



Cairo University
Faculty of Graduate Studies for Statistical
Research

The 54th Annual Conference on Statistics,
Computer Sciences and Operations Research

Mathematical Statistics

9-11 Dec. 2019

Index
MATHEMTICAL STATISTICS

1	Bivariate Generalized Rayleigh Distribution Based on Clayton Copula El-Sayed A.El-Sherpieny and Ehab M.Almetwally	1-19
2	Estimation of Chen Distribution Based on Record Values Mohamed Mubarak and Shimaa Ashraf	20-33
3	Estimation of the Lifetime Performance Index with Burr Type III Distribution Under Type II Censoring Amal S.Hassan, Salwa M.Assar and Amany S.Selmy	34-49
4	Modified Topp-Leone-Chen Distribution: Properties and Estimation Based on Progressive Type-II Censoring Scheme AL-Sayed, N.T.,Swielum,E.M.,AL-Dayian,G.R. and EL-Helbawy,A.A.	50-75
5	Estimation and Prediction Based on Constant Stress-Partially Accelerated Life Testing for Topp Leone-Inverted Kumaraswamy Distribution Behairy, S.M., Refaey,R.M., AL-Dayian,G.R. and El-Helbawy, A.A.	76- 93
6	Estimation for Exponentiated Generalized Xgamma Distribution from New General Class Based on Dual Generalized Order Statistics A.M.Abd AL-Fattah, R.E.Abd EL-Kader, G.R.AL-Dayian and A. A.El-Helbawy	94-117
7	Discrete Inverted Kumaraswamy Distribution: Properties and Estimation Hegazy,M.A.,Abd El-Kader,R.E.,AL-Dayian,G.R. and El-Helbawy,A. A.	118-150
8	Parameter Estimation of Xgamma Distribution based on Grouped Data Marwa Abd-Alla Abd-Elghafar and Eman Al-Saeed Shehata Sharaf	151- 166
9	New Transformation for Generating Families of Continuous Distributions with Application to Weibull Distribution Samah Ali El-Taweel	167-184
10	Right Truncated Lomax Distribution: Properties and Estimation Mohamed A.H.Sabry and A.Elsehetry	185-199
11	Three Parameters Estimation of Log-Logistic Distribution Using Algorithm of Percentile Roots Doaa Abd El-Shafi Abd El-Rahman and Mohammed Mohammed El Genidy	200-210

Bivariate Generalized Rayleigh Distribution Based on Clayton Copula

El-Sayed A. El-Sherpieny¹, Ehab M. Almetwally¹

¹Department of Mathematical Statistics, Faculty of Graduate Studies for Statistical Research, Cairo University, Egypt.

Abstract: The bivariate generalized Rayleigh distribution is an important lifetime distribution in survival analysis. In this paper, Clayton copula and generalized Rayleigh marginal distribution are used for creating bivariate distribution which are called bivariate generalized Rayleigh based on Clayton copula (Clayton-BGR) distribution. The reliability function and hazard function were obtained for bivariate generalized Rayleigh distribution. Two different estimation methods for the unknown parameters of the Clayton-BGR model were discussed. To evaluate the performance of the estimators, a Monte Carlo simulation study was conducted to compare the efficiency between the Clayton-BGR and another models. Also, a real data set is analyzed to investigate the models and useful results are obtained for illustrative purposes.

Keywords: Generalized Rayleigh Distribution; Clayton copula; Maximum Likelihood Estimation; Inference Function for Margins, Monte Carlo Simulations.

1. Introduction

Burr (1942) introduced twelve different forms of cumulative distribution functions for modeling lifetime data. Among those distributions, Burr Type-X which is called the two-parameter generalized Rayleigh (GR) distribution is the most popular one. Several authors consider different aspects of the generalized Rayleigh distribution, see for example: In Kundu and Raqab (2005), different estimation procedures have been used to estimate the unknown parameter of GR distribution. A random variable X has the two-parameter GR distribution with shape and scale parameters α, λ respectively, the cumulative distribution function (cdf) of GR distribution is given by

$$F(x; \alpha, \lambda) = (1 - e^{-(\lambda x)^2})^\alpha; \quad x > 0, \alpha, \lambda > 0, \quad (1.1)$$

and the probability density function (pdf) of GR distribution is given by

$$f(x; \alpha, \lambda) = 2\alpha\lambda^2 x e^{-(\lambda x)^2} (1 - e^{-(\lambda x)^2})^{\alpha-1}. \quad (1.2)$$

Many authors introduced many applications for GR distribution as follows: Raqab and Madi (2011) discussed the maximum likelihood and Bayesian approach based on progressive Type-II censoring for GR distribution. Tomer et al. (2014) have discussed the maximum likelihood and Bayesian estimation of the parameters for GR distribution and reliability function under Type-I progressive hybrid censoring scheme. Kundu and Raqab (2015) considered the estimation of the stress-strength parameter $R = P[Y < X]$, when X and Y are both three-parameter generalized Rayleigh distributions with the same scale and locations parameters but different shape parameters. Almetwally et al. (2019b) discussed parameter estimation for the GR distribution under the adaptive Type-II progressive censoring schemes based on maximum product spacing and maximum likelihood estimation. There are many uses and multiple applications for this distribution.

The copula is a convenient approach to describe a bivariate distribution with dependence structure. Nelsen (2006) showed the definition of copulas as following; copula is a function that join bivariate distribution functions with uniform $[0, 1]$ margins. Sklar (1973) introduced the joint pdf and joint cdf for the two dimension copula as follows, He considered the two random variables X_1 and X_2 , with distribution functions $F_1(x_1)$ and $F_2(x_2)$ respectively, then the joint cdf and joint pdf for bivariate copula are respectively given as

$$F(x_1, x_2) = C(F_1(x_1), F_2(x_2)), \quad (1.3)$$

Where C is denoted the cdf of copula and c is denoted the pdf of copula.

$$f(x_1, x_2) = f_1(x_1)f_2(x)c(F_1(x_1), F_2(x_2)). \quad (1.4)$$

In Archimedian copulas, there are three copulas in common use: the Clayton, Frank and Gumbel. Nadarajah et al. (2018) discussed Clayton in general form as

$$C(u_1, u_2, \dots, u_p) = \left[\sum_{i=1}^p u_i^{-\alpha} - p + 1 \right]^{-1/\alpha}, \quad \alpha > 0, p \text{ is a number of variable}$$

In this paper, we will be discussed the 2-dimintion of Clayton copula, this is an asymmetric Archimedean copula, exhibiting greater dependence in the negative tail than in the positive, that have been introduced by Clayton (1978). The joint cdf of Clayton copula is given by:

$$C(u, v) = (u^{-\theta} + v^{-\theta} - 1)^{-1/\theta}, \quad (1.5)$$

where $u, v \in [0, 1]$ and the joint pdf of Clayton copula is given by:

$$c(u, v) = \frac{\partial^2 C(u, v)}{\partial u \partial v} = (\theta + 1)(uv)^{-\theta-1}(u^{-\theta} + v^{-\theta} - 1)^{\frac{-2\theta-1}{\theta}}. \quad (1.6)$$

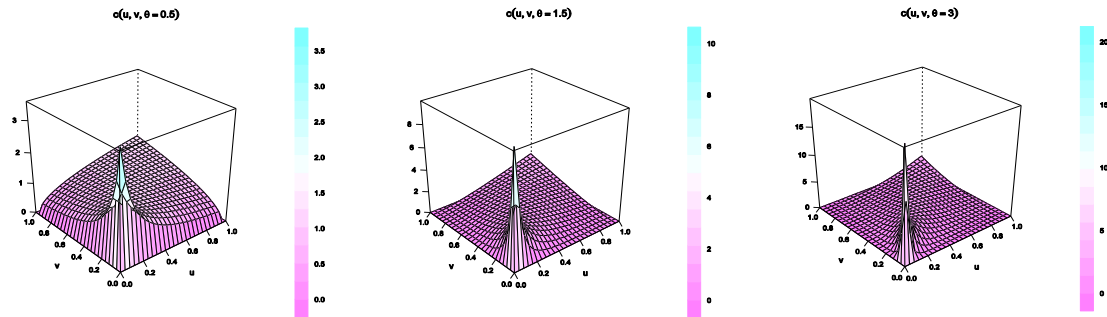


Figure 1. The pdf of Clayton copula with Various Value of the Copula Parameter (θ).

Flores (2009) presented six different bivariate Weibull distributions derived from different copula functions specially Bivariate Weibull based on Clayton copula can be denoted Clayton-BW distribution. Almetwally (2019) showed Clayton copula and introduced bivariate Weibull distribution based on Farlie–Gumbel–Morgenstern (FGM) copula. Fredricks and Nelsen (2007) derived the formula for Spearman's and Kendall's correlation coefficient. In case of Clayton copula, the Spearman's and Kendall's correlation as follows

$$\rho_{sperman} = \left(12 \int \int (u^{-\theta} + v^{-\theta} - 1)^{-1/\theta} du dv \right) - 3, \quad (1.7)$$

and

$$\rho_{Kendall} = 1 - 4 \int \int \frac{\partial C}{\partial u} C(u, v) \frac{\partial C}{\partial v} C(u, v) du dv = \frac{\theta}{\theta+2}, \quad (1.8)$$

where $\frac{\partial C}{\partial u} C(u, v) = u^{-\theta-1} (u^{-\theta} + v^{-\theta} - 1)^{\frac{-\theta-1}{\theta}}$, $\frac{\partial C}{\partial v} C(u, v) = v^{-\theta-1} (u^{-\theta} + v^{-\theta} - 1)^{\frac{-\theta-1}{\theta}}$.

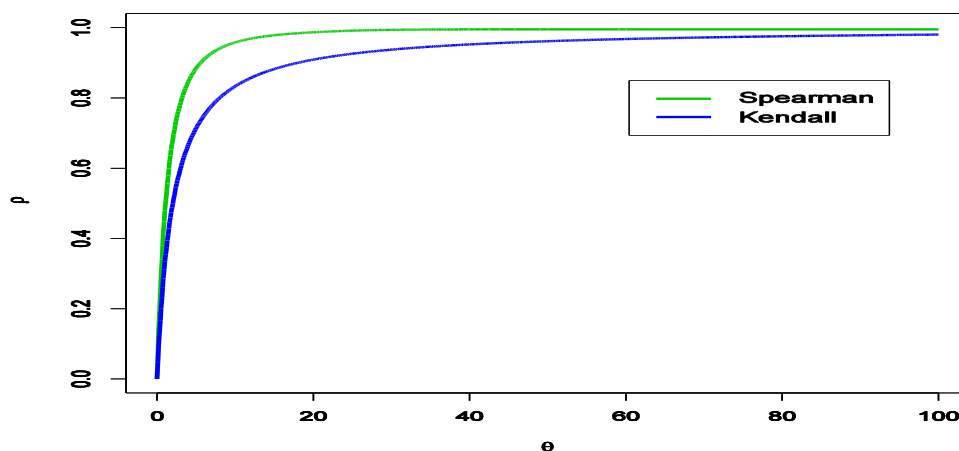


Figure 2. Correlation of Clayton Copula with Various Value of Copula Parameter.

In Bivariate Generalized Rayleigh Distribution, Abdel-Hady (2013) introduced Marshall and Olkin's bivariate exponential model to the generalized bivariate Rayleigh (GBR) Distribution. Sarhan (2019) introduced bivariate generalized Rayleigh distribution of type of

Marshall-Olkin (Shock model), this can be denoted BMOGR and the maximum likelihood and Bayes methods were applied to estimate the unknown parameters for BMOGR distribution. Shubhashree (2019) discussed parameter estimation of multicomponent stress-strength system reliability for a BGR distribution.

The main aim of this paper is to use the similar idea as of El-Sherpieny et al. (2018) to introduce a new extension for bivariate generalized Rayleigh distribution, using the GR distributions and Clayton copula function. The proposed distribution has five parameters. The joint cdf, the joint pdf, the joint survival function and the joint hazard rate function of the Clayton-BGR distribution are derived in closed forms. Parameter estimation for Clayton bivariate GR distribution is introduced by two estimation methods which are: the MLE and inference functions for margins (IFM). Therefore, numerical methods as Monte Carlo Simulation are required to calculate them.

The rest of this paper is organized as follows: Clayton bivariate GR distribution is obtained in Section 2. Parameter estimation methods for the Clayton bivariate GR distribution are obtained in Section 3. In Section 4, the potentiality of the new model is illustrated by Monte Carlo simulation study. In Section 5, Application of a real data set is discussed. Finally, Conclusions are addressed in Section 6.

2. Clayton Bivariate GR Distribution

According to Sklar theorem, using (1.1) and (1.2) in (1.3) and (1.4), the joint cdf of the bivariate GR distribution for any copula is as follows

$$F(x_1, x_2) = C \left((1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1}, (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2} \right), \quad (2.1)$$

and the joint pdf of the bivariate GR distribution for any copula is as follows

$$\begin{aligned} f(x_1, x_2) = & 4\alpha_1 \lambda_1^2 x_1 e^{-(\lambda_1 x_1)^2} (1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1 - 1} \alpha_2 \lambda_2^2 x_2 e^{-(\lambda_2 x_2)^2} (1 \\ & - e^{-(\lambda_2 x_2)^2})^{\alpha_2 - 1} \\ & c \left((1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1}, (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2} \right). \end{aligned} \quad (2.2)$$

In the following, The Equations (2.1) and (2.2) are used to define the Clayton bivariate GR distribution. The cdf, pdf, reliability or survival function and hazard function for Clayton-BGR distribution are obtained as follows:

$$F(x_1, x_2) = \left[\left((1 - e^{-(\lambda_1 x_1)^2})^{-\theta \alpha_1} + (1 - e^{-(\lambda_2 x_2)^2})^{-\theta \alpha_2} - 1 \right)^{-1/\theta} \right], \quad (2.3)$$

$$\begin{aligned} f(x_1, x_2) = & 4\alpha_1 \lambda_1^2 x_1 e^{-(\lambda_1 x_1)^2} (1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1 - 1} \alpha_2 \lambda_2^2 x_2 e^{-(\lambda_2 x_2)^2} (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2 - 1}, \\ & (\theta + 1) \left((1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1} (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2} \right)^{-\theta - 1} \left((1 - e^{-(\lambda_1 x_1)^2})^{-\theta \alpha_1} \right. \\ & \left. + (1 - e^{-(\lambda_2 x_2)^2})^{-\theta \alpha_2} - 1 \right)^{\frac{-2\theta - 1}{\theta}}, \end{aligned} \quad (2.4)$$

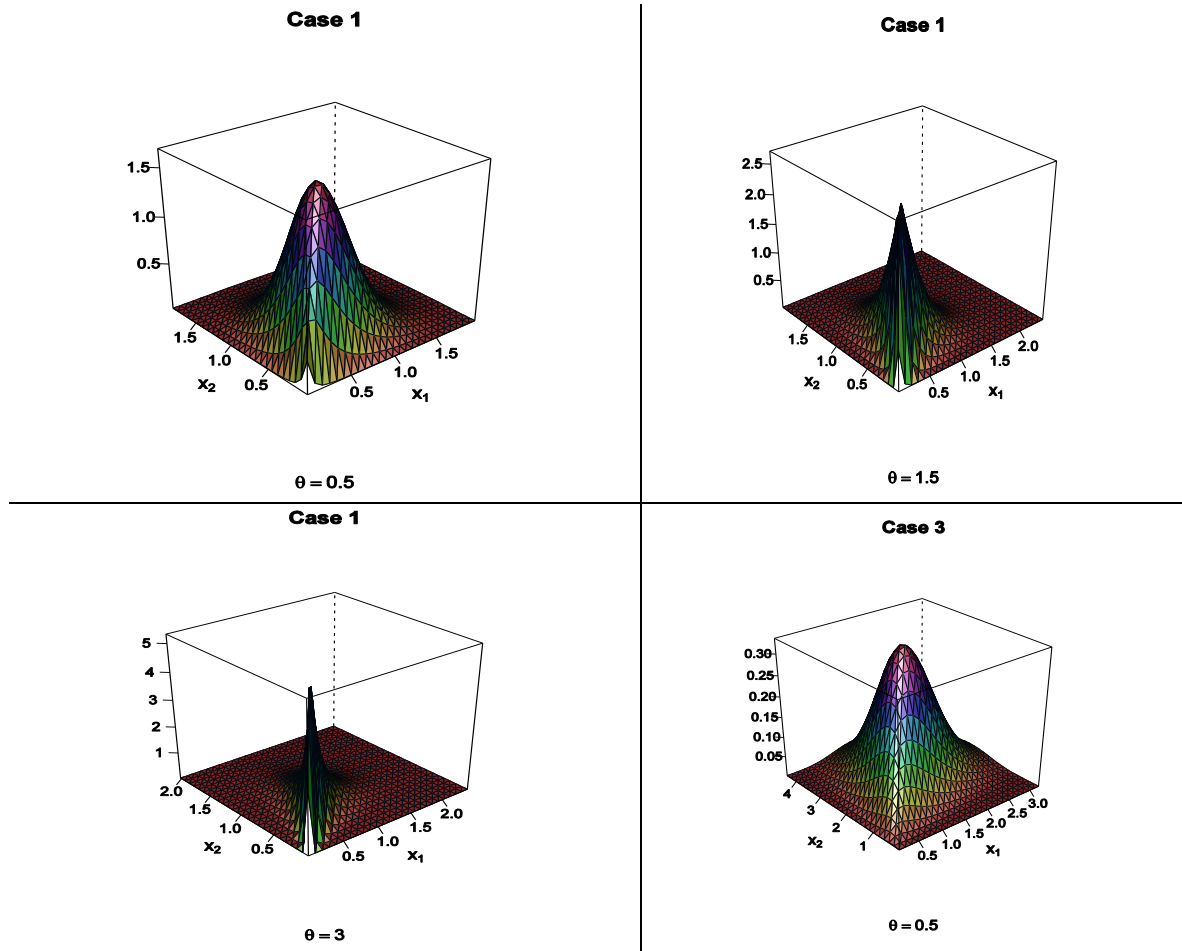
$$R(x_1, x_2) = \left(\left(1 - (1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1} \right)^{-\theta} + \left(1 - (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2} \right)^{-\theta} - 1 \right)^{-1/\theta}, \quad (2.5)$$

and

$$h(x_1, x_2) = \frac{f(x_1, x_2)}{\left(\left(1 - (1 - e^{-(\lambda_1 x_1)^2})^{\alpha_1} \right)^{-\theta} + \left(1 - (1 - e^{-(\lambda_2 x_2)^2})^{\alpha_2} \right)^{-\theta} - 1 \right)^{-1/\theta}}, \quad (2.6)$$

respectively, where $\alpha_1, \lambda_1, \alpha_2, \lambda_2, \theta > 0$ and $x_1, x_2 > 0$

Figure 3 shows the plot 3-dimension for the pdf of Clayton-BGR distribution with different parameters values.



