUMAR, C.S. and DHARMAJA, S.H.S. (2013). On reduced Kies distribution, **Collection of Recent Statistical Method and Applications**, 111-123, Department of Statistics, University of Kerala Publishers, Trivandrum.
- [9] KUMAR, C.S. and DHARMAJA, S.H.S. (2014). On some properties of Kies distribution, **Metron**, 72,(1):97-122.
- [10] KUMAR, C.S. and DHARMAJA, S.H.S. (2017). On Modified Kies Distribution and its Applications, **Journal of Statistical Research**, 51(1):41-60.
- [11] KUMAR, C.S. and DHARMAJA, S.H.S. (2017). The Exponentiated Reduced Kies Distribution: Properties and Applications, **Communications in Statistics-Theory and Methods**, 46(17):8778-8790. *Envirometrics* 15, 701-710.
- [12] KUNDU, D. and HOWLADER, H. (2010). Bayesian Inference and Prediction of the Inverse Weibull Distribution for Type-II Censored Data, **Computational Statistics and Data Analysis** 54, 1547-1558.
- [13] KUNDU, D. and RAQAB, M.Z. (2012). Bayesian Inference and Prediction for a Type-II censored Weibull Distribution, **Journal of Statistical Planning and Inference** 142(1), 41-47.
- [14] NESREEN and el at (2021). Record data from Kies distribution and related statistical inferences, **Statistics in Transition new series**, 22(2): 153-170.
- [15] PRADHAN, B. and KUNDU, D. (2011). Bayes Estimation and Prediction of the Two-Parameter Gamma Distribution, **Journal of Statistical Computation and Simulation** 81:9, 1187-1198.
- [16] XU K., XIE M., TANG L. C., HO S. L. (2003). Application of neural networks in forecasting engine systems reliability. **Applied Soft Computing** 2, 255-268.