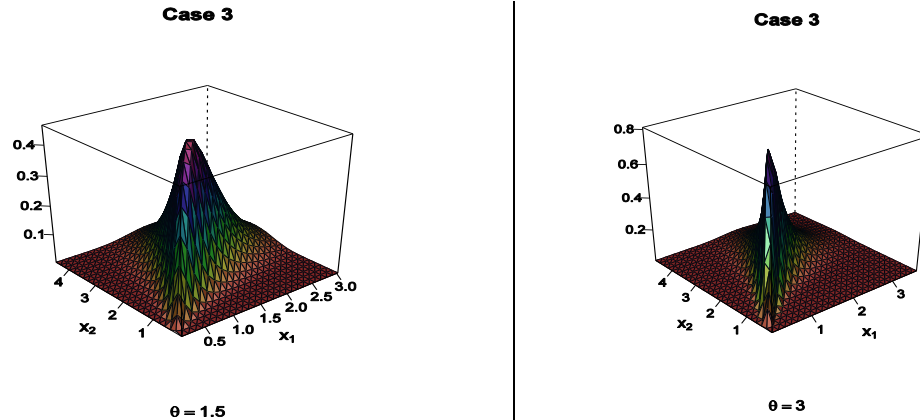
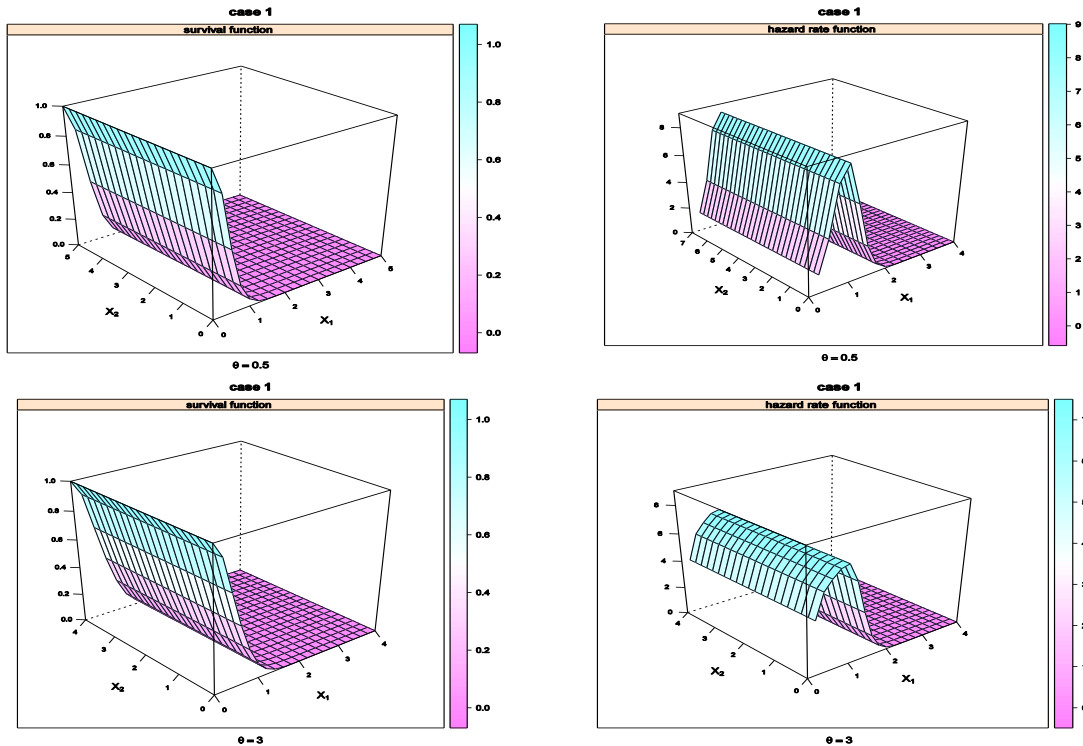


Figure 3. The pdf of Clayton-BGR Distribution with Various Value of the Parameters.

Figure 4 shows the plot 3-dimension for the survival function and hazard rate function of Clayton-BGR distribution with different parameters values.



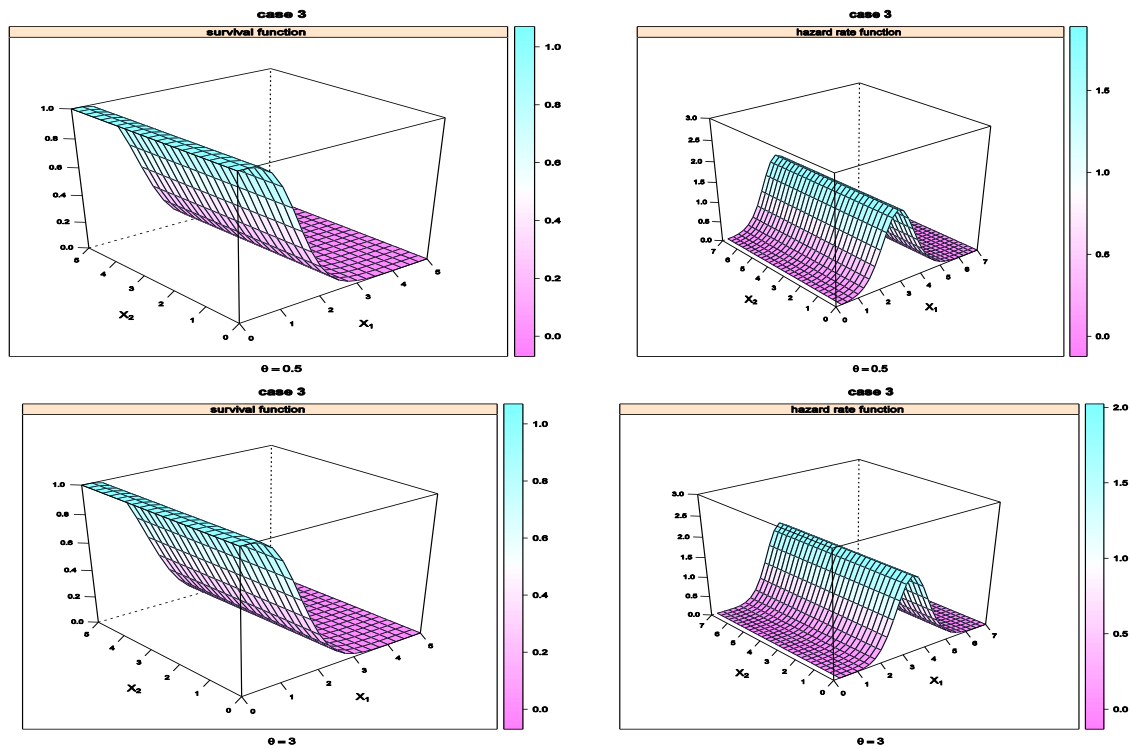


Figure 4. Survival Function and Hazard Rate Function of Clayton-BGR Distribution with Various Value of the Parameters.

As shown in the Figure 3, 4, the curve of Clayton-BGR distribution behavior have more than one direction where it takes increasing and decreasing direction, which will have many applications in life testing. Sub-model can be derived from Clayton-BGR distribution as Clayton bivariate Rayleigh (Clayton-BR) distribution (new) when $\alpha_1 = 1, \lambda_1 = \lambda_1, \alpha_2 = 1, \lambda_2 = \lambda_2, \theta = \theta$.

3. Parameter Estimation Methods

In the section, we introduce two estimation methods that used to estimate the unknown parameters of Clayton-BGR distribution, such as: MLE and IFM. To more information about these methods see Kim et al. (2007) and to more example see Elaal and Jarwan (2017) and Almetwally et al. (2019_a).

3.1. Maximum Likelihood Estimation (MLE)

The log-likelihood function of Clayton-BGR distribution can be written as

$$\begin{aligned}
l(\Omega) = & 2n(\ln(\lambda_1) + \ln(\lambda_2)) + n(\ln(4) + \ln(\alpha_1) + \ln(\alpha_2) + \ln(\theta + 1)) \\
& + \sum_{i=1}^n (\ln(x_{1i}) + \ln(x_{2i})) - \sum_{i=1}^n (\lambda_1 x_{1i})^2 - \sum_{i=1}^n (\lambda_2 x_{2i})^2 \\
& + (-1 - \theta\alpha_1) \sum_{i=1}^n \ln(1 - e^{-(\lambda_1 x_{1i})^2}) \\
& + (-1 - \theta\alpha_2) \sum_{i=1}^n \ln(1 - e^{-(\lambda_2 x_{2i})^2}) \\
& - \frac{2\theta + 1}{\theta} \sum_{i=1}^n \ln \left((1 - e^{-(\lambda_1 x_{1i})^2})^{-\theta\alpha_1} + (1 - e^{-(\lambda_2 x_{2i})^2})^{-\theta\alpha_2} \right. \\
& \left. - 1 \right)
\end{aligned} \tag{3.1}$$

where Ω is a vector of parameters.

To obtain the normal likelihood equations for the unknown parameters, we differentiate equation (3.1) partially with Ω vector of parameters and equate them to zero. The estimators for $\hat{\Omega}_{MLE}$ can be obtained as in the solution of the following equations:

$$\begin{aligned}
\frac{\partial l(\Omega)}{\partial \alpha_j} &= \frac{n}{\alpha_j} - \theta \sum_{i=1}^n \ln \left(1 - e^{-(\lambda_j x_{ji})^2} \right) - \\
& \frac{2\theta + 1}{\theta} \sum_{i=1}^n \frac{-\theta \left(1 - e^{-(\lambda_j x_{ji})^2} \right)^{-\theta\alpha_j} \ln \left(1 - e^{-(\lambda_j x_{ji})^2} \right)}{\left(1 - e^{-(\lambda_1 x_{1i})^2} \right)^{-\theta\alpha_1} + \left(1 - e^{-(\lambda_2 x_{2i})^2} \right)^{-\theta\alpha_2} - 1} = 0, \\
\frac{\partial l(\Omega)}{\partial \lambda_j} &= \frac{2n}{\lambda_j} + 2\lambda_j \sum_{i=1}^n x_{ji}^2 - (1 + \theta\alpha_j) \sum_{i=1}^n \frac{2\lambda_j x_{ji}^2 e^{-(\lambda_j x_{ji})^2}}{1 - e^{-(\lambda_j x_{ji})^2}} - \\
& \frac{2\theta + 1}{\theta} \sum_{i=1}^n \frac{-\theta\alpha_j 2\lambda_j x_{ji}^2 \left(1 - e^{-(\lambda_1 x_{1i})^2} \right)^{-\theta\alpha_j - 1}}{\left(1 - e^{-(\lambda_1 x_{1i})^2} \right)^{-\theta\alpha_1} + \left(1 - e^{-(\lambda_2 x_{2i})^2} \right)^{-\theta\alpha_2} - 1} = 0,
\end{aligned}$$

and

$$\begin{aligned}
\frac{\partial l(\Omega)}{\partial \theta} &= \frac{n}{\theta + 1} - \alpha_j \sum_{i=1}^n \ln \left(1 - e^{-(\lambda_j x_{ji})^2} \right) + \frac{1}{\theta^2} \sum_{i=1}^n \ln \left(\left(1 - e^{-(\lambda_1 x_{1i})^2} \right)^{-\theta\alpha_1} + \right. \\
& \left. \left(1 - e^{-(\lambda_2 x_{2i})^2} \right)^{-\theta\alpha_2} - 1 \right) - \\
& \frac{2\theta + 1}{\theta} \sum_{i=1}^n \frac{-\alpha_j \left(1 - e^{-(\lambda_j x_{ji})^2} \right)^{-\theta\alpha_j} \ln \left(1 - e^{-(\lambda_j x_{ji})^2} \right) - \alpha_l \left(1 - e^{-(\lambda_l x_{li})^2} \right)^{-\theta\alpha_l} \ln \left(1 - e^{-(\lambda_l x_{li})^2} \right)}{\left(1 - e^{-(\lambda_1 x_{1i})^2} \right)^{-\theta\alpha_1} + \left(1 - e^{-(\lambda_2 x_{2i})^2} \right)^{-\theta\alpha_2} - 1} = 0,
\end{aligned}$$

where $j = l = 1, 2 ; j \neq l$. The MLEs $\hat{\Omega}_{MLE}$ are obtained by performing numerically using a nonlinear optimization algorithm.

3.2. Estimation by Inference Functions for Margins (IFM)

Kim et al. (2007) discussed this parametric method with two-step of estimation. Firstly, each marginal of GR distribution is estimated separately, by using Equation (1.2), the log-likelihood function of GR distribution can be written as

$$\begin{aligned} \ln L_j(\alpha_j, \lambda_j) &= \sum_{i=1}^n \ln f_j(x_{ji}, \alpha_j, \lambda_j); j = 1, 2 \\ \ln L_j(\alpha_j, \lambda_j) &= n(\ln(2) + \ln(\alpha_j) + 2 \ln(\lambda_j)) + \sum_{i=1}^n (\ln(x_{ji})) - \sum_{i=1}^n (\lambda_j x_{ji})^2 \\ &\quad + (\alpha_j - 1) \sum_{i=1}^n \ln(1 - e^{-(\lambda_j x_{ji})^2}) \end{aligned} \quad (3.2)$$

To obtain the normal likelihood equations for the unknown parameters, we differentiate equation (3.2) partially with α_j, λ_j vector of parameters and equate them to zero. The estimators for $\hat{\alpha}_j, \hat{\lambda}_j$ can be obtained as in the solution of the following equations:

$$\begin{aligned} \frac{\partial \ln L_j(\alpha_j, \lambda_j)}{\partial \alpha_j} &= \frac{n}{\alpha_j} + \sum_{i=1}^n \ln(1 - e^{-(\lambda_j x_{ji})^2}), \\ \frac{\partial \ln L_j(\alpha_j, \lambda_j)}{\partial \lambda_j} &= \frac{2n}{\lambda_j} + 2\lambda_j \sum_{i=1}^n x_{ji}^2 + (\alpha_j - 1) \sum_{i=1}^n \frac{2\lambda_j x_{ji}^2 e^{-(\lambda_j x_{ji})^2}}{1 - e^{-(\lambda_j x_{ji})^2}}. \end{aligned}$$

The MLEs($\hat{\alpha}_j, \hat{\lambda}_j$) can be obtained by solving simultaneously the likelihood equations

$$\left. \frac{\partial \ln L_j(\alpha_j, \lambda_j)}{\partial \lambda_j} \right|_{\lambda_j = \hat{\lambda}_j} = 0, \quad \left. \frac{\partial \ln L_j(\alpha_j, \lambda_j)}{\partial \alpha_j} \right|_{\alpha_j = \hat{\alpha}_j} = 0, \quad j = 1, 2$$

In the second step, the copula parameter is estimated by maximizing the log-likelihood function of the copula density using the MLE estimates of the marginal $\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1)$ and $\hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2)$. The log-likelihood function by using IFM method estimation of any copula function is defined as

$$l(\theta) = \ln L_{IFM}(\theta) = \sum_{i=1}^n \ln \left(c \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1), \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2) \right) \right).$$

By using Equation (1.6) then, the log-likelihood function of Clayton copula function can be written as

$$\begin{aligned} l(\theta) &= n \ln(\theta + 1) - (\theta + 1) \sum_{i=1}^n \ln \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1) \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2) \right) \\ &\quad - \frac{2\theta + 1}{\theta} \sum_{i=1}^n \ln \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1)^{-\theta} + \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2)^{-\theta} - 1 \right). \end{aligned} \quad (3.3)$$

Basing on this, differentiating the log-likelihood function with respect to θ for Clayton-BGR distribution is given as

$$\begin{aligned} \frac{\partial l(\theta)}{\partial \theta} &= \frac{n}{\theta + 1} - \sum_{i=1}^n \ln \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1) \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2) \right) \\ &+ \frac{1}{\theta^2} \sum_{i=1}^n \ln \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1)^{-\theta} + \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2)^{-\theta} - 1 \right) \\ &- \frac{2\theta + 1}{\theta} \sum_{i=1}^n \frac{-\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1)^{-\theta} \ln \left(\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1) \right) - \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2)^{-\theta} \ln \left(\hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2) \right)}{\hat{F}_{1i}(x_{1i}, \hat{\alpha}_1, \hat{\lambda}_1)^{-\theta} + \hat{F}_{2i}(x_{2i}, \hat{\alpha}_2, \hat{\lambda}_2)^{-\theta} - 1} \end{aligned}$$

The estimates of copula parameter is handled numerically simultaneously of the likelihood equations

$$\left. \frac{\partial l(\theta)}{\partial \theta} \right|_{\theta=\hat{\theta}} = 0$$

There is no closed-form expression for the MLE $\hat{\theta}$ and its computation has to be performed numerically using a nonlinear optimization algorithm.

4. Simulation Study

In this section; a Monte Carlo simulation is done for comparison between estimation methods based on copula such as: MLE and IFM. For estimating Clayton-BGR distribution parameters by R packages.

To generate random variables: Nelsen (2006) discussed generating a sample from a specified joint distribution. By conditional distribution method, the joint distribution function is as follows

$$f(x_1, x_2) = f(x_1)f(x_2|x_1)$$

We used the R packages in copula package to generate the random variables of Clayton-BGR distribution.

Simulation Algorithm: A simulation experiments were carried out based on the following data-generated form Clayton-BGR distributions, the values of the parameters $\alpha_1, \lambda_1, \alpha_2, \lambda_2$ and θ is chosen as the following cases for the random variables generating:

Case 1: $(\alpha_1 = 1.2, \lambda_1 = 1.4, \alpha_2 = 1.1, \lambda_2 = 1.5, \theta = (0.5 \text{ or } 1.5 \text{ or } 3 \text{ or } 6))$,

Case 2: $(\alpha_1 = 0.5, \lambda_1 = 1.4, \alpha_2 = 0.6, \lambda_2 = 1.5, \theta = (0.5 \text{ or } 1.5 \text{ or } 3 \text{ or } 6))$,

Case 3: $(\alpha_1 = 2, \lambda_1 = 0.7, \alpha_2 = 2.5, \lambda_2 = 0.5, \theta = (0.5 \text{ or } 1.5 \text{ or } 3 \text{ or } 6))$.

For different sample size $n = 50, 100, 200$. The simulation methods are compared using the criteria of parameters estimation, the comparison is performed by calculate in the Bias and the Root mean square error (RMSE) for each method of estimation as following

$$Bias = (\hat{\Omega} - \Omega) . \quad (4.1)$$

where $\hat{\Omega}$ is the estimated value of Ω , and RMSE is as follows

$$RMSE = \text{mean}\sqrt{(\hat{\Omega} - \Omega)^2}. \quad (4.2)$$

We restricted the number of repeated-samples to 1000.

From Table (1) to Table(9) we can conclude the following :

- i. By increasing the sample size, the RMSE and Bias of the considered parameters of Clayton-BGR distribution are decreases.
- ii. If the copula parameter θ increases for any cases, then the value of the Bias and RMSE increase for the parameters especially θ of Clayton-BGR distribution.
- iii. In IFM method, we note that the values of the $\hat{\alpha}_1, \hat{\lambda}_1, \hat{\alpha}_2, \hat{\lambda}_2$ parameters estimation are repeated with changing of θ due to the properties of the first-step in IFM method, which parameters estimated of marginal of GR distribution. They are not affected by the change of θ .
- iv. The IFM estimates has more efficiency than MLE for copula parameter of Clayton-BGR distribution.

The IFM method is better than MLE method, it is noted that result, IFM method is the best method because it is a two steps of estimation, first, the marginal distribution parameters estimated and second the copula parameter is estimated, taking into consideration of previous parameter estimates of marginal distribution.

Table 1: Estimation of the Parameters of Clayton-BGR Distribution: Case 1, $n=50$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	1.2658	0.0658	0.2736	1.2686	0.0686	0.2792
	$\hat{\lambda}_1$	1.4299	0.0299	0.1340	1.4307	0.0307	0.1361
	$\hat{\alpha}_2$	1.1545	0.0545	0.2303	1.1543	0.0543	0.2318
	$\hat{\lambda}_2$	1.5296	0.0296	0.1422	1.5294	0.0294	0.1422
	$\hat{\theta}$	0.5368	0.0368	0.2548	0.5292	0.0292	0.2485
1.5	$\hat{\alpha}_1$	1.2644	0.0644	0.2639	1.2686	0.0686	0.2792
	$\hat{\lambda}_1$	1.4301	0.0301	0.1286	1.4307	0.0307	0.1361
	$\hat{\alpha}_2$	1.1506	0.0506	0.2195	1.1518	0.0518	0.2291
	$\hat{\lambda}_2$	1.5305	0.0305	0.1377	1.5297	0.0297	0.1433
	$\hat{\theta}$	1.5787	0.0787	0.4506	1.5446	0.0446	0.4392
3	$\hat{\alpha}_1$	1.2644	0.0644	0.2595	1.2686	0.0686	0.2792
	$\hat{\lambda}_1$	1.4305	0.0305	0.1226	1.4307	0.0307	0.1361
	$\hat{\alpha}_2$	1.1526	0.0526	0.2220	1.1545	0.0545	0.2384
	$\hat{\lambda}_2$	1.5323	0.0323	0.1331	1.5313	0.0313	0.1465
	$\hat{\theta}$	3.1345	0.1345	0.7527	3.0370	0.0370	0.7268
6	$\hat{\alpha}_1$	1.2639	0.0639	0.2554	1.2686	0.0686	0.2792
	$\hat{\lambda}_1$	1.4308	0.0308	0.1173	1.4307	0.0307	0.1361
	$\hat{\alpha}_2$	1.1548	0.0548	0.2247	1.1564	0.0564	0.2440
	$\hat{\lambda}_2$	1.5339	0.0339	0.1290	1.5320	0.0320	0.1489
	$\hat{\theta}$	6.2542	0.2542	1.3608	5.9452	-0.0548	1.2999

Table 2: Estimation of the Parameters of Clayton-BGR Distribution: Case 1, $n=100$

		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	1.2315	0.0315	0.1794	1.2315	0.0315	0.1804
	$\hat{\lambda}_1$	1.4109	0.0109	0.0900	1.4110	0.0110	0.0907
	$\hat{\alpha}_2$	1.1246	0.0246	0.1519	1.1258	0.0258	0.1535
	$\hat{\lambda}_2$	1.5101	0.0101	0.0953	1.5106	0.0106	0.0967
	$\hat{\theta}$	0.5155	0.0155	0.1808	0.5120	0.0120	0.1790
1.5	$\hat{\alpha}_1$	1.2298	0.0298	0.1735	1.2315	0.0315	0.1804
	$\hat{\lambda}_1$	1.4103	0.0103	0.0859	1.4110	0.0110	0.0907
	$\hat{\alpha}_2$	1.1248	0.0248	0.1521	1.1267	0.0267	0.1590
	$\hat{\lambda}_2$	1.5102	0.0102	0.0922	1.5109	0.0109	0.0986
	$\hat{\theta}$	1.5303	0.0303	0.3173	1.5125	0.0125	0.3140
3	$\hat{\alpha}_1$	1.2283	0.0283	0.1698	1.2315	0.0315	0.1804
	$\hat{\lambda}_1$	1.4097	0.0097	0.0817	1.4110	0.0110	0.0907
	$\hat{\alpha}_2$	1.1245	0.0245	0.1502	1.1261	0.0261	0.1603
	$\hat{\lambda}_2$	1.5101	0.0101	0.0883	1.5104	0.0104	0.0991
	$\hat{\theta}$	3.0555	0.0555	0.5205	3.0043	0.0043	0.5132
6	$\hat{\alpha}_1$	1.2277	0.0277	0.1668	1.2315	0.0315	0.1804
	$\hat{\lambda}_1$	1.4096	0.0096	0.0780	1.4110	0.0110	0.0907
	$\hat{\alpha}_2$	1.1247	0.0247	0.1491	1.1256	0.0256	0.1620
	$\hat{\lambda}_2$	1.5106	0.0106	0.0855	1.5098	0.0098	0.0995
	$\hat{\theta}$	6.1031	0.1031	0.9186	5.9378	-0.0622	0.9039

Table 3: Estimation of the Parameters of Clayton-BGR Distribution: Case 1, $n=200$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	1.2137	0.0137	0.1157	1.2142	0.0142	0.1166
	$\hat{\lambda}_1$	1.4066	0.0066	0.0614	1.4068	0.0068	0.0621
	$\hat{\alpha}_2$	1.1127	0.0127	0.1046	1.1131	0.0131	0.1055
	$\hat{\lambda}_2$	1.5073	0.0073	0.0676	1.5074	0.0074	0.0684
	$\hat{\theta}$	0.5161	0.0161	0.1292	0.5143	0.0143	0.1284
1.5	$\hat{\alpha}_1$	1.2139	0.0139	0.1127	1.2142	0.0142	0.1166
	$\hat{\lambda}_1$	1.4068	0.0068	0.0587	1.4068	0.0068	0.0621
	$\hat{\alpha}_2$	1.1108	0.0108	0.1021	1.1118	0.0118	0.1049
	$\hat{\lambda}_2$	1.5065	0.0065	0.0649	1.5070	0.0070	0.0686
	$\hat{\theta}$	1.5273	0.0273	0.2218	1.5186	0.0186	0.2199
3	$\hat{\alpha}_1$	1.2137	0.0137	0.1111	1.2142	0.0142	0.1166
	$\hat{\lambda}_1$	1.4067	0.0067	0.0558	1.4068	0.0068	0.0621
	$\hat{\alpha}_2$	1.1110	0.0110	0.1005	1.1118	0.0118	0.1043
	$\hat{\lambda}_2$	1.5068	0.0068	0.0620	1.5071	0.0071	0.0683
	$\hat{\theta}$	3.0434	0.0434	0.3638	3.0186	0.0186	0.3590
6	$\hat{\alpha}_1$	1.2137	0.0137	0.1098	1.2142	0.0142	0.1166
	$\hat{\lambda}_1$	1.4068	0.0068	0.0530	1.4068	0.0068	0.0621
	$\hat{\alpha}_2$	1.1116	0.0116	0.0996	1.1124	0.0124	0.1046
	$\hat{\lambda}_2$	1.5073	0.0073	0.0594	1.5076	0.0076	0.0683
	$\hat{\theta}$	6.0741	0.0741	0.6461	5.9911	-0.0089	0.6330

Table 4: Estimation of the Parameters of Clayton-BGR Distribution: Case 2, $n=50$

θ	n=50	MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	0.5191	0.0191	0.0937	0.5199	0.0199	0.0952
	$\hat{\lambda}_1$	1.4433	0.0433	0.1772	1.4444	0.0444	0.1794
	$\hat{\alpha}_2$	0.6232	0.0232	0.1098	0.6233	0.0233	0.1104
	$\hat{\lambda}_2$	1.5375	0.0375	0.1707	1.5375	0.0375	0.1708
	$\hat{\theta}$	0.5367	0.0367	0.2530	0.5299	0.0299	0.2470
1.5	$\hat{\alpha}_1$	0.5188	0.0188	0.0914	0.5199	0.0199	0.0952
	$\hat{\lambda}_1$	1.4430	0.0430	0.1719	1.4444	0.0444	0.1794
	$\hat{\alpha}_2$	0.6212	0.0212	0.1058	0.6220	0.0220	0.1100
	$\hat{\lambda}_2$	1.5383	0.0383	0.1662	1.5383	0.0383	0.1729
	$\hat{\theta}$	1.5776	0.0776	0.4455	1.5466	0.0466	0.4350
3	$\hat{\alpha}_1$	0.5190	0.0190	0.0911	0.5199	0.0199	0.0952
	$\hat{\lambda}_1$	1.4431	0.0431	0.1657	1.4444	0.0444	0.1794
	$\hat{\alpha}_2$	0.6218	0.0218	0.1067	0.6231	0.0231	0.1138
	$\hat{\lambda}_2$	1.5396	0.0396	0.1599	1.5404	0.0404	0.1770
	$\hat{\theta}$	3.1321	0.1321	0.7454	3.0437	0.0437	0.7217
6	$\hat{\alpha}_1$	0.5191	0.0191	0.0909	0.5199	0.0199	0.0952
	$\hat{\lambda}_1$	1.4433	0.0433	0.1607	1.4444	0.0444	0.1794
	$\hat{\alpha}_2$	0.6226	0.0226	0.1080	0.6238	0.0238	0.1157
	$\hat{\lambda}_2$	1.5406	0.0406	0.1534	1.5411	0.0411	0.1797
	$\hat{\theta}$	6.2489	0.2489	1.3529	5.9675	-0.0325	1.2948

Table 5: Estimation of the Parameters of Clayton-BGR Distribution: Case 2, $n=100$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	0.5101	0.0101	0.0634	0.5099	0.0099	0.0634
	$\hat{\lambda}_1$	1.4165	0.0165	0.1178	1.4166	0.0166	0.1182
	$\hat{\alpha}_2$	0.6114	0.0114	0.0745	0.6118	0.0118	0.0751
	$\hat{\lambda}_2$	1.5140	0.0140	0.1150	1.5144	0.0144	0.1164
	$\hat{\theta}$	0.5148	0.0148	0.1791	0.5118	0.0118	0.1774
1.5	$\hat{\alpha}_1$	0.5097	0.0097	0.0620	0.5099	0.0099	0.0634
	$\hat{\lambda}_1$	1.4159	0.0159	0.1138	1.4166	0.0166	0.1182
	$\hat{\alpha}_2$	0.6114	0.0114	0.0748	0.6121	0.0121	0.0777
	$\hat{\lambda}_2$	1.5138	0.0138	0.1111	1.5145	0.0145	0.1180
	$\hat{\theta}$	1.5286	0.0286	0.3131	1.5132	0.0132	0.3102
3	$\hat{\alpha}_1$	0.5093	0.0093	0.0613	0.5099	0.0099	0.0634
	$\hat{\lambda}_1$	1.4150	0.0150	0.1093	1.4166	0.0166	0.1182
	$\hat{\alpha}_2$	0.6112	0.0112	0.0741	0.6117	0.0117	0.0783
	$\hat{\lambda}_2$	1.5132	0.0132	0.1053	1.5136	0.0136	0.1182
	$\hat{\theta}$	3.0523	0.0523	0.5162	3.0081	0.0081	0.5092
6	$\hat{\alpha}_1$	0.5092	0.0092	0.0609	0.5099	0.0099	0.0634
	$\hat{\lambda}_1$	1.4147	0.0147	0.1055	1.4166	0.0166	0.1182
	$\hat{\alpha}_2$	0.6112	0.0112	0.0735	0.6114	0.0114	0.0786
	$\hat{\lambda}_2$	1.5135	0.0135	0.1007	1.5128	0.0128	0.1187
	$\hat{\theta}$	6.0985	0.0985	0.9164	5.9511	-0.0489	0.9009

Table 6: Estimation of the Parameters of Clayton-BGR Distribution: Case 2, $n=200$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	0.5041	0.0041	0.0420	0.5043	0.0043	0.0423
	$\hat{\lambda}_1$	1.4095	0.0095	0.0804	1.4098	0.0098	0.0810
	$\hat{\alpha}_2$	0.6056	0.0056	0.0517	0.6058	0.0058	0.0520
	$\hat{\lambda}_2$	1.5095	0.0095	0.0816	1.5097	0.0097	0.0823
	$\hat{\theta}$	0.5161	0.0161	0.1282	0.5145	0.0145	0.1275
1.5	$\hat{\alpha}_1$	0.5043	0.0043	0.0414	0.5043	0.0043	0.0423
	$\hat{\lambda}_1$	1.4098	0.0098	0.0781	1.4098	0.0098	0.0810
	$\hat{\alpha}_2$	0.6046	0.0046	0.0508	0.6052	0.0052	0.0518
	$\hat{\lambda}_2$	1.5084	0.0084	0.0784	1.5092	0.0092	0.0822
	$\hat{\theta}$	1.5270	0.0270	0.2194	1.5191	0.0191	0.2175
3	$\hat{\alpha}_1$	0.5042	0.0042	0.0413	0.5043	0.0043	0.0423
	$\hat{\lambda}_1$	1.4096	0.0096	0.0753	1.4098	0.0098	0.0810
	$\hat{\alpha}_2$	0.6047	0.0047	0.0502	0.6051	0.0051	0.0516
	$\hat{\lambda}_2$	1.5084	0.0084	0.0743	1.5093	0.0093	0.0818
	$\hat{\theta}$	3.0428	0.0428	0.3612	3.0204	0.0204	0.3564
6	$\hat{\alpha}_1$	0.5043	0.0043	0.0413	0.5043	0.0043	0.0423
	$\hat{\lambda}_1$	1.4097	0.0097	0.0725	1.4098	0.0098	0.0810
	$\hat{\alpha}_2$	0.6049	0.0049	0.0499	0.6054	0.0054	0.0517
	$\hat{\lambda}_2$	1.5088	0.0088	0.0704	1.5098	0.0098	0.0818
	$\hat{\theta}$	6.0729	0.0729	0.6448	5.9980	-0.0020	0.6314

Table 7: Estimation of the Parameters of Clayton-BGR Distribution: Case 3, $n=50$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	2.1379	0.1379	0.5268	2.1435	0.1435	0.5383
	$\hat{\lambda}_1$	0.7104	0.0104	0.0847	0.7133	0.0133	0.0605
	$\hat{\alpha}_2$	2.6776	0.1776	0.6593	2.6770	0.1770	0.6652
	$\hat{\lambda}_2$	0.5071	0.0071	0.0504	0.5079	0.0079	0.0396
	$\hat{\theta}$	0.5371	0.0371	0.2559	0.5291	0.0291	0.2497
1.5	$\hat{\alpha}_1$	2.1362	0.1362	0.5073	2.1435	0.1435	0.5383
	$\hat{\lambda}_1$	0.7121	0.0121	0.0710	0.7133	0.0133	0.0605
	$\hat{\alpha}_2$	2.6654	0.1654	0.6145	2.6691	0.1691	0.6487
	$\hat{\lambda}_2$	0.5065	0.0065	0.0573	0.5079	0.0079	0.0397
	$\hat{\theta}$	1.5791	0.0791	0.4532	1.5436	0.0436	0.4421
3	$\hat{\alpha}_1$	2.1370	0.1370	0.4971	2.1435	0.1435	0.5383
	$\hat{\lambda}_1$	0.7138	0.0138	0.0548	0.7133	0.0133	0.0605
	$\hat{\alpha}_2$	2.6708	0.1708	0.6162	2.6787	0.1787	0.6792
	$\hat{\lambda}_2$	0.5089	0.0089	0.0363	0.5084	0.0084	0.0405
	$\hat{\theta}$	3.1354	0.1354	0.7559	3.0334	0.0334	0.7303
6	$\hat{\alpha}_1$	2.1359	0.1359	0.4865	2.1435	0.1435	0.5383
	$\hat{\lambda}_1$	0.7140	0.0140	0.0526	0.7133	0.0133	0.0605
	$\hat{\alpha}_2$	2.6757	0.1757	0.6182	2.6863	0.1863	0.7025
	$\hat{\lambda}_2$	0.5093	0.0093	0.0349	0.5086	0.0086	0.0413
	$\hat{\theta}$	6.2577	0.2577	1.3655	5.9348	-0.0652	1.3033

Table 8: Estimation of the Parameters of Clayton-BGR Distribution: Case 3, $n=100$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	2.0631	0.0631	0.3377	2.0635	0.0635	0.3405
	$\hat{\lambda}_1$	0.7047	0.0047	0.0400	0.7048	0.0048	0.0404
	$\hat{\alpha}_2$	2.5736	0.0736	0.4144	2.5775	0.0775	0.4210
	$\hat{\lambda}_2$	0.5026	0.0026	0.0263	0.5027	0.0027	0.0268
	$\hat{\theta}$	0.5160	0.0160	0.1818	0.5122	0.0122	0.1800
1.5	$\hat{\alpha}_1$	2.0595	0.0595	0.3256	2.0635	0.0635	0.3405
	$\hat{\lambda}_1$	0.7044	0.0044	0.0382	0.7048	0.0048	0.0404
	$\hat{\alpha}_2$	2.5742	0.0742	0.4102	2.5814	0.0814	0.4366
	$\hat{\lambda}_2$	0.5027	0.0027	0.0253	0.5029	0.0029	0.0274
	$\hat{\theta}$	1.5313	0.0313	0.3194	1.5118	0.0118	0.3163
3	$\hat{\alpha}_1$	2.0570	0.0570	0.3185	2.0635	0.0635	0.3405
	$\hat{\lambda}_1$	0.7043	0.0043	0.0365	0.7048	0.0048	0.0404
	$\hat{\alpha}_2$	2.5733	0.0733	0.4007	2.5806	0.0806	0.4405
	$\hat{\lambda}_2$	0.5027	0.0027	0.0241	0.5028	0.0028	0.0277
	$\hat{\theta}$	3.0566	0.0566	0.5220	3.0015	0.0015	0.5149
6	$\hat{\alpha}_1$	2.0565	0.0565	0.3120	2.0635	0.0635	0.3405
	$\hat{\lambda}_1$	0.7044	0.0044	0.0351	0.7048	0.0048	0.0404
	$\hat{\alpha}_2$	2.5741	0.0741	0.3958	2.5795	0.0795	0.4499
	$\hat{\lambda}_2$	0.5029	0.0029	0.0232	0.5026	0.0026	0.0278
	$\hat{\theta}$	6.1043	0.1043	0.9188	5.9299	-0.0701	0.9050

Table 9: Estimation of the Parameters of Clayton-BGR Distribution: Case 3, $n=200$

θ		MLE			IFM		
		Mean	Bias	RMSE	Mean	Bias	RMSE
0.5	$\hat{\alpha}_1$	2.0282	0.0282	0.2139	2.0292	0.0292	0.2162
	$\hat{\lambda}_1$	0.7029	0.0029	0.0273	0.7030	0.0030	0.0277
	$\hat{\alpha}_2$	2.5396	0.0396	0.2811	2.5408	0.0408	0.2850
	$\hat{\lambda}_2$	0.5019	0.0019	0.0186	0.5019	0.0019	0.0189
	$\hat{\theta}$	0.5162	0.0162	0.1298	0.5142	0.0142	0.1290
1.5	$\hat{\alpha}_1$	2.0285	0.0285	0.2076	2.0292	0.0292	0.2162
	$\hat{\lambda}_1$	0.7030	0.0030	0.0260	0.7030	0.0030	0.0277
	$\hat{\alpha}_2$	2.5344	0.0344	0.2710	2.5376	0.0376	0.2837
	$\hat{\lambda}_2$	0.5008	0.0008	0.0360	0.5018	0.0018	0.0191
	$\hat{\theta}$	1.5275	0.0275	0.2234	1.5182	0.0182	0.2216
3	$\hat{\alpha}_1$	2.0283	0.0283	0.2038	2.0292	0.0292	0.2162
	$\hat{\lambda}_1$	0.7030	0.0030	0.0248	0.7030	0.0030	0.0277
	$\hat{\alpha}_2$	2.5350	0.0350	0.2637	2.5375	0.0375	0.2817
	$\hat{\lambda}_2$	0.5019	0.0019	0.0169	0.5019	0.0019	0.0190
	$\hat{\theta}$	3.0436	0.0436	0.3652	3.0174	0.0174	0.3604
6	$\hat{\alpha}_1$	2.0286	0.0286	0.2005	2.0292	0.0292	0.2162
	$\hat{\lambda}_1$	0.7031	0.0031	0.0237	0.7030	0.0030	0.0277
	$\hat{\alpha}_2$	2.5366	0.0366	0.2585	2.5394	0.0394	0.2824
	$\hat{\lambda}_2$	0.5020	0.0020	0.0160	0.5020	0.0020	0.0190
	$\hat{\theta}$	6.0743	0.0743	0.6463	5.9872	-0.0128	0.6326

5. Application of Real Data

In this section, we analyze a real data set of economic data in Egypt for World Bank national accounts data that analyzed by El-Sherpieny et al. (2018). We study the parameter estimation of the appropriate distribution of each data, where the correlation between the two variables (bivariate data) is positive correlation. And through this access to a fit model specialized in the study of weak relations and the extent of their impact and effectiveness. The economic data set: Exports of goods and services (x_1) and GDP growth (x_2). In this data, we cannot used this data in Shock model, because $x_1 > x_2$.

Genest et al. (2013) introduced Multiplier bootstrap-based goodness-of-fit test, this consideration leads naturally to Anderson–Darling-type statistics such as

$$R_n = n \int_0^1 \left[\frac{\hat{C}_n(u_1, u_2) - C_{\theta_n}(u_1, u_2)}{[C_{\theta_n}(u_1, u_2) (1 - C_{\theta_n}(u_1, u_2)) + \delta_m]^m} \right]^2 d\hat{C}_n(u_1, u_2)$$

Involving a consistent, rank-based estimator θ_n of θ , and tuning parameters $m \geq 0$ and $\delta_m \geq 0$.

We use the concludes of Genest to fit of Clayton copulas by R package, then $R_n = 0.2658$, $\hat{\theta} = 0.3912$ and $p - value = 0.2257$, while $p - value > 0.05$, then the data fit of the Clayton copula. This by using a parametric bootstrap N=10000 time and the empirical copula estimate.

Table 10: parameter estimation of marginal, standard error and goodness of fit

	x_1	x_2
$\hat{\alpha}$	5.47579 (1.8474)	1.7395 (0.4537)
$\hat{\lambda}$	0.06494 (0.0056)	0.2134 (0.0219)
K-S (p-value)	0.0979 (0.8993)	0.0547 (0.9999)

Table 11: The Estimates and the Corresponding Standard Deviation of Parameters of Clayton-BGR Distribution

	MLE		IFM	
	Coef	std	coef	Std
$\hat{\alpha}_1$	5.4509	1.847	5.47579	(1.847)
$\hat{\lambda}_1$	0.065	0.0055	0.06494	(0.0056)
$\hat{\alpha}_2$	1.7186	0.4487	1.7395	(0.4537)
$\hat{\lambda}_2$	0.2118	0.0220	0.2134	(0.0219)
$\hat{\theta}$	0.2458	0.2903	0.2407	0.2716

In the Table (11), the best method for of Estimation of the Clayton-BGR distribution is IFM method since it has the least standard deviation of the copula parameter.

Table 12: The Estimates of Parameters of Clayton-BGR and anther Distributions

	$\hat{\alpha}_1$	$\hat{\beta}_1$	$\hat{\alpha}_2$	$\hat{\beta}_2$	$\hat{\theta}$	AIC	BIC	HQIC	CAIC
Clayton-BGR	5.4509	0.065	1.7186	0.2118	0.2458	335.5796	342.7495	337.9168	337.9796
Clayton-BR	-	0.04285	-	0.19398	1.2567	355.0116	359.3136	356.4139	355.9005
Clayton-BW	4.4445	0.0397	2.6535	0.1726	0.3523	336.1476	343.3176	338.4848	338.5476
FGMBW	4.5225	0.0395	2.6953	0.1712	0.6712	335.617	342.789	337.9545	338.0172

The new distribution Clayton-BGR is more appropriate as compared to the Clayton-BR, Clayton-BW and FGM bivariate Weibull (FGMBW) distributions, where the Clayton-BGR distribution gives the lowest value for the AIC, BIC, AICC and H AIC statistics compared other models.

6. Conclusion

In this paper, we have proposed a class of bivariate GR distributions based on Clayton copula function. Moreover, we obtained the reliability functions for Clayton-BGR distributions; therefore, it can be used quite effectively in life testing data. Additionally, the new BGR model based on Clayton copula can be used as an alternative to different distributions as Clayton-BR, Clayton-BW and FGMBW. A comparison between different estimation methods of the bivariate GR distributions are concluded. The result show that the best method of estimation is IFM method.

References

- Abdel-Hady, D. H. (2013). Bivariate Generalized Rayleigh Distribution. *Journal of Applied Sciences Research*, 9(9), 5403-5411.
- Almetwally, E. M. (2019). Parameter Estimation of Bivariate Models under Some Censoring Schemes. Master Thesis, Faculty of Graduate Studies for Statistical Research, Cairo University.
- Almetwally, E. M., Almongy, H. M., & El-Sherpieny, E. S. A. (2019_b). Adaptive Type-II Progressive Censoring Schemes based on Maximum Product Spacing with Application of Generalized Rayleigh Distribution. *Journal of Data Science*, 17(4), 696-723.
- Almetwally, E. M., Muhammed, H. Z., & El-Sherpieny, E. S. A. (2019_a). Bivariate Weibull Distribution: Properties and Different Methods of Estimation. *Annals of Data Science*, 1-31.
- Burr, I. W. (1942). Cumulative frequency functions. *The Annals of Mathematical Statistics*, 13(2), 215-232.
- Clayton, D. G. (1978). A model for association in bivariate life tables and its application in epidemiological studies of familial tendency in chronic disease incidence. *Biometrika*, 65(1), 141-151.
- Elaal, M. K. A., & Jarwan, R. S. (2017). Inference of Bivariate Generalized Exponential Distribution Based on Copula Functions. *Applied Mathematical Sciences*, 11(24), 1155-1186.

El-Sherpieny, E. A., Muhammed, H. Z., & Almetwally, E. M. (2018). FGM Bivariate Weibull Distribution. In Proceedings of the Annual Conference in Statistics (53rd), Computer Science, and Operations Research, Institute of Statistical Studies and Research, Cairo University, 55-77.

Flores, A. Q. (2009). Testing copula functions as a method to derive bivariate Weibull distributions. In American Political Science Association (APSA), Annual Meeting.

Fredricks, G. A., & Nelsen, R. B. (2007). On the relationship between Spearman's rho and Kendall's tau for pairs of continuous random variables. Journal of Statistical Planning and Inference, 137(7), 2143-2150.

Genest, C., Huang, W., & Dufour, J.-M. (2013). A regularized goodness-of-fit test for copulas. Journal de la Société Française de Statistique 154, 64–77.

Kim, G., Silvapulle, M. J., & Silvapulle, P. (2007). Comparison of semiparametric and parametric methods for estimating copulas. Computational Statistics & Data Analysis, 51(6), 2836-2850.

Kundu, D., & Raqab, M. Z. (2005). Generalized Rayleigh distribution: different methods of estimations. Computational Statistics & Data Analysis, 49(1), 187-200.

Kundu, D., & Raqab, M. Z. (2015). Estimation of $R = P[Y < X]$ for three-parameter generalized Rayleigh distribution. Journal of Statistical Computation and Simulation, 85(4), 725-739.

Nadarajah, S., Afuecheta, E., & Chan, S. (2018). A compendium of copulas. Statistica, 77(4), 279-328.

Nelsen, R.B., (2006). *An Introduction to Copulas*. Springer, New York.

Ochi, M. K. (1978). Generalization of Rayleigh probability distribution and its application. Journal of Ship Research, 22(04), 259-265.

Raqab, M. Z., & Madi, M. T. (2011). Inference for the generalized Rayleigh distribution based on progressively censored data. Journal of Statistical Planning and Inference, 141(10), 3313-3322.

Sarhan, A. (2019). The Bivariate Generalized Rayleigh Distribution. Journal of Mathematical Sciences and Modelling, 2(2), 99-111.

Shubhashree, J. (2019). Estimation of Multicomponent System Reliability for a Bivariate Generalized Rayleigh Distribution. *International Journal of Computational and Theoretical Statistics*, 6(1), 51-63.

Sklar A (1973) Random variables, joint distributions, and copulas. *Kybernetika* 9:449-460.

Tomer, S. K., Gautam, R. S., & Kumar, J. (2014). Bayesian Estimation for the Generalized Rayleigh Distribution under Type-I Progressive Hybrid Censoring Scheme. In *Special Issue on Recent Statistical Methodologies and Applications* (2), 84-93.

Estimation of Chen distribution based on record values

Mohamed Mubarak¹ Shimaa Ashraf^{2,3}

Abstract According to record values and the associated statistics are interested in many real life applications involving data relating to auction and reverse auction, dynamic purchasing systems and production lines. In this article, we consider the problem of estimating unknown parameters of Chen distribution is obtained based on record values as upper record, lower record . We present the maximum likelihood (MLE) and Bayes estimators for two parameter of Chen distribution. We have examined Bayes estimates under symmetric loss function based on Monte carlo simulation of record values and numerical Computations. Furthermore the comparisons between the different estimators are given.

Keywords: Point and Bayesian estimation, Chen distribution, Record values, Monte Carlo Markov Chain, Gibbs-Sampling.

1 Introduction

Many authors have been studied record values and associated statistics. Record values and the associated statistics are interest and important in many real life applications involving data relating to weather, economic, Sports data "Olympic records or world records in sport" and life time testing studies. Chandler [9] he formulated the theory of record values as a model for successive extremes in a sequence of independently and identically random variables. Feller [11] gave some examples of record values. Resnick [15] discussed the asymptotic theory of records. Mubarak [13] estimate parameters of frechet distribu-

¹Mathematics department. Faculty of science. El-Minia University. El-Minia. Egypt. E-mail: mubarak_7061111@yahoo.com

²Mathematics department. Faculty of science. El-Minia University. El-Minia. Egypt.

³Corresponding author. E-mail address:sh_ashraf1991@yahoo.com

tion based on record. Interested readers may refer to [2].

For other recent examples, we refer the readers to [16]. For a review of developments in this characterizing distribution via their two parameters has along history. Some recently published example include: statistical inference about the shape parameter of new two parameter Bathtub-Shaped life time distribution. Bayesian estimation based on progressive Type-II Censoring from two parameter bathtub-shaped life time model: an Markov chain Monte carlo approach [1]. Monte Carlo estimation of Bayesian credible and HPD intervals. Estimating the parameters of bathtub-shaped distribution under progressive type-II censoring [14]. Estimation of two-parameter bathtub-shaped life time distribution with progressive censoring [17]. Finally, exponentiated Chen distribution: Properties and estimation [18], Estimation and predication for Chen distribution with bathtub-shape under progressive censoring [19].

Let $(X_{1:n}, X_{2:n}, X_{3:n}, \dots, X_{n:n})$ be a sequence of independent and identically distributed (iid), random variables with Cumulative distribution function $F(x)$ and probability density $f(x)$. With simplify, we rewrite $(X_1, X_2, \dots, X_n)$. Set $Y_n = \min(X_1, X_2, \dots, X_n)$ for $n \geq 1$, we say Y_j is a lower record value of this sequence if $Y_j < Y_{j-1}$ for $j > 1$. Therefore, Y_1 is a lower record value. Also, $Y_n = \max(X_1, X_2, \dots, X_n)$ for $n \geq 1$, we say Y_j is a upper record value of this sequence if $Y_j > Y_{j-1}$ for $j > 1$. Hence, Y_1 is a upper record value. Let $X_{L_n} = \min\{j | j > L_{n-1}, Y_j < Y_{L_{n-1}}, n \geq 2\}$ with L_1 denote the times of lower record values. Similarly, Let $X_{U_n} = \min\{j | j > U_{n-1}, Y_j > Y_{U_{n-1}}, n \geq 2\}$ with $U_1 = 1$ denote the times of upper record values. For comprehensive accounts of the theory and applications of record values, we refer the readers to [2]. The sequence $\{X_{U_n}\}_{n=1}^{\infty}$ is known as upper record values and the sequence $\{U_n\}_{n=1}^{\infty}$ is known as record times sequence [3]. In the following, we give some preliminaries. Let X_{U_m} and X_{U_n} for $m < n$ denote the upper record statistics from a given family. Let X_{L_m} and X_{L_n} for $m < n$ denote the corresponding lower record values. The joint probability density function of X_{U_m} and X_{U_n} is given by Arnold et al. [4]. In this article, we will discuss the estimation unknown two parameters from Chen distribution based on three different records. Let $(X_1, X_2, \dots, X_n)$ be a sequence of (iid) random variables from Chen distribution with (cdf) and (pdf) are given, respectively, as follows,

$$F(X) = 1 - e^{\lambda(1-e^{X^\beta})}, \quad X > 0, \quad (1)$$

and

$$f(X, \lambda, \beta) = \lambda \beta X^{\beta-1} e^{X^\beta + \lambda(1-e^{X^\beta})}, \quad 0 < X < \infty, \lambda, \beta > 0, \quad (2)$$

where λ and β are the shape and scale parameters respectively.

2 Estimation of parameters based on upper Record:

According to the upper record values it will be used in auction for private and general sectors. In this section, we used maximum likelihood estimation and Bayesian estimation with square loss function to estimate the parameters λ and β for Chen distribution. See table (1,4).

2.1 Maximum Likelihood Estimation

Let $\underline{X} = (X_{u_1}, X_{u_2}, X_{u_3}, \dots, X_{u_n})$ be the observation of n upper record value from the Chen distribution. The joint distribution of $\underline{X} = (X_{u_1}, X_{u_2}, \dots, X_{u_n})$ is defined according to Arnold et al.[6] as

$$L(\underline{X}, \lambda, \beta) = f(X_{u_n}) \prod_{i=1}^{n-1} \frac{f(X_{u_i})}{1 - F(X_{u_i})}. \quad (3)$$

The likelihood function with simplify of upper observation records $(X_1, X_2, \dots, X_n)$ will be written as follows

$$L(\underline{X}, \lambda, \beta) = \lambda^n \beta^n \left(\prod_{i=1}^n X_i \right)^{\beta-1} e^{\lambda(1-e^{X_n^\beta}) + \sum_{i=1}^n X_i^\beta}. \quad (4)$$

Hence, the log-likelihood function becomes,

$$L^*(\underline{X}, \lambda, \beta) = n \ln \lambda \beta + (\beta - 1) \sum_{i=1}^n \ln X_i + \sum_{i=1}^n X_i^\beta + \lambda(1 - e^{X_n^\beta}). \quad (5)$$

By differentiation with respect to λ and β and equal to zero, to obtain the estimator's λ and β as follows,

$$\hat{\lambda} = \frac{n}{e^{X_n^\beta} - 1}, \quad (6)$$

and

$$\frac{n}{\beta} + \sum_{i=1}^n (1 + X_i^\beta) \ln X_i - \lambda e^{X_n^\beta} X_n^\beta \ln X_n = 0. \quad (7)$$

Then, $\hat{\beta}$ can be obtained as the solution of the non linear above equation. Then we can rewrite as

$$\hat{\beta} = -n \left(\sum_{i=1}^n (1 + X_i^\beta) \ln X_i + \frac{n e^{X_n^\beta} X_n^\beta \ln X_n}{(1 - e^{X_n^\beta})} \right)^{-1}. \quad (8)$$

So, we solve it by simulation of fixed point for function $\psi(\beta) = \beta$. We applying algorithm of fixed point $\psi(\beta_k) = \beta_{k+1}$ and so on which stope when $|\beta_{k+1} - \beta_k| < \epsilon$ such that $\lim_{n \rightarrow \infty} \epsilon = 0$ (sufficiently small), where

$$\psi(\beta) = -n \left[\sum_{i=1}^n (1 + X_i^\beta) \ln X_i + \frac{ne^{X_n^\beta} X_n^\beta \ln X_n}{(1 - e^{X_n^\beta})} \right]^{-1}. \quad (9)$$

Once to obtain $\hat{\beta}$ and substituting in equation (6) to obtain $\hat{\lambda}$.

One can see that, the statistic $T_n = e^{X_n^\beta} - 1$ where X_n is maximum of random sample (upper record) satisfied the following propositions

Proposition 1:

Let $T_n = e^{X_n^\beta} - 1$ be a statistic which is sufficient and complete for the parameter λ . Then it distributed as exponential with density function as follows

$$f(T_n | \lambda) = \lambda e^{-\lambda T_n}, \quad T_n > 0.$$

Proposition 2:

Let $\hat{\lambda} = \frac{n}{T_n}$ and $T_n = e^{X_n^\beta} - 1$ then $\hat{\lambda}^* = \frac{\hat{\lambda}}{n\varphi(\beta)}$ is unbiased estimator where

$$\varphi(\beta) = \int_0^\infty \frac{e^{-\lambda T_n}}{T_n} dT_n.$$

Proposition 3:

Let $\frac{1}{\lambda}$ and $\frac{1}{\lambda^2}$ be mean and variance respectively for the statistic T_n . Then $n\lambda$ and $n^2\lambda^2$ are mean and variance for the estimator $\hat{\lambda}$ respectively and the mean square error of the estimator $\hat{\lambda}$ is $(2n - 1)\lambda^2$.

2.2 Bayesian Estimation

In this section, we study Bayes estimators for the unknown parameters λ and β under square error loss function (SELF). Assuming that, the parameters λ and β have the following independent Gamma conjugate prior distribution of the form

$$\Pi_1(\lambda) \propto \lambda^{b-1} e^{-a\lambda}, \quad \lambda > 0,$$

$$\Pi_2(\beta) \propto \beta^{d-1} e^{-c\beta}, \quad \beta > 0.$$

So that, the joint Prior distribution of λ and β have the following form,

$$\Pi(\lambda, \beta) \propto \lambda^{b-1} \beta^{d-1} e^{-a\lambda} e^{-c\beta}, \quad \lambda, \beta, a, b, c, d > 0. \quad (10)$$

The posterior distribution of λ and β can be written as,

$$\Pi^*(\lambda, \beta | \underline{X}) = \frac{L(\lambda, \beta | \underline{X}) \Pi(\lambda, \beta)}{\int_0^\infty \int_0^\infty L(\lambda, \beta | \underline{X}) \Pi(\lambda, \beta) d\lambda d\beta}.$$

Then, applying Bayes theorem, the joint posterior density function of (λ, β) where $\lambda, \beta > 0$, become as follows,

$$\Pi^*(\lambda, \beta | \underline{X}) = \frac{\lambda^{n+b-1} \beta^{n+d-1} (\prod_{i=1}^n X_i)^{\beta-1} e^{-c\beta - \lambda(e^{X_n^\beta} - 1 + a) + \sum_{i=1}^n X_i^\beta}}{\Gamma(n+b) \Psi(a, b, c, d, X_n)}, \quad (11)$$

where

$$\Psi(a, b, c, d, X_n) = \int_0^\infty \frac{\beta^{n+d-1} (\prod_{i=1}^n X_i)^{\beta-1} e^{\sum_{i=1}^n X_i^\beta} e^{-c\beta}}{(1 - a - e^{X_n^\beta})^{n+b}} d\beta. \quad (12)$$

The Bayes estimate is the mean of posterior density. Therefore, from (8) the Bayes estimators of λ and β are given

$$\begin{aligned} \hat{\lambda}_{BS} &= E(\lambda) = E(\lambda | \underline{X}) = \int_0^\infty \int_0^\infty \lambda \pi^*(\lambda, \beta | \underline{X}) d\beta d\lambda \\ &= \frac{(n+b) \Psi(b+1, c, d, a, X_n)}{\Psi(a, b, c, d, X_n)}. \end{aligned} \quad (13)$$

And also, with similar method, we obtain

$$\hat{\beta}_{BS} = \frac{\Psi(a, b, c, d+1, X_n)}{\Psi(a, b, c, d, X_n)}. \quad (14)$$

But, we can not obtain result for above integration in closed form. So, we use numerical integration technique in maple software in order to solve the above integration.

3 Estimation of parameters based on lower record

It has been observed that the lower record values which used as reverse auction and dynamic purchasing systems for some fields. In this section, we used maximum likelihood estimation and Bayesian estimation with square loss function to estimate the parameters λ and β for Chen distribution. See table (2,4).

3.1 Maximum Likelihood Estimation

Let $\underline{X} = (X_{L_1}, X_{L_2}, X_{L_3}, \dots, X_{L_n})$ be the lower record value of size n from the Chen distribution with parameters λ and β . The likelihood function for observed lower records

$\underline{X} = (X_{L_1}, X_{L_2}, \dots, X_{L_n})$ was given by see Arnold et al. [7]

$$L(\underline{X}, \lambda, \beta) = f(X_{L_n}) \prod_{i=1}^{n-1} \frac{f(X_{L_i})}{F(X_{L_i})}.$$

The likelihood function with simplify of lower observation records $(X_1, X_2, \dots, X_n)$ will be written as follows

$$L(\underline{X}, \lambda, \beta) = \lambda^n \beta^n \left(\prod_{i=1}^n X_i \right)^{\beta-1} \prod_{i=1}^{n-1} \left(1 - e^{\lambda(1-e^{X_i^\beta})} \right)^{-1} e^{\sum_{i=1}^n X_i^\beta + \lambda(1-e^{X_i^\beta})}. \quad (15)$$

Therefore, the log-likelihood function becomes

$$L^*(\underline{X}, \lambda, \beta) = n \log \lambda + n \log \beta + (\beta-1) \sum_{i=1}^n \log X_i + \sum_{i=1}^n \left(X_i^\beta + \lambda(1 - e^{X_i^\beta}) \right) - \sum_{i=1}^{n-1} \log \left(1 - e^{\lambda(1-e^{X_i^\beta})} \right). \quad (16)$$

By differentiation with respect to λ and β and equal to zero, to obtain the estimator's λ and β as follows,

$$\frac{n}{\lambda} + \sum_{i=1}^n (1 - e^{X_i^\beta}) + \sum_{i=1}^{n-1} \frac{e^{\lambda(1-e^{X_i^\beta})} (1 - e^{X_i^\beta})}{(1 - e^{\lambda(1-e^{X_i^\beta})})}, \quad (17)$$

$$\frac{n}{\beta} + \sum_{i=1}^n \ln X_i + (1 - \lambda) \sum_{i=1}^n X_i^\beta \ln X_i - \lambda \sum_{i=1}^{n-1} \frac{X_i^\beta e^{X_i^\beta + \lambda(1-e^{X_i^\beta})}}{1 - e^{\lambda(1-e^{X_i^\beta})}} = 0. \quad (18)$$

Then, $\hat{\beta}$ and $\hat{\lambda}$ can be obtained as the solution of the non linear above equations. Then we can rewrite as

$$\hat{\lambda} = -n \left[1 - e^{X_n^\beta} + \sum_{i=1}^{n-1} \frac{1 - e^{X_i^\beta}}{1 - e^{\lambda(1-e^{X_i^\beta})}} \right]^{-1}, \quad (19)$$

$$\hat{\beta} = -n \left[\sum_{i=1}^n \ln X_i + (1 - \lambda) \sum_{i=1}^n X_i^\beta \ln X_i - \lambda \sum_{i=1}^{n-1} \frac{X_i^\beta e^{X_i^\beta + \lambda(1-e^{X_i^\beta})}}{1 - e^{\lambda(1-e^{X_i^\beta})}} \right]^{-1}. \quad (20)$$

We solve it by simulation of fixed point for function $\psi(\theta) = \theta$ where $\theta = (\lambda, \beta)$. We applying algorithm of fixed point $\psi(\theta_k) = \theta_{k+1}$ and so on which stoppe when $|\theta_{k+1} - \theta_k| < \epsilon$ such that $\lim_{n \rightarrow \infty} \epsilon = 0$ (sufficiently small), where

$$\begin{aligned} \psi(\lambda) &= -n \left[1 - e^{X_n^\beta} + \sum_{i=1}^{n-1} \frac{1 - e^{X_i^\beta}}{1 - e^{\lambda(1-e^{X_i^\beta})}} \right]^{-1}, \\ \psi(\beta) &= -n \left[\sum_{i=1}^n \ln X_i + (1 - \lambda) \sum_{i=1}^n X_i^\beta \ln X_i - \lambda \sum_{i=1}^{n-1} \frac{X_i^\beta e^{X_i^\beta + \lambda(1-e^{X_i^\beta})}}{1 - e^{\lambda(1-e^{X_i^\beta})}} \right]^{-1}. \end{aligned} \quad (21)$$

Once to obtain $\hat{\lambda}$ and $\hat{\beta}$ with (λ_0, β_0) and repeat the fixed point algorithm T items. Then,

$$\hat{\lambda}^* = \frac{1}{T} \sum_{i=1}^T \hat{\lambda}_i, \quad \hat{\beta}^* = \frac{1}{T} \sum_{i=1}^T \hat{\beta}_i. \quad (22)$$

3.2 Bayesian Estimation

In this section, we study Bayes estimators for the unknown parameters λ and β under symmetric loss function. Assuming that, the parameters λ and β have the following independent Gamma Conjugate prior distribution as

$$\Pi_1(\lambda) \propto \lambda^{b-1} e^{-a\lambda}, \quad \lambda > 0,$$

$$\Pi_2(\beta) \propto \beta^{d-1} e^{-c\beta}, \quad \beta > 0.$$

So, the joint Prior distribution of λ and β have the following form,

$$\Pi(\lambda, \beta) \propto \lambda^{b-1} e^{-a\lambda} \beta^{d-1} e^{-c\beta}, \quad \lambda, \beta, a, b, c, d > 0. \quad (23)$$

Then the Posterior distribution λ and β can be written as,

$$\Pi^*(\lambda, \beta | \underline{X}) = \frac{L(\underline{X}, \lambda, \beta) \Pi(\lambda, \beta)}{\int_0^\infty \int_0^\infty L(\underline{X}, \lambda, \beta) \Pi(\lambda, \beta) d\lambda d\beta}.$$

Therefore,

$$\Pi^*(\lambda, \beta | \underline{X}) \propto \lambda^{n+b-1} \beta^{n+d-1} \left(\prod_{i=1}^n X_i \right)^{\beta-1} \prod_{i=1}^{n-1} \left(1 - e^{\lambda(1-e^{X_i^\beta})} \right)^{-1} e^{-c\beta - a\lambda + \sum_{i=1}^n X_i^\beta + \lambda(1-e^{X_i^\beta})}. \quad (24)$$

Then the Bayes estimate of any function of λ and β say $\eta(\lambda, \beta)$ under the squared error loss function is given by

$$\hat{\eta}_{BS}(\lambda, \beta) = \int_0^\infty \int_0^\infty \eta(\lambda, \beta) \Pi^*(\lambda, \beta | \underline{X}) d\lambda d\beta. \quad (25)$$

Unfortunately, from equation (25) this integration can not be computed explicitly. Therefore, we adopt the Gibbs sampling technique to compute the Bayes estimates and the corresponding credible intervals of the unknown parameters. The conditional posterior functions of λ and β are respectively, as follows:

$$\begin{aligned} \Pi_1^*(\beta | \lambda, \underline{X}) &\propto \beta^{n+d-1} \left(\prod_{i=1}^n X_i \right)^{\beta-1} \prod_{i=1}^{n-1} \left(1 - e^{\lambda(1-e^{X_i^\beta})} \right)^{-1} e^{-c\beta + \sum_{i=1}^n X_i^\beta + \lambda(1-e^{X_i^\beta})}, \\ \Pi_2^*(\lambda | \beta, \underline{X}) &\propto \lambda^{n+b-1} \prod_{i=1}^{n-1} \left(1 - e^{\lambda(1-e^{X_i^\beta})} \right)^{-1} e^{-a\lambda + \sum_{i=1}^n \lambda(1-e^{X_i^\beta})}. \end{aligned} \quad (26)$$

It is clear that, direct generation of a pseudo random number from the posterior density functions of λ and β is not easy, instead, we use the Metropolis Hastings method (MCMC). Therefore, the algorithm of Gibbs sampling is as follows:

1. Take some initial value of λ and β such as λ_0 and β_0 .
2. Set $t = 1$.
3. Generate λ_t from $\Pi_2^*(\lambda|\beta_{t-1}, \underline{X})$ and β_t from $\Pi_1^*(\beta|\lambda_{t-1}, \underline{X})$.
4. Set $t = t + 1$.
5. Repeat steps 2-4 into T times in order to obtain Bayes estimate of λ and β as $\hat{\lambda}_{BS} = \frac{1}{T} \sum_{t=1}^T \lambda_t$ and $\hat{\beta}_{BS} = \frac{1}{T} \sum_{t=1}^T \beta_t$.

4 Asymptotic variance and covariance matrix under record values

The asymptotic variance and covariance matrix of maximum likelihood estimates are given by the elements of the inverse of Fisher information matrix as

$$I_{ij}(\underline{\theta}) \cong E \left(-\frac{\partial^2 \ln L}{\partial \theta_i \partial \theta_j} \right)$$

Unfortunately, the exact mathematical expressions for the previous expectation are very difficult to obtain, therefore Fisher information matrix is given by

$$I_{ij}(\underline{\theta}) \cong \left(-\frac{\partial^2 \ln L}{\partial \theta_i \partial \theta_j} \right)$$

which is obtained by dropping the expectation on operation E and replaced λ, β with $\hat{\lambda}, \hat{\beta}$ respectively, Cohen [10]. The asymptotic variance and covariance matrix F for the maximum likelihood estimates can be written as follows

$$F = \begin{pmatrix} -\frac{\partial^2 \ln L}{\partial \lambda^2} & -\frac{\partial^2 \ln L}{\partial \lambda \partial \beta} \\ -\frac{\partial^2 \ln L}{\partial \beta \partial \lambda} & -\frac{\partial^2 \ln L}{\partial \beta^2} \end{pmatrix}$$

4.1 lower record

$$L_{\lambda\lambda}^*(\lambda, \beta|\underline{X}) = -\frac{n}{\lambda^2} - \sum_{i=1}^{n-1} \frac{(1 - e^{x_i^\beta})^2 e^{\lambda(1-e^{x_i^\beta})}}{(1 - e^{\lambda(1-e^{x_i^\beta})})^2}. \quad (27)$$

$$\begin{aligned} L_{\beta\beta}^*(\lambda, \beta|\underline{X}) &= -\sum_{i=1}^n x_i^\beta e^{x_i^\beta} \ln(x_i) \\ &\quad - \sum_{i=1}^{n-1} \frac{\lambda(1 - e^{x_i^\beta}) x_i \ln(x_i) e^{\lambda(1-e^{x_i^\beta})+x_i^\beta}}{1 - e^{\lambda(1-e^{x_i^\beta})}} \\ &\quad + \frac{\lambda(1 - e^{x_i^\beta}) x_i \ln(x_i) e^{2\lambda(1-e^{x_i^\beta})+x_i^\beta}}{(1 - e^{\lambda(1-e^{x_i^\beta})})^2} + \frac{x_i \ln(x_i) e^{\lambda(1-e^{x_i^\beta})+x_i^\beta}}{1 - e^{\lambda(1-e^{x_i^\beta})}}. \end{aligned} \quad (28)$$

$$\begin{aligned} L_{\lambda\beta}^*(\lambda, \beta|\underline{X}) &= L_{\beta\lambda}^*(\lambda, \beta|\underline{X}) = -\frac{n}{\beta^2} + \sum_{i=1}^n x_i^\beta \ln^2(x_i) \\ &\quad - \lambda \sum_{i=1}^n e^{x_i^\beta} \ln^2(x_i) (x_i^{2\beta} - x_i^\beta) - \sum_{i=1}^{n-1} -\frac{\lambda^2 x_i^{2\beta} \ln^2(x_i) e^{2\lambda(1-e^{x_i^\beta})+2x_i^\beta}}{(1 - e^{\lambda(1-e^{x_i^\beta})})^2} \\ &\quad + \frac{\lambda x_i^\beta \ln^2(x_i) e^{\lambda(1-e^{x_i^\beta})+x_i^\beta}}{1 - e^{\lambda(1-e^{x_i^\beta})}} \\ &\quad + \frac{\lambda x_i^\beta \ln(x_i) e^{\lambda(1-e^{x_i^\beta})+x_i^\beta} (x_i^\beta \ln(x_i) - \lambda e^{x_i^\beta} x_i^\beta \ln(x_i))}{1 - e^{\lambda(1-e^{x_i^\beta})}}. \end{aligned} \quad (29)$$

4.2 upper record

$$L_{\lambda\lambda}^*(\lambda, \beta|\underline{X}) = -\frac{n}{\lambda^2}. \quad (30)$$

$$L_{\beta\beta}^*(\lambda, \beta|\underline{X}) = -\frac{n}{\beta^2} + \sum_{i=1}^n x_i^\beta \ln(x_i) - \lambda x_n^{2\beta} e^{x_n^\beta} \ln^2(x_n) - \lambda x_n^\beta e^{x_n^\beta} \ln^2(x_n). \quad (31)$$

$$L_{\lambda\beta}^*(\lambda, \beta|\underline{X}) = L_{\beta\lambda}^*(\lambda, \beta|\underline{X}) = -x_n^\beta e^{x_n^\beta} \ln(x_n). \quad (32)$$

The asymptotic normality of maximum likelihood estimates can be used to compute the approximate confidence intervals for parameters λ and β therefore $(1 - \alpha)100$ confidence intervals for parameters become as

$$\hat{\lambda} \pm Z_{\frac{\alpha}{2}} \sqrt{Var(\hat{\lambda})}, \quad \hat{\beta} \pm Z_{\frac{\alpha}{2}} \sqrt{Var(\hat{\beta})}$$

where $Z_{\frac{\alpha}{2}}$ is the percentile, $Var(\hat{\lambda}) = L_{\lambda\lambda}^*(\lambda, \beta|\underline{X})$ and $Var(\hat{\beta}) = L_{\beta\beta}^*(\lambda, \beta|\underline{X})$. See table (1-3).

5 Simulation study

In previous section we consider point, interval and Bayes estimation based on different upper and lower record values. In this section we employ a monte Carlo simulation study to compare the performance of proposed methods of estimation namely likelihood and Bayesian estimation methods. We simulate a generating data from Chen distribution with parameters $\lambda = 2$ and $\beta = 1.5$ based on upper and lower record values for different choice of n and various record values k . We suppose 10000 number of simulation and generating record values $k = 5, 7, 15, 25, \dots$ and prior parameter of λ, β : $a = 0.5, b = 4, c = 0.5, d = 3$ the average estimate, interval estimate (Asymptotic confidence interval) and mean square error (MSE) are reported for both type method of estimation and different record values. See tables (1-4).

6 Conclusion

In this paper we have studied a two parameter distribution with bathtub shape or increasing hazard rate function (Chen distribution) based on record values, we computed different estimates for its unknown parameters using maximum likelihood and Bayesian approaches under different record values. The asymptotic confidence intervals are constructed from observed Fisher information matrix obtained. The Bayes estimates of unknown parameters are obtained with respect to gamma prior distribution. Simulation results indicated that Bayes estimates perform quite good compared to the corresponding MLEs. Also, we observed that better estimates are obtained with increase in duration of the experimentation. We analyze a real data set originally reported in Hand et al. (1994)[12]. The data represent graft survival times in months of 148 renal transplant patients. We first verify whether a Chen distribution with bathtub shape is appropriate for making inference from the considered data. We fitted three more distribution namely generalized exponential, Weibull and exponential to this data set for comparison purposes the maximum likelihood estimates of unknown parameters of competing models along with values of different model selection Criteria like Kolmogorov-Smirnov (K-S) statistics. The tabulated values indicate that proposed model fit the data set really good compared to the other distributions.

Distribution	λ	β	$K - S$
Chen	0.2650	0.6358	0.0603
GExponential	0.8927	0.5348	0.1334
Weibull	1.0260	0.5615	0.1189
Exponential		0.5744	0.1232

References

- Ahmed EA.(2014), " Bayesian estimation based on progressive Type-II censoring from two-parameter bathtub-shapedlifetime model: an Markov chain Monte Carlo approach". J Appl Stat. 2014;41:752-768.
- Ahsanullah M. (1995), "Record Statistics". New York: Nova Science Publishers. MR1443904.
- Ahsanullah M. (2004), "Record Values-Theory and Applications". New York: Univ. Press of America.
- Ahsanullah M. and Raqab M. Z. (2006), "Bounds and Characterizations of Record Statistics". Hauppauge, NY: Nova Science Publishers. MR2294882.
- Arnold B.C. and Balakrishnan N. (1989), " Relations bounds and approximations for Order Statistics". LectureNotes in Statistics. 53 Springer -Verlag. New York.
- Arnold B.C., Balakrishnan N. and Nagaraja H. N. (1992), " A first course in order statistics". Wiley, New York.
- Arnold B. C., Balakrishnan N. and Nagaraja H. N. (1998), "Records". New York:Wiley. MR1628157.
- Arslan G., Ahsanullah M. and Bairamov I. G. (2005a), "On characteristic properties of the uniform distribution". Sankhy (a) 67, 715-721. MR2283011.
- Chandler K. N. (1952), "The distribution and frequency of record values". Journal of the Royal StatisticalSociety, Ser. B 14, 220-228. MR0053463.
- Cohen, A. C. (1965)," Maximum likelihood estimation in the Weibull distribution based on complete and on censored samples". Technometrics, 5, 579-588.

Feller W. (1966), "An introduction to probability theory and its applications". Vol. 2, John Wiley and Sons, NewYork.

Hand, D. J., Daly, F., Lunn, A. D., Mcconway, K. J. and Ostrowski, E., (1994), " A Handbook of small data sets". London: Chapman 8 Hall, 255.

Mubarak M. (2011), " Estimation of the frechet distribution parameters based on record values". Arabian Journal for science and engineering (AJSE). vol.36, No 8, pp.4067-4073.

Manoj Kumar Rastogi, Yogesh Mani Tripathi and Shuo-Jye Wu (2012), " Estimatingthe parameters of a bathtub-shaped distribution under progressive type-II censoring". Journal of Applied Statistics, 39:11, 2389-2411.

Resnick S.I. (1973), "Record values and maxima". Ann. Probab. 650-662.

Sultan K. S. and Abd El-Mougod G. A. (2006), " Characterizations of general classes of doubly truncated distributions based on record values". JPSS Journal of Probability and Statistical Sciences4, 65-72. MR2240302.

Shuo-Jye Wu (2008)," Estimation of the two-parameter bathtub-shaped lifetimedistribution with progressive censoring". Journal of Applied Statistics, 35:10, 1139-1150.

Sanku Dey, Devendra Kumar, Pedro L. Ramos and Francisco Louzada (2016),"Exponentiated Chen distribution: Properties and estimation". Communications in Statistics - Simulation and Computation.

Tanmay Kayal, Yogesh Mani Tripathi, Devendra Pratap Singh and Manoj Kumar Rastogi (2016)," Estimation and prediction for Chen distribution with bathtub shapeunder progressive censoring". Journal of Statistical Computation and Simulation.

Table 1: Average estimated values, interval estimates and MSEs of Chen distribution with $\lambda = 2$ and $\beta = 1.5$ under upper record

MLE					Bayes		
(n,k)	MSE			Asy. CI	MSE		Asy. CI
(20,5)	λ	3.4930	(10.5090)	(0.8177, 11.9458)	-	-	-
	β	2.6763	(4.07834)	(1.0825, 7.1734)	-	-	-
(28,7)	λ	2.8774	(6.3112)	(1.2553, 7.5394)	2.8903	(6.3073)	(1.2934, 7.5393)
	β	1.5456	(0.4014)	(0.9738, 7.6615)	2.8704	(7.4574)	(1.1853, 7.5484)
(40,10)	λ	2.7317	(11.5336)	(0.8707, 13.3460)	4.5689	(81.5015)	(0.1995, 27.5593)
	β	2.5943	(3.9313)	(1.2683, 8.1524)	3.6374	(55.2975)	(0.9979, 19.7312)
(60,15)	λ	1.7183	(1.2045)	(0.8019, 4.3422)	2.0694	(2.0689)	(0.2971, 6.3887)
	β	2.3094	(2.7914)	(1.2043, 7.1124)	1.7333	(1.1367)	(0.9615, 4.3380)
(100,25)	λ	1.9721	(0.2179)	(1.4250, 2.9232)	0.8077	(1.8076)	(0.0876, 2.1415)
	β	0.9892	(0.0331)	(0.4968, 1.5585)	1.2475	(0.1749)	(0.8934, 2.1978)

Table 2: Average estimated values, interval estimates and MSEs of Chen distribution with $\lambda = 2$ and $\beta = 1.5$ under lower record

MLE					Bayes		
(n,k)	MSE			Asy. CI	MSE		Asy. CI
(20,5)	λ	3.5413	(10.307)	(1.1743, 12.064)	-	-	-
	β	2.2370	(2.8102)	(0.8281, 6.2856)	-	-	-
(28,7)	λ	2.8773	(6.3113)	(0.6244, 3.1376)	2.8903	(6.3074)	(1.2934, 7.5482)
	β	1.5456	(0.4014)	(1.2084, 4.9394)	2.8704	(7.4573)	(1.1851, 7.5483)
(40,10)	λ	2.4503	(3.0194)	(1.2084, 4.9393)	2.4714	(3.0082)	(1.2661, 4.9424)
	β	1.2896	(0.3237)	(0.5507, 2.5157)	2.4213	(3.7083)	(1.1103, 4.9332)
(60,15)	λ	2.0421	(0.3447)	(1.2641, 3.4406)	2.0687	(0.3328)	(1.4175, 3.4469)
	β	1.0768	(0.3186)	(0.5219, 1.8963)	1.9973	(0.6074)	(1.3113, 3.4164)
(100,25)	λ	1.9721	(0.2179)	(1.4250, 2.9232)	2.0025	(0.2061)	(1.4802, 2.9338)
	β	0.9892	(0.0331)	(0.4968, 1.5585)	1.9184	(0.4042)	(1.3052, 2.8874)

Table 3: MLE and Bayesian estimates for unknown parameters respectively under different censoring with the same string and different record for each censoring

(n,m,k)	$\hat{\lambda}_0 \quad \hat{\beta}_0$		<u>U-record</u> $\hat{\lambda} \quad \hat{\beta}$		<u>L-record</u> $\hat{\lambda} \quad \hat{\beta}$	
	$\hat{\lambda}_0$	$\hat{\beta}_0$	$\hat{\lambda}$	$\hat{\beta}$	$\hat{\lambda}$	$\hat{\beta}$
(1000,100,3)	0.5	1.5	1.0111	1.1537	0.9839	1.9338
	1.5	0.5	1.0102	0.9983	1.5154	1.0067
	1	2	0.9721	1.8487	1.4425	1.8677
	2	1	2.0880	1.1537	1.9837	1.6371
(1000,100,5)	0.5	1.5	0.8111	1.4537	0.7839	1.7338
	1.5	0.5	1.2102	0.8973	1.5954	0.8067
	1	2	0.9421	1.9987	1.5625	2.2677
	2	1	1.9110	1.1207	1.9967	1.3371
(1000,100,7)	0.5	1.5	1.1413	2.4633	1.7078	2.9541
	1.5	0.5	1.3091	0.7968	2.0896	0.7450
	1	2	0.9391	2.1968	1.6896	2.7450
	2	1	1.9003	1.0634	2.0896	1.2350
(1000,100,3)	0.5	1.5	0.9981	1.1529	0.7839	1.8548
	1.5	0.5	1.3102	0.8983	1.6154	0.7085
	1	2	0.9926	1.9847	1.3828	2.4273
	2	1	1.9186	0.9098	2.2837	1.4372
(1000,100,5)	0.5	1.5	0.8719	1.3437	0.6837	1.7933
	1.5	0.5	1.3982	0.7548	1.5957	0.6787
	1	2	0.9721	1.8487	1.4425	2.2677
	2	1	1.9891	1.1537	2.0893	1.2971
(1000,100,7)	0.5	1.5	0.7113	1.3043	0.6058	1.6581
	1.5	0.5	1.3091	0.1968	1.7089	0.6050
	1	2	0.9903	1.9198	1.3086	2.0450
	2	1	2.1410	1.0734	1.9896	1.0745

Estimation of the Lifetime Performance Index with Burr Type III Distribution Under Type II Censoring

Amal S. Hassan¹, Salwa M. Assar² and Amany S. Selmy³

^{1,2,3}

Department of Mathematical Statistics, Faculty of Graduate Studies for Statistical Research, Cairo University, Egypt.

Abstract

Assessing the lifetime performance index of products is one of the most important topics in the manufacturing industries. Uniformly minimum variance unbiased estimator (UMVUE) of the lifetime performance index based on the right Type II censored sample for Burr Type III distribution is obtained. Then, the UMVUE of lifetime performance index is utilized to develop the new hypothesis testing procedure in the condition of known specification limit. One practical example and Monte Carlo simulation are given to illustrate the use of testing procedure under given significance level.

Keywords: Process capability index, Lifetime performance index, Burr Type III distribution, Type-II censoring, Uniformly minimum variance unbiased estimator.

1. Introduction

The globalization of business activities has intensified competitions among business organizations and resulted in quality consciousness, manufacture defect-free products and offer greater reliability of their products. These results are then compared with those of industry leaders for competitive bench marking. One metric popularly used is the Process capability index (PCI). Process capability analysis has the following benefits: (i) continuously monitoring the process quality through process capability indices in order to assure the products manufactured are conforming to the specifications; (ii) supplying information on product design and process quality improvement for engineers and designer; and (iii) providing the basis for reducing the cost of product failures (see Pan and Wu 1997).

Controlling and improving quality of products is important for many companies. In particular, lifetime of products is a relevant quality characteristic that producers have to monitor. As a measure of quality for life times, Montgomery (1985) and Kane (1986) proposed the process capability index, C_L for evaluating the lifetime performance of electronic components, where L is the lower specification limit. Statistical inferences for C_L on the basis of complete sample for different lifetime distributions have been considered in the literature. For example, Tong et al. (2002) obtained the UMVUE for C_L and considered hypothesis testing procedure for the one parameter exponential distribution. Chen et al. (2002) also used the UMVUE of C_L to develop the confidence interval under an exponential distribution.

Statistical studies for C_L based on Type II censored samples have been considered by several researchers. For instance, Hong et al. (2007) constructed a maximum likelihood estimator of C_L under the Pareto distribution based on the right Type II censored sample. Lee et al. (2010) constructed the UMVUE of C_L and developed a testing procedure for the performance index of products for the exponential distribution based on the Type II right censored sample. Lee et al. (2015) derived maximum likelihood estimator and constructed a hypothesis testing procedure for lifetime performance index from the normal distribution but sample data modeled by fuzzy numbers.

In life testing experiments, there are time limits and/or other restrictions including money, material resources, mechanical or experimental difficulties on data collection. In these situations, the exact lifetime of products put on the test is not observable and the censored data is used. There are several types of censoring schemes in survival analysis and the Type-II censoring is one of the most common for consideration. So, in this paper, we construct a UMVUE of C_L for Burr Type III distribution with Type II right censored sample. The UMVUE of C_L is then utilized to develop a new hypothesis testing procedure under the condition of a known lower specification limit L . The new testing procedure can be employed to assess whether the lifetime of products (or items) adheres to the required level in the condition of known L .

The rest of this paper is organized as follows. Section 2 introduces some properties of the lifetime performance index for lifetime of product (or item) and discusses the relationship between the lifetime performance index and conforming rate. Section 3 presents the UMVUE of the lifetime performance index and its statistical properties. Section 4 develops a new hypothesis testing procedure for the lifetime performance index. In the same section, the power of the statistical test and the one sided confidence interval of C_L are obtained. A Monte Carlo simulation algorithm of power function and real data example are provided in Sections 5 and 6 respectively. The article ends with concluding remarks.

2. Lifetime Performance Index and Conforming Rate

In 1942, Burr suggested twelve types of cumulative distribution function, one of them, was Burr Type III distribution. It is a very important lifetime model in the analysis of equipment failure data, vehicle and informational studies, as well as in models of stress and durability.

Suppose that the lifetime X of product has a two-parameter Burr Type III distribution with the cumulative distribution function (cdf) and probability distribution function (pdf) defined respectively as

$$F(x; \lambda, \beta) = \left[1 + x^{-\lambda} \right]^{-\beta}, \quad x > 0, \quad \lambda, \beta > 0, \quad (1)$$

where λ and β are shape parameters. The corresponding pdf of Burr Type III distribution is,

$$f(x; \lambda, \beta) = \lambda \beta x^{-\lambda-1} [1 + x^{-\lambda}]^{-\beta-1}; x > 0. \quad (2)$$

To obtain the performance index C_L for the Burr Type III distribution, a random variable X will be converted to exponential distribution through the transformation $Y = \log_e(1 + X^{-\lambda})$. Hence, the pdf and cdf of Y are obtained respectively as follows,

$$f(y; \beta) = \beta e^{-\beta y}, y > 0, \beta > 0, \quad (3)$$

and

$$F(y; \beta) = 1 - e^{-\beta y}, \beta > 0.$$

Also, it is known that the mean and variance of Y are $1/\beta$ and $1/\beta^2$ respectively.

The failure rate function $r(y; \beta)$ is defined by

$$r(y; \beta) = \frac{f(y; \beta)}{1 - F(y; \beta)} = \frac{\beta e^{-\beta y}}{e^{-\beta y}} = \beta, \quad \beta > 0. \quad (4)$$

So, the life time performance index C_L is obtained by substituting mean and variance as follows

$$C_L = \frac{\mu - L}{\sigma} = \frac{1/\beta - L}{1/\beta} = 1 - \beta L, \quad -\infty < C_L < 1, \quad (5)$$

where the process mean $\mu = 1/\beta$, the process standard deviation $\sigma = \sqrt{\text{Var}(Y)} = 1/\beta$, and L is the lower specification limit.

When the mean lifetime of products $(1/\beta) > L$, then the life time performance index $C_L > 0$. From Equations (4) and (5), it can be seen that the larger the mean $1/\beta$, the smaller the failure rate and the larger the lifetime performance index C_L . Therefore, the life time performance index C_L reasonably and accurately represents the lifetime performance of products.

If the lifetime of a product Y ; which $Y \geq L$ and $\beta > 0$ exceeds the lower specification limit L , then the product is defined as a conforming product. The ratio of conforming products is known as the conforming rate P_r and can be defined as

$$P_r = P(Y \geq L) = \int_L^{\infty} f(y) dy = \int_L^{\infty} \beta e^{-\beta y} dy = e^{-\beta L}.$$

From Equation (5), $\beta L = 1 - C_L$, then P_r is as follows

$$P_r = e^{(C_L - 1)}, \quad -\infty < C_L < 1. \quad (6)$$

Obviously, a strictly increasing relationship exists between the conforming rate P_r and the lifetime performance index C_L . Thus, the larger the index value C_L , the larger conforming rate P_r . Table 1 lists various C_L values and the corresponding P_r . For the C_L values which are not listed in Table1, the conforming rate P_r can be obtained through

interpolation. The conforming rate can be calculated by dividing the number of conforming products by the total number of products sampled.

To accurately estimate the conforming rate, Montgomery (1985) suggested that the sample size must be large. However, a large sample size is usually not practical from the perspective of cost, since collecting the lifetime data of new products (or items) need many monies. In addition, a complete sample is also not practical due to time limitation and /or other restrictions (such as lack of funds, lack of material resources, mechanical or experimental difficulties, etc.) on data collection.

Since a one-to one mathematical relationship exists between the conforming rate P_r and the lifetime performance index C_L . Therefore, utilizing the one-to-one relationship between P_r and C_L , lifetime performance index can be a flexible and effective tool, not only evaluating product quality, but also for estimating the conforming rate P_r .

Table1. The lifetime performance index versus the conforming rate

C_L	P_r	C_L	P_r
$-\infty$	0.00000	0.15	0.42741
-9.00	0.00004	0.20	0.44933
-8.00	0.00012	0.25	0.47237
-7.00	0.00033	0.30	0.49659
-6.00	0.00091	0.35	0.52205
-5.00	0.00248	0.40	0.54881
-4.50	0.00409	0.45	0.57695
-4.00	0.00673	0.50	0.60653
-3.50	0.01111	0.55	0.63763
-3.00	0.01832	0.60	0.67032
-2.50	0.03019	0.65	0.70469
-2.00	0.04979	0.70	0.74082
-1.50	0.08208	0.75	0.77880
-1.00	0.13534	0.80	0.81873
-0.50	0.22313	0.85	0.86071
0.00	0.36788	0.90	0.90484
0.05	0.38674	0.95	0.95123
0.10	0.40657	1.00	1.00000

3. UMVUE of Lifetime Performance Index

In this section, Type II censored is considered, let n denote the sample size of the test, m denote the number of failures in the test. Consider that $y_{(1)}, y_{(2)}, \dots, y_{(m)}$ is the corresponding Type II censored sample. The likelihood function for observed samples $y_{(1)}, y_{(2)}, \dots, y_{(m)}$, is obtained as follows

$$\begin{aligned} L &= \frac{n!}{(n-m)!} \prod_{i=1}^m f(y_{(i)}, \beta) [1 - F(y_{(m)})]^{n-m} \\ &= \frac{n!}{(n-m)!} \prod_{i=1}^m \left(\beta e^{-\beta y_{(i)}} \right) \left[1 - \left(1 - e^{-\beta y_{(m)}} \right) \right]^{n-m} \\ &= \frac{n!}{(n-m)!} \beta^m e^{-\beta \left(\sum_{i=1}^m y_{(i)} + (n-m)y_{(m)} \right)} \end{aligned} \quad (7)$$

To show that $S(y) = \sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}$ is a complete and sufficient statistic from Equation (7), it must be rewritten as the exponential family as follows

$$L = \exp \left[\left(\ln \frac{n!}{(n-m)!} \right) + m \ln \beta - \left(\beta \left(\sum_{i=1}^m y_{(i)} + (n-m)y_{(m)} \right) \right) \right], \quad (8)$$

where, $c(\beta) = -\beta$, $S \equiv S(y) = \sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}$, $d(\beta) = m \ln \beta$, $A = \ln \left[\frac{n!}{(n-m)!} \right]$,

$$I_A(y) = \prod_{i=1}^m (1, \infty), y = y_{(1)}, y_{(2)}, \dots, y_{(m)}.$$

Hence, based on Lehmann (1983), and Hogg et al. (2005), $S = \sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}$ is a complete and sufficient statistic for β . In addition, by using Theorem 4.1.1 and Corollary 4.1.1 of Lawless (2003) also can obtain that $2\beta S \sim \chi^2_{(2m)}$.

By using the invariance property of UMVUE, hence the UMVUE estimator of C_L say $\hat{C}_L$, can be written as

$$\begin{aligned} \hat{C}_L &= 1 - \hat{\beta} L, \\ &= 1 - \frac{(m-1)L}{\sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}}. \end{aligned} \quad (9)$$

Meanwhile, the k^{th} moment of $\hat{C}_L$ is obtained as follows

$$E \left(\hat{C}_L \right)^k = E \left[1 - \frac{(m-1)L}{\sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}} \right]^k$$

$$= E \left[1 - \frac{2\beta(m-1)L}{2\beta S} \right]^k,$$

where $2\beta S$ has a chi-square distribution with $2m$ degrees of freedom.

$$E \left(\hat{C}_L \right)^k = \sum_{i=0}^k \binom{k}{i} (-1)^i [2\beta(m-1)L]^i E (2\beta S)^{-i},$$

where,

$$E (2\beta S)^{-i} = \frac{\Gamma(m-i)}{\Gamma(m)2^i}.$$

Hence, the k^{th} moment of $\hat{C}_L$ takes the form

$$E \left(\hat{C}_L \right)^k = \sum_{i=0}^k \binom{k}{i} (-1)^i \frac{[(m-1)L]^i \beta^i \Gamma(m-i)}{\Gamma(m)}. \quad (10)$$

In particular, the expected value and the variance of $\hat{C}_L$ can be obtained respectively as

$$E \left(\hat{C}_L \right) = 1 - \beta L, \quad (11)$$

and

$$Var(\hat{C}_L) = \frac{\beta^2 L^2}{m-2}, m > 2. \quad (12)$$

4. Testing Procedure for the Lifetime Performance Index

In this section, the power of the statistical test and at given specified significance level γ , the level $(1-\gamma)$ one-sided confidence interval for C_L are discussed.

To assess whether the lifetime performance index adheres to the required level, a statistical testing procedure is constructed. Assuming that the required index value of lifetime performance is larger than c^* , where c^* denotes the target value, the null hypothesis

$$H_0: C_L \leq c^* \text{ (the product is unreliable),}$$

the alternative hypothesis

$$H_1: C_L > c^* \text{ (the product is reliable),}$$

are constructed. By using $\hat{C}_L$, the UMVUE of C_L as the test statistic, the rejection region can be expressed as $\{\hat{C}_L \mid C_L > C_0^*\}$. Given the specified significance level γ , the critical value C_0^* can be calculated as follows

$$\begin{aligned} P(\hat{C}_L > C_0^* \mid C_L \leq c^*) &= \gamma, \\ P\left(1 - \frac{(m-1)L}{S} > C_0^* \mid C_L \leq c^*\right) &= \gamma, \\ P\left(1 - \frac{2\beta(m-1)L}{2\beta S} > C_0^* \mid C_L \leq c^*\right) &= \gamma, \\ \left(\text{Since } P\left(2\beta S > \frac{2(m-1)(1-c^*)}{1-C_0^*}\right) \text{ as } C_L \leq c^* \right), \\ P\left(2\beta S \leq \frac{2(m-1)(1-c^*)}{1-C_0^*}\right) &= 1 - \gamma, \end{aligned} \quad (13)$$

where $2\beta S \sim \chi^2_{(2m)}$. From Equation (13), utilizing inverse-chi-square (INVCHI) $(1-\gamma, 2m)$ function which represents the lower $1-\gamma$ percentile of $\chi^2_{(2m)}$, then

$$\frac{2(m-1)(1-c^*)}{1-C_0^*} = \text{INVCHI}(1-\gamma, 2m),$$

is obtained. Thus, the critical value can be derived as;

$$C_0^* = 1 - \frac{2(m-1)(1-c^*)}{\text{INVCHI}(1-\gamma, 2m)}, \quad (14)$$

where c^* , γ and m are denoted to the target value, the specified significance level and the number of observed failures before termination, respectively.

4.1 Power Function of the Test

The power of a statistical test is the probability of correctly rejecting a false null hypothesis. The procedure used in practice is to limit the probability of Type I error to some pre-assigned level γ (usually 0.01 or 0.05) that is small and to maximize the power of a statistical test. The null hypothesis $H_0: C_L \leq c^*$, and the alternative

hypothesis $H_1: C_L > c^*$, are constructed. The power of a statistical test is derived as follows:

Under Type II censoring scheme, we get a size γ test with the rejection region

$$\left\{ \hat{C}_L \left| \hat{C}_L > 1 - \frac{2(m-1)(1-c^*)}{\text{INVCHI}(1-\gamma, 2m)} \right. \right\}, \text{ for the number of observed failures before}$$

termination m and sample size n ($m \leq n$). The power $P(C_L)$ of the test at this point $C_L > c^*$ defined as;

$$\begin{aligned} P(C_L) &= P \left(\hat{C}_L > 1 - \frac{2(m-1)(1-c^*)}{\text{INVCHI}(1-\gamma, 2m)} \right) \\ &= P \left(1 - \frac{(m-1)L}{S} > 1 - \frac{2(m-1)(1-c^*)}{\text{INVCHI}(1-\gamma, 2m)} \left| \beta = \frac{1-C_L}{L} \right. \right) \\ &= P \left(2\beta S > \frac{\beta L \text{INVCHI}(1-\gamma, 2m)}{(1-c^*)} \right) \\ &= P \left(2\beta S > \frac{(1-C_L) \text{INVCHI}(1-\gamma, 2m)}{(1-c^*)} \right), \end{aligned} \quad (15)$$

where $2\beta S \sim \chi^2_{(2m)}$.

4.2 Confidence Interval for C_L

In this section, given the specified significance level γ , the level $(1-\gamma)$ one-sided confidence interval for C_L can be derived as follows:

Since the pivotal quantity $2\beta S$, where $2\beta S \sim \chi^2_{2m}$ and $\text{INVCHI}(1-\gamma, 2m)$ which represents the lower $1-\gamma$ percentile of χ^2_{2m} . So

$$P(2\beta S \leq \text{INVCHI}(1-\gamma, 2m)) = 1-\gamma, \quad (16)$$

where $C_L = 1 - \beta L$ and $\hat{C}_L = 1 - \frac{(m-1)L}{S}$. Multiply inequality (16) by $\frac{\hat{\beta} L}{2(m-1)}$, then

$$P \left(\frac{2\beta S \hat{\beta} L}{2(m-1)} \leq \frac{\hat{\beta} L \text{INVCHI}(1-\gamma, 2m)}{2(m-1)} \right) = 1-\gamma. \quad (17)$$

But, $\hat{\beta} = \frac{(m-1)}{S}$,

$$P \left(\frac{\beta \hat{\beta} L}{\hat{\beta}} \leq \frac{\hat{\beta} L \text{ INVCCHI}(1-\gamma, 2m)}{2(m-1)} \right) = 1-\gamma,$$

$$P \left(1 - \beta L \geq 1 - \frac{\hat{\beta} L \text{ INVCCHI}(1-\gamma, 2m)}{2(m-1)} \right) = 1-\gamma, \quad (18)$$

From Equation (18), then

$$C_L \geq 1 - \frac{(1 - \hat{C}_L) \text{ INVCCHI}(1-\gamma, 2m)}{2(m-1)} \quad (19)$$

is the level $(1-\gamma)$ one-sided confidence interval for C_L . Thus, the level $(1-\gamma)$ lower confidence bound (LB) for C_L can be written as

$$LB = 1 - \frac{(1 - \hat{C}_L) \text{ INVCCHI}(1-\gamma, 2m)}{2(m-1)}, \quad (20)$$

where $\hat{C}_L$, γ and m denote the UMVUE of C_L , the specified significance level and the number of observed failures before termination, respectively.

In addition, the proposed testing procedure can be constructed with the one-sided confidence interval too. The decision rule of statistical test is “If performance index value $c^* \notin (LB, \infty)$, it is concluded that the lifetime performance index of products meets the required level”.

5. Numerical Study

The numerical study is designed to obtain the critical values of the power test. The simulated procedures are described as follows

Step 1:

(a) Generate a random sample $u_1, u_2, \dots, u_m$ of size n from a uniform (0,1), then the uniform random numbers can be transformed to Burr III random numbers by using the following transformation

$$x_i = \left(u_i^{\frac{-1}{\beta}} - 1 \right)^{\frac{-1}{\lambda}}, \quad i = 1, 2, \dots, m.$$

The generated data for $y_i = \log_e(1 + x_i^{-\lambda})$, $i = 1, \dots, m$, $(x_1, x_2, \dots, x_m)$ is a random sample from the Burr Type III distribution

(b) Set $(y_{(1)}, y_{(2)}, \dots, y_{(m)})$ as a Type II right censored sample from a one-parameter exponential distribution with pdf (3).

(c) Given $c^* = 0.1, 0.5, 0.9$, $C_L = 0.1, 0.5, 0.9$, $\lambda = 1.7$, $L=0.5$, and the indicated significance levels of $\gamma = 0.01$ and 0.05 . The selected values of (n, m) are $(n, m)=(10,2)$, $(20,4)$, $(30,6)$, $(50,10)$, where $c^* < C_L < 1$ and $m < n$.

(d) The value of $\hat{C}_L$ is calculated by
$$\hat{C}_L = 1 - \frac{(m-1)L}{\sum_{i=1}^m y_{(i)} + (n-m)y_{(m)}}$$
.

(e) If $\hat{C}_L > c_0^*$, the critical value is obtained by using $C_0^* = 1 - \frac{2(m-1)(1-c^*)}{\text{INVCCHI}(1-\gamma, 2m)}$.

Step 2: (a) Step1 is repeated 1000 times.

(b) The estimation of the power $p(C_L)$, is $\hat{P}(C_L) = \frac{\text{Total Count}}{1000}$.

Step 3: (a) Step 2, based on 100 estimations of the power $P(C_L)$ can be obtained as follows: $\hat{P}_1(C_L), \hat{P}_2(C_L), \dots, \hat{P}_{100}(C_L)$.

(b) The mean $\overline{\hat{P}(C_L)}$ of $\hat{P}_1(C_L), \hat{P}_2(C_L), \dots, \hat{P}_{100}(C_L)$, that is $\overline{\hat{P}(C_L)} = \frac{\sum_{i=1}^{100} \hat{P}_i(C_L)}{100}$, is calculated.

(c) The MSE of $\hat{P}_1(C_L), \hat{P}_2(C_L), \dots, \hat{P}_{100}(C_L)$ is obtained as follows

$$\text{MSE} = \frac{\sum_{i=1}^{100} (\hat{P}_i(C_L) - P(C_L))^2}{100}.$$

Step 4: The results of simulated data are listed in Tables 2 and 3.

Table 2 The values of $P(C_L)$, $\overline{\hat{P}(C_L)}$ and MSE for the test at and $\gamma = 0.01$

(n,m)	C_L	$P(C_L)$	$\overline{\hat{P}(C_L)}$	MSE
(10,2)	0.1	0.102	0.11464	0.00008
	0.2	0.158	0.16223	0.00015
	0.3	0.225	0.22444	0.00011
	0.4	0.311	0.30162	0.00018
	0.5	0.419	0.39794	0.00018
	0.6	0.513	0.50769	0.00031
	0.7	0.64	0.62132	0.00019
	0.8	0.733	0.73375	0.00019
	0.9	0.815	0.83224	0.00011
(20,4)	0.1	0.036	0.03791	0.00003
	0.2	0.066	0.06969	0.00007
	0.3	0.11	0.12378	0.00007
	0.4	0.206	0.20793	0.00017
	0.5	0.32	0.33202	0.00023
	0.6	0.505	0.50244	0.00022
	0.7	0.686	0.68711	0.00018
	0.8	0.859	0.84906	0.00011
	0.9	0.955	0.95012	0.00004
(30,6)	0.1	0.009	0.0083	0.00000
	0.2	0.03	0.01897	0.00001
	0.3	0.058	0.04483	0.00003
	0.4	0.111	0.10084	0.00008
	0.5	0.226	0.20733	0.00015
	0.6	0.409	0.39547	0.00022
	0.7	0.654	0.64797	0.00022
	0.8	0.843	0.86947	0.00011
	0.9	0.983	0.97845	0.00002
(50,10)	0.1	0.002	0.00028	0.00000
	0.2	0.002	0.00102	0.00000
	0.3	0.003	0.00368	0.00000
	0.4	0.01	0.01401	0.00001
	0.5	0.05	0.05054	0.00004
	0.6	0.157	0.17075	0.00014

(n,m)	C_L	$P(C_L)$	$\hat{P}(C_L)$	MSE
	0.7	0.452	0.45792	0.00022
	0.8	0.853	0.83797	0.00013
	0.9	0.989	0.99154	0.00000

Table 3 The values of $P(C_L)$, $\hat{P}(C_L)$ and MSE for the test at $\gamma = 0.05$

(n, m)	C_L	$P(C_L)$	$\hat{P}(C_L)$	MSE
(10,2)	0.1	0.287	0.29048	0.00015
	0.2	0.348	0.35118	0.00028
	0.3	0.417	0.41815	0.00022
	0.4	0.453	0.49504	0.00025
	0.5	0.583	0.57597	0.00021
	0.6	0.636	0.65109	0.00019
	0.7	0.749	0.7288	0.00018
	0.8	0.801	0.79704	0.00018
	0.9	0.858	0.85668	0.00009
(20,4)	0.1	0.132	0.13298	0.00012
	0.2	0.176	0.20031	0.00013
	0.3	0.273	0.29113	0.00018
	0.4	0.395	0.40436	0.00024
	0.5	0.504	0.53828	0.00028
	0.6	0.676	0.67845	0.00018
	0.7	0.801	0.80758	0.00018
	0.8	0.887	0.90554	0.00009
	0.9	0.964	0.96457	0.00003
(30,6)	0.1	0.031	0.03931	0.00003
	0.2	0.08	0.07566	0.00006
	0.3	0.132	0.13826	0.00012
	0.4	0.25	0.24643	0.00018
	0.5	0.425	0.40484	0.00022
	0.6	0.605	0.60371	0.00022
	0.7	0.774	0.79464	0.00018
	0.8	0.927	0.92893	0.00005
	0.9	0.986	0.98654	0.00001
(50,10)	0.1	0.002	0.00199	0.00002
	0.2	0.004	0.00591	0.00000
	0.3	0.016	0.0181	0.00001

(n, m)	C_L	$P(C_L)$	$\overline{\hat{P}(C_L)}$	MSE
	0.4	0.064	0.05213	0.00004
	0.5	0.131	0.14097	0.00010
	0.6	0.342	0.34203	0.00022
	0.7	0.635	0.65975	0.00027
	0.8	0.906	0.92025	0.00006
	0.9	0.998	0.9961	0.00000

From Tables 2 and 3, we observe the following:

- For fixed value of m , the simulation power $\overline{\hat{P}(C_L)}$ and the power $P(C_L)$ increase as C_L increases (see Figure 1).

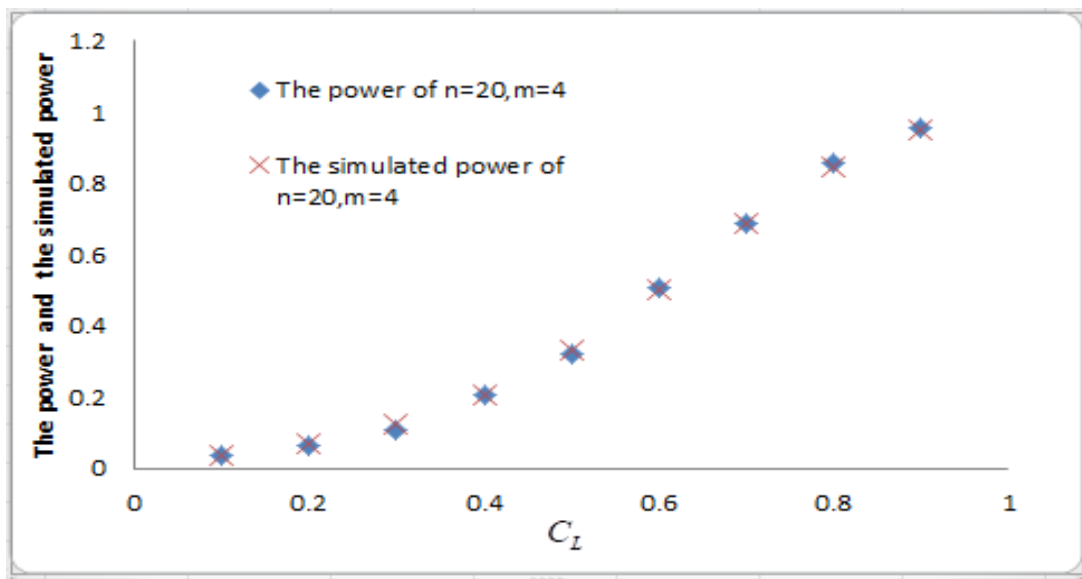


Figure 1 The power function and simulated power function at $\gamma = 0.01$ for $n=20$, $m=4$ and

- All of the simulated power function $\overline{\hat{P}(C_L)}$ close to the power $P(C_L)$, for any value of C_L .
- The power values, $P(C_L)$ increases when C_L increases for different values of n and m (see Figure 2).

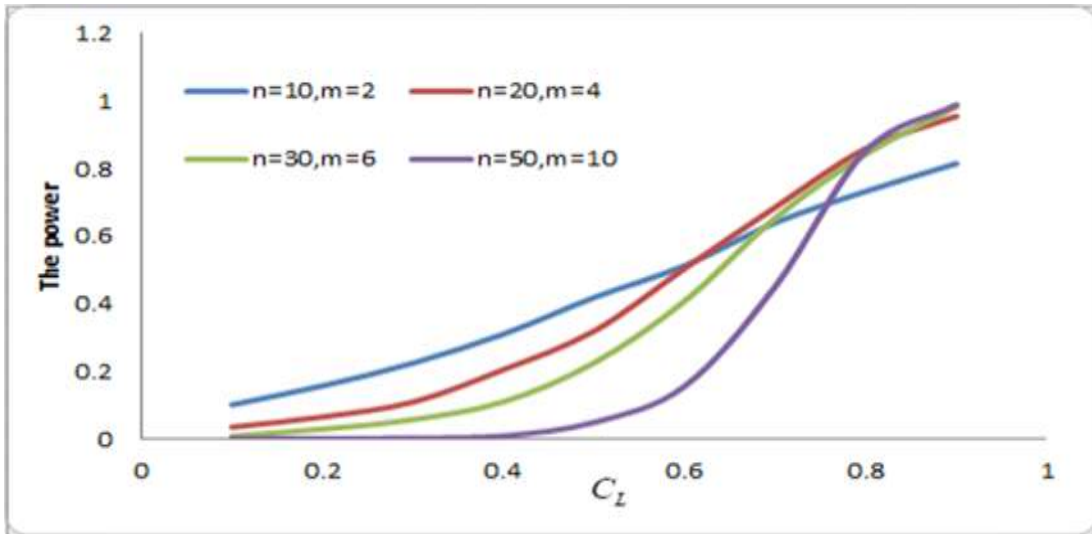


Figure 2. The power function of test at $\gamma = 0.01$ for different values of n and m

- Values of MSE are enough small and the scope of MSE is between 0.00000 and 0.00029.
- These results illustrate that the performance of our proposed method is acceptable.

6. Real Data Example

A real data set given by Mann and Fertig (1973) by (Real life data) is considered here. This data set is also studied by Altindag et al. (2017) who proposed the estimation and prediction problems for the Burr Type III distribution under Type II censored data.

The data set contains 10 failure times of airplane components of total 13 items. The censoring scheme is Type II censoring with $m = 10$ and $n = 13$. The observed failure times and the Type II censoring scheme are given in Table 4.

Then by using the transformed Type II censored sample with transformation $y = \log_e(1 + x^{-\lambda})$ for $\hat{\lambda} = 1.4990$, the transformed values are calculated.

Table 4. The observed failure times and censoring scheme

i	1	2	3	4	5
x_i	0.22	0.50	0.88	1.00	1.32
i	6	7	8	9	10
x_i	1.33	1.54	1.76	2.50	3.00

Step 1: The lower lifetime limit $L=0.1$. To deal with concerns about the lifetime performance, the conforming rate P_r of products is required to exceed 80%. Referring to Table 1, the C_L value is required to exceed 0.80. Thus, the performance index value is set at $c^* = 0.80$. The testing hypothesis $H_0: C_L \leq c^*$ versus $H_1: C_L > c^*$ is constructed.

Step 2: Specify a significance level $\gamma = 0.05$.

Step 3: Calculate the value of test statistic

$$\hat{C}_L = 1 - \frac{(10 - 1) \times 0.1}{7.39 + (13 - 10) 2.37} = 0.94.$$

Step 4: Obtain the critical value $C_0^* = 0.885$ by using Equation (14), according to $c^* = 0.80$, $m=10$ and the significance level $\gamma = 0.05$.

Step5. Because of $\hat{C}_L = 0.94 > C_0^* = 0.885$, so we reject the null hypothesis

$$H_0: C_L \leq c^*$$

The lower confidence bound $(0.90, \infty)$ for Type II censored. So, the performance index value to $c^* = 0.80$ don't belong to lower confidence bound, it is also concluded that the lifetime performance index of airplane components meets the required level.

7. Concluding Remarks

This study constructs a UMVUE estimator of C_L under the assumption of Burr Type III distribution in case of Type II censored sample. The UMVUE estimator of C_L is then utilized to develop the new hypothesis testing procedure in the condition of known L . The proposed testing procedure is easily applied and can effectively evaluate whether the lifetime of products meets requirements. In addition, this study provides a table of the lifetime performance index with its corresponding conforming rate. Hence, for any specified conforming rate, a corresponding C_L can be obtained, and the hypotheses of the proposed testing procedure can also be expressed in terms of the conforming rate under L is known limit.

References

Altindag, O.; Ankaya, M. N.; Yalinkaya, A. and Aydogdu, H. (2017). Statistical inference for the Burr Type III distribution under Type II censored data. *Communication Science*. 66(2), 297-310.

Burr, W. I. (1942). Cumulative frequency functions. *The Annals of Mathematical Statistics*, 13(2), 215-232.

Chen, H. T.; Tong, L. I. and Chen, K. S. (2002). Assessing the lifetime performance of electronic components by confidence interval. *Journal of the Chinese Institute of Industrial Engineers*, 19(2), 53–60.

Hogg, R. V., McKean, J. W. and Craig, A. T. (2005). *Introduction to Mathematical Statistics*, 6th edn. Pearson Prentice Hall, Pearson Education, Ltd., London.

Hong, C. W., Wu .J. W. and Cheng C. H. (2007). Computational procedure of performance assessment of lifetime index of businesses for the Pare to lifetime model with the right type II censored sample. *Journal of Computational and Applied Mathematics*, 184, 336–350.

Kane, V.E. (1986). Process capability indices. *Journal of Quality Technology*, 18, 41-52.

Lawless, J. F.(1982). *Statistical Models and Methods for Life Time Data*. John Wiley, New York.

Lee, W. C. , Hong, C. W. and Wu, J. W. (2015). Computational procedure of performance assessment of lifetime index of normal products with fuzzy data under the Type II right censored sampling plan. *Journal of Intelligent & Fuzzy Systems*, 28(4), 1755-1773.

Lee, W.C., Wu, J.W. and Lei, L.C. (2010). Evaluating the lifetime performance index for the exponential lifetime products. *Applied Mathematical Modeling* , 34, 1217–1224.

Lehmann, E. L. (1983). *Theory of Point Estimation*. Wiley, New York.

Mann, N. R. and Fertig, K. W. (1973). Tables for obtaining Weibull condence bounds and tolerance bounds based on best linear invariant estimates of parameters of the extreme-value distribution. *Technometrics* 15(1), 87-101.

Montgomery, D.C. (1985). *Introduction to Statistical Quality Control*. John Wiley and Sons Inc., New York.

Pan, J.N. and Wu, S.L. (1997). Process capability analysis for non-normal relay test data, *Microelectronics and Reliability*, 37, 421-428.

Tong, L. I., Chen, K. S. and Chen, H. T. (2002). Statistical testing for assessing the performance of lifetime index of electronic components with exponential distribution. *International Journal of Quality & Reliability Management*, 19, 812-824.

Modified Topp-Leone-Chen Distribution: Properties and Estimation Based on Progressive Type-II Censoring Scheme

AL-Sayed, N. T., Swielum, E. M., AL-Dayian, G. R. and EL-Helbawy, A. A.

Statistics Department, Faculty of Commerce

AL-Azhar University (Girls' Branch), Cairo, Egypt

Abstract

In recent years, composite models based on Topp-Leone distribution have become popular in actuarial sciences and related areas. On the other hand the modeling and analysis of lifetimes are important aspects of statistical work in a wide variety of scientific and technological fields. In this paper, a modified Topp-Leone Chen distribution is introduced and studied as a composite distribution. Some properties of the proposed distribution are obtained. The maximum likelihood method is used under progressive Type-II censored samples for estimating the model parameters, reliability and hazard rate functions. A numerical example is given to investigate the precision of the maximum likelihood estimates and an application using real data sets are used to demonstrate how the results can be used in practice.

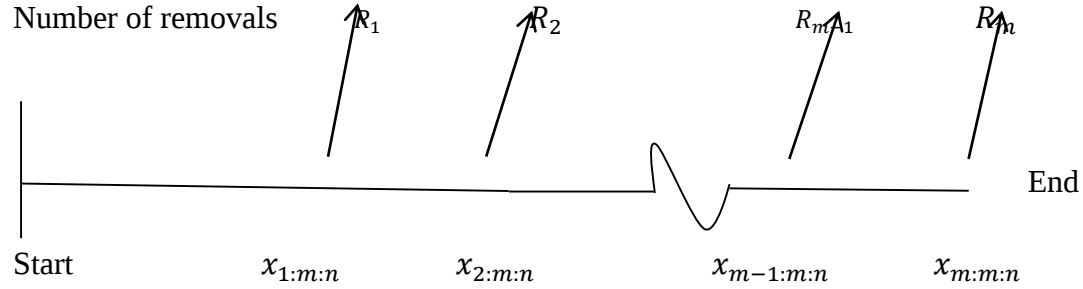
Keywords: *Topp-Leone distribution; Chen distribution; progressive Type-II censored samples; maximum likelihood method; asymptotic Fisher information matrix.*

1. Introduction

In context of reliability analysis such data are often known to be censored. Several methods of censoring have been suggested in literature to analyze lifetime data. Type-I and Type-II censoring schemes have probably found the most extensive applications in these situations. In Type-I censoring a life test is conducted for a fixed-time period while in Type-II censoring an experiment terminates when a prescribed number of units fail. One drawback of these schemes is that live units can be removed only at the end of the experiment. However in many life testing experiment, it is desired to withdraw live units from the experiment at time points other than the final termination point of the test. The progressive censoring possesses such flexibility and thus allows in between removals of units as well. Different progressive censoring schemes have been introduced in the literature. The most popular one is known as the progressive Type-II censoring scheme and it can be briefly described as follows:

Considering n identical units are put to test and the lifetime distribution of the n units are denoted by $X_1, X_2, \dots, X_n$. The integer $m (< n)$ is fixed at the beginning of the experiment and $R_1, R_2, \dots, R_m$ are m pre-fixed integers satisfying $R_1 + R_2 + \dots + R_m + m = n$. At the time of the first failure $X_{1:m:n}$, R_1 units are chosen randomly from the remaining $n - 1$ units and they are removed from the experiment. Similarly at the time of the second failure $X_{2:m:n}$, R_2 of the remaining $n - R_1 - 2$ units are removed from the test and so on. Finally, when the m -th failure is observed the experiment is

terminated and the remaining surviving units R_m with $R_m = n - R_1 - R_2 - \dots - R_{m-1} - m$ are removed. Here $(R_1, R_2, \dots, R_m)$ is known as the censoring scheme and it is prefixed before the experiment starts. [For more details about the progressive Type-II censoring scheme, see, Balakrishnan and Aggarwala (2000), Balakrishnan and Cramer (2014), Almetwally *et al.* (2018), Almetwally and Almongy (2018) and Karakoca and Pekgör (2019)].



Topp-Leone (TL) distribution was proposed by Topp and Leone (1955); as an alternate model failure data. It is a continuous unimodal distribution with bounded support; therefore it is appropriate for modeling lifetime of distributions with finite support such as limited power supply, maintenance/repair resource, or design life of the system. Nadarajah and Kotz (2003) showed that TL distribution have bathtub failure rate function with widespread applications in reliability. Furthermore, there were numerous authors who were interested in this distribution. For example, EL-Helbawy (2018 a), EL-Helbawy (2018 b), Arshad and Jamal (2019), Reyad *et al.* (2019) and Oguntunde *et al.* (2019).

The *probability density function* (pdf) of TL distribution is given by

$$f(y; \alpha, \theta) = \frac{2\alpha}{\theta} \left(1 - \frac{y}{\theta}\right) \left[2\frac{y}{\theta} - \left(\frac{y}{\theta}\right)^2\right]^{\alpha-1}, \quad 0 < y < \theta; (\alpha, \theta > 0). \quad (1)$$

The *cumulative distribution function* (cdf) of TL distribution is as follows:

$$F(y; \alpha, \theta) = \begin{cases} 0, & y \leq 0, \\ \left[2\frac{y}{\theta} - \left(\frac{y}{\theta}\right)^2\right]^\alpha, & 0 < y < \theta, \\ 1, & y \geq \theta, \end{cases} \quad (2)$$

where α is a shape parameter, θ is a scale parameter and the pdf in (1) is J-shaped distribution. For $\theta = 1$, then the standard TL distribution is obtained.

Chen (2000) introduced a two-parameter lifetime distribution with bathtub shaped or increasing failure rate function which enables it to fit real lifetime data sets.

The pdf and cdf of Chen distribution are, respectively, given by

$$g(t; \lambda, b) = \lambda b t^{b-1} \exp \left[\lambda (1 - e^{t^b}) + t^b \right], \quad t > 0; (\lambda, b > 0), \quad (3)$$

and

$$G(t; \lambda, b) = 1 - \exp \left[\lambda (1 - e^{t^b}) \right], \quad t > 0; (\lambda, b > 0), \quad (4)$$

where λ and b are shape parameters. The parameter b governs the shape of hazard function for this distribution. Many practical investigations such as observing strength of specific materials, mortality rate and latency period of a deadly disease often lead to the study of distributions that have bathtub-shaped hazard functions. Numerous authors who were interested in this distribution, for example, Chaubey and Zhang (2015), Kayal *et al.* (2016) and Dey *et al.* (2017).

1.1 Topp-Leone Chen distribution

Considering $\theta = 1$ in (2); without any loss of generality, a random variable Y is distributed as the TL distribution with parameter α denoted by $Y \sim TL(\alpha)$ with a cdf

$$H(y) \equiv H(y; \alpha) = [2y - y^2]^\alpha, \quad 0 < y < 1; (\alpha > 0). \quad (5)$$

A composition of H , given by (5) and a cdf G , with positive support, yields a new cdf, which is given below

$$F(t) = H(G(t)) = [2G(t) - (G(t))^2]^\alpha, \quad (6)$$

where $G(t)$ is a cdf of Chen distribution; denoted by $T \sim Ch(\lambda)$, and is as follows:

$$G(t) \equiv G(t; \lambda) = 1 - \exp[\lambda(1 - e^t)], \quad t > 0; (\lambda > 0). \quad (7)$$

[For more details on composition of distribution functions, see, AL-Hussaini (2012)].

Substituting (7) in (6), the cdf for *Topp-Leone Chen* (TLCh) distribution is given by

$$F(t) \equiv F(t; \lambda, \alpha) = [1 - \exp[2\lambda(1 - e^t)]]^\alpha, \quad t > 0; (\lambda, \alpha > 0). \quad (8)$$

The pdf, corresponding to the cdf given in (8), is as follows:

$$f(t; \lambda, \alpha) = 2\lambda\alpha e^{[t+2\lambda(1-e^t)]} [1 - \exp(2\lambda(1 - e^t))]^{\alpha-1}, \quad t > 0; (\lambda, \alpha > 0). \quad (9)$$

The rest of this paper is organized as follows: In Section 2, the *Modified Topp-Leone Chen* (MTLCh) distribution is introduced as a composite distribution and some statistical properties are studied, such as, quantile function, *mean residual life* (MRL), *mean past lifetime* (MPL) and order statistics. Some relations between the MTLCh and other distributions are presented in Section 3. In Section 4, *Maximum likelihood* (ML) estimation of the parameters, *reliability function* (rf), *hazard rate function* (hrf) and *reversed hazard rate function* (rhrf) based on progressive Type-II censoring scheme are derived, *asymptotic confidence intervals* (ACIs) for the parameters, rf, hrf and rhrf are

discussed and the asymptotic Fisher information matrix is determined. A numerical study simulation and two real data sets are performed to investigate the precision of ML estimates in Section 5. Finally general conclusion is given in Section 6.

2. Modified Topp-Leone Chen Distribution: Description and Properties

In this section, the MTLCh distribution with three parameters is obtained and some statistical properties are derived.

2.1 Modified Topp-Leone Chen distribution

Adding one or more parameters to a distribution makes the resulting distribution richer and more flexible for analyzing and modeling data.

Let $x = \frac{t}{\beta}$, then the pdf and cdf of MTLCh distribution are, respectively, given by

$$f(x; \lambda, \alpha, \beta) = 2\lambda\alpha\beta e^{[\beta x + 2\lambda(1 - e^{\beta x})]} \left[1 - \exp\left(2\lambda(1 - e^{\beta x})\right)\right]^{\alpha-1},$$

$$x > 0; (\lambda, \alpha, \beta > 0), \quad (10)$$

and

$$F(x; \lambda, \alpha, \beta) = \left[1 - \exp\left(2\lambda(1 - e^{\beta x})\right)\right]^{\alpha}, \quad x > 0; (\lambda, \alpha, \beta > 0), \quad (11)$$

where λ, α are shape parameters and β is a scale parameter.

Many well-known families of distributions can be written as infinite (or finite) weighted series of their parent distributions [see, Eugene *et al.* (2002) and Cordeiro *et al.* (2014)]. This also means that the properties and inferences can be obtained from the same measures of its parent distribution. Therefore, it is useful to derive expansion of the cdf and pdf.

For real value of α , using following series representation of Prudnikov *et al.* (1986)

$$(1 + x)^{\alpha} = \sum_{j=0}^{\infty} \binom{\alpha}{j} x^j = \sum_{j=0}^{\infty} \frac{\Gamma(\alpha+1)}{j! \Gamma(\alpha-j+1)} x^j, \quad (12)$$

then, the cdf in (11) is expressed as infinite sum given as follows:

$$F(x; \lambda, \alpha, \beta) = \sum_{j=0}^{\infty} (-1)^j \binom{\alpha}{j} \exp\left(2\lambda j(1 - e^{\beta x})\right). \quad (13)$$

By using the power series for the exponential function, hence

$$\exp(-2\lambda j e^{\beta x}) = \sum_{k=0}^{\infty} \frac{(-1)^k (2\lambda j)^k e^{\beta k x}}{k!}, \quad (14)$$

substituting (14) in (13) yields

$$F(x; \lambda, \alpha, \beta) = \sum_{(j,k)=0}^{\infty} \varphi_{j,k}(\lambda, \alpha) e^{\beta k x}, \quad (15)$$

$$\text{where } \varphi_{j,k}(\lambda, \alpha) = \frac{(-1)^{j+k} \binom{\alpha}{j} (2\lambda j)^k}{k!} e^{2\lambda j}. \quad (16)$$

The pdf of MTLCh distribution can be obtained by differentiating the cdf given in (15) as follows:

$$f(x; \lambda, \alpha, \beta) = \sum_{(j,k)=0}^{\infty} \varphi_{j,k}(\lambda, \alpha) \beta k e^{\beta k x}, \quad (17)$$

where $\varphi_{j,k}(\lambda, \alpha)$ is defined in (16).

2.2 Some statistical properties

This subsection is devoted to illustrate some statistical properties of the MTLCh distribution.

1. Reliability function and some stress-strength models

The rf of the MTLCh distribution, based on (11), is given by

$$R(x; \lambda, \alpha, \beta) = 1 - \left[1 - \exp \left(2\lambda(1 - e^{\beta x}) \right) \right]^{\alpha}, \quad x > 0; (\lambda, \alpha, \beta > 0). \quad (18)$$

The stress-strength model is a measure of the reliability of a component. In this study, two formulae of stress-strength reliability model are given below:

I. First formula

Suppose that X is the strength of a component which is subjected to stress Y , where X and Y are independent and have MTLCh $(\lambda, \alpha_1, \beta)$ and MTLCh $(\lambda, \alpha_2, \beta)$, then the rf is given by

$$R_1(x; \alpha_1, \alpha_2) = P(Y < X) = \frac{\alpha_1}{\alpha_1 + \alpha_2}. \quad (19)$$

II. Second formula

Let Y and Z are two independent random stress variables with known cdfs $F_Y(y)$ and $F_Z(z)$ which have MTLCh $(\lambda, \alpha_2, \beta)$ and MTLCh $(\lambda, \alpha_3, \beta)$, respectively, considering X is independent of Y and Z is a random strength variable with known cdf, $F_X(x)$, which has MTLCh $(\lambda, \alpha_1, \beta)$, then the rf is given by

$$R_2(x; \alpha_1, \alpha_2, \alpha_3) = P(Y < X < Z) = \frac{\alpha_1 \alpha_3}{(\alpha_1 + \alpha_2)(\alpha_1 + \alpha_2 + \alpha_3)}. \quad (20)$$

2. Hazard and reversed hazard rate functions

The hrf of the MTLCh distribution is given by

$$h(x; \lambda, \alpha, \beta) = \frac{2\lambda\alpha\beta e^{\left[\beta x + 2\lambda(1-e^{\beta x})\right]} \left[1 - \exp\left(2\lambda(1-e^{\beta x})\right)\right]^{\alpha-1}}{1 - \left[1 - \exp\left(2\lambda(1-e^{\beta x})\right)\right]^{\alpha}}, \quad x > 0; (\lambda, \alpha, \beta > 0). \quad (21)$$

The rhrf, which is known by the dual of the hrf; describes the probability of an immediate past failure, given that the unit has already failed at time x , as opposed to the immediate future failure. The rhrf is given by

$$rh(x; \lambda, \alpha, \beta) = \frac{2\lambda\alpha\beta e^{\left[\beta x + 2\lambda(1-e^{\beta x})\right]}}{\left[1 - \exp\left(2\lambda(1-e^{\beta x})\right)\right]}, \quad x > 0; (\lambda, \alpha, \beta > 0). \quad (22)$$

3. Graphical description

The plots of the pdf, hrf and rhrf are provided for different values of parameters in Figures 1-5.

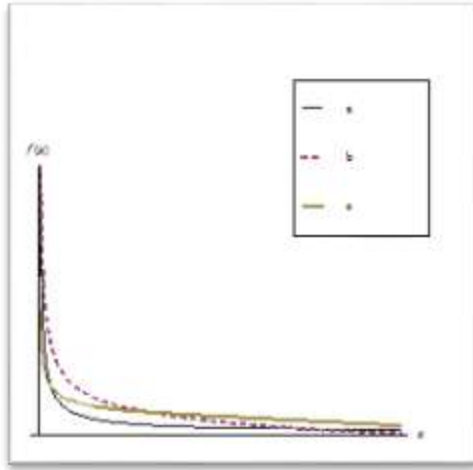


Figure 1

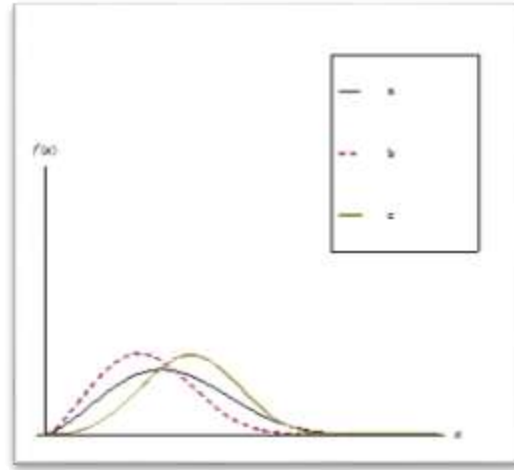


Figure 2

Plots of the probability density function of MTLCh at different values of the parameters

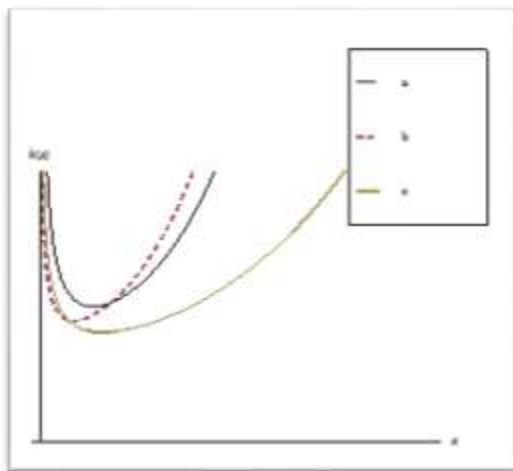


Figure 3

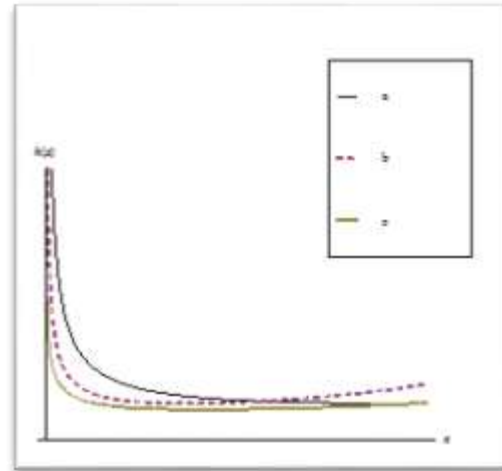


Figure 4

Plots of the hazard rate function of MTLCh at different values of the parameters

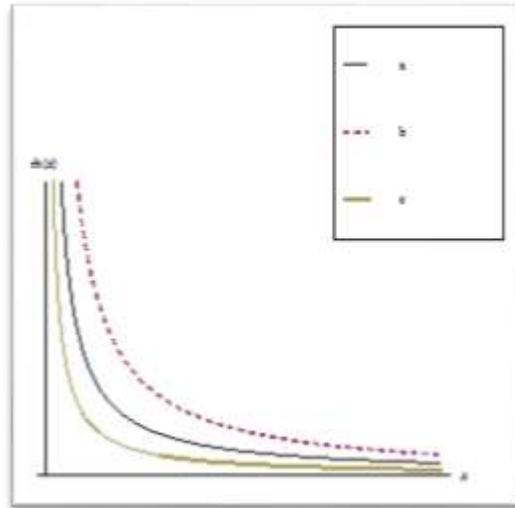


Figure 5

**Plots of the reversed hazard rate function of MTLCh
at different values of the parameters**

In Figure 1, the MTLCh density is monotonically decreasing at $\lambda = 0.21, 0.44, 0.63$, $\alpha = 0.07, 0.61, 0.31$, $\beta = 0.08, 0.13, 0.12$, while in Figure 2 the curves are approximately symmetric at $\lambda = 2.37, 2.39, 3.12$, $\alpha = 0.47, 0.48, 0.23$, and $\beta = 0.26, 0.32, 0.37$.

Figure 3 and 4, show that the hrf of MTLCh distribution for different values of the parameters is quite flexible for modeling survival data. In Figure 3, the hrf has bathtub shape at $\lambda = 0.23, 0.52, 0.46$, $\alpha = 0.24, 0.27, 0.60$, $\beta = 0.35, 0.37, 0.18$, while in Figure 4 the curves are monotonically decreasing at $\lambda = 0.10, 0.30, 0.52$, $\alpha = 0.87, 0.14, 0.39$ and $\beta = 0.04, 0.13, 0.07$.

In Figure 5, the rhrf of MTLCh distribution is monotonically decreasing at $\lambda = 0.04, 0.08, 0.05$, $\alpha = 0.32, 0.01, 0.17$ and $\beta = 0.09, 1.71, 0.41$.

4. Quantile function

It can be obtained by inverting (11) as follows $q = F^{-1}(u)$. Then,

$$q = \frac{1}{\beta} \ln \left[1 - \frac{1}{2\lambda} \ln \left(1 - u^{1/\alpha} \right) \right], \quad 0 < u < 1. \quad (23)$$

One can easily generate x by taking u as a uniform random variable in $(0, 1)$.

Special quantiles may be obtained using (23). For example, if $u = 0.5$, the median of the MTLCh (λ, α, β) is

$$\text{Median} = F^{-1}(1/2) = \frac{1}{\beta} \ln \left[1 - \frac{1}{2\lambda} \ln \left(1 - (0.5)^{1/\alpha} \right) \right], \quad (\lambda, \alpha, \beta > 0). \quad (24)$$

5. Mean residual life

The MRL is the expected remaining life, $X - x_0$, given that the item has survived to time x_0 . It is denoted by $m(x_0)$ and is defined as

$$m(x_0) = E(X - x_0 | X \geq x_0) = \int_{x_0}^{\infty} (X - x_0) P(X \leq t | X > x_0) dx, \quad (25)$$

see, Nassar *et al.* (2018). It can be easily seen that the first moment is equivalent to the MRL function at $x_0 = 0$.

$$m(0) = \frac{\int_0^{\infty} t f(t) dt}{R(0)} = \mu. \quad (26)$$

Alternatively, the MRL can be written as,

$$m(x_0) = \frac{\int_{x_0}^{\infty} R(t) dt}{R(x_0)} = \frac{\int_0^{\infty} R(t) dt - \int_0^{x_0} R(t) dt}{R(x_0)} = \frac{\mu - \int_0^{x_0} R(t) dt}{R(x_0)}. \quad (27)$$

Then the MRL of MTLCh distribution can be obtained as follows:

$$m(x_0) = \frac{\mu - \left[x_0 - \sum_{(j,k)=0}^{\infty} \frac{\varphi_{j,k}(\lambda, \alpha)}{\beta k} (e^{\beta k x_0} - 1) \right]}{R(x_0)}, \quad (28)$$

where $\varphi_{j,k}(\lambda, \alpha)$ is given by (16).

6. Mean past lifetime

In a real life situation, where systems often are not monitored continuously, one might be interested in getting inference more about the history of the system, e.g. when the individual components have failed. Assume that a component with lifetime X has failed at or some time before x_0 , $x_0 \geq 0$. Consider the conditional random variable $X - x_0 | X \leq x_0$. This conditional random variable shows, in fact, the time elapsed from the failure of the component given that its lifetime is less than or equal to x_0 . Hence, the MPL of the component can be defined as

$$\mu(x_0) = E(X - x_0 | X \leq x_0) = \frac{\int_0^{x_0} F(t) dt}{F(x_0)} = \frac{1}{F(x_0)} \sum_{(j,k)=0}^{\infty} \frac{\varphi_{j,k}(\lambda, \alpha)}{\beta k} (e^{\beta k x_0} - 1), \quad (29)$$

where $\varphi_{j,k}(\lambda, \alpha)$ is given by (16).

7. Order statistics

Order statistics are one of the most fundamental tools in non-parametric statistics and inference. It is used in the problems of estimation and hypothesis tests in different ways.

Suppose that $x_1, x_2, \dots, x_n$ is a random sample with pdf, $f(x)$ and cdf, $F(x)$. Let $x_{(1)} < x_{(2)} < \dots < x_{(n)}$ denote the corresponding order statistics, then $x_{(i)}$ is called the i^{th} order statistic, $i = 1, 2, \dots, n$. The pdf of the i^{th} order statistics $x_{(i)}$ of MTLCh distribution, based on (15) and (17), is

$$g(x_{(i)}) = D(i) \sum_{l_1=0}^{n-i} \sum_{(l_2, l_3)=0}^{\infty} \delta_{l_2, l_3}(\lambda, \alpha, \beta) (-1)^{(l_1 + l_2 + l_3)} \binom{n-i}{l_1} \binom{\alpha(i+l_1)-1}{l_2} e^{\beta(l_3+1)x_{(i)}},$$

$$x_{(i)} > 0; \quad i = 1, 2, \dots, n, \quad (30)$$

where

$$D(i) = \frac{n!}{(i-1)!(n-i)!} \text{ and } \delta_{l_2, l_3}(\lambda, \alpha, \beta) = \alpha\beta(2\lambda)^{(l_3+1)} \frac{(l_2+1)^{l_3}}{l_3!} e^{2\lambda(l_2+1)}. \quad (31)$$

Special Cases:

- If $i = 1$, in (30), the pdf of the smallest order statistic can be obtained.
- If $i = n$, in (30), the pdf of the largest order statistic can be obtained.
- If $i = \frac{n+1}{2}$, in (30), the pdf of the median observable in the odd sample size case can be obtained.

3. Some Related Distributions

It can be shown that the MTLCh (λ, α, β) distribution is related to a wide range of well-known distributions through variables transformation of the variable X or by considering special values of the parameters $(\lambda, \alpha, \beta)'$, such as exponentiated Chen, Chen, standard Chen, Kumaraswamy Weibull, Kumaraswamy exponential, Kumaraswamy Rayleigh, exponentiated Weibull, exponentiated exponential, Weibull, Burr Type X, left truncated exponentiated exponential, TL-beta Type II, Beta Type I, Kumaraswamy power function, exponentiated power function, power function, TL-Weibull [which was introduced by EL-Helbawy (2018 a)], TL-exponential [which was presented by Al-Shomrani *et al.* (2016)], TL-Rayleigh, TL-F, F, log-logistic, TL-Dagum, Dagum, TL-Burr Type III and Burr Type III distributions.

4. Estimation Based on Progressive Type-II Censored Sample

Let $X_{1:m:n}, X_{2:m:n}, \dots, X_{m:m:n}$ denote a progressive Type-II censored sample obtained from MTLCh (λ, α, β) distribution. The *likelihood function* (LF) is given by

$$L(\underline{\theta}; \underline{x}) = C(n, m-1) \prod_{i=1}^m f(x_{(i)}; \underline{\theta}) [1 - F(x_{(i)}; \underline{\theta})]^{R_i}, \quad (32)$$

where $\underline{\theta} = (\lambda, \alpha, \beta)'$, $\underline{x} = (x_{1:m:n}, x_{2:m:n}, \dots, x_{m:m:n})$ denotes an observed value of $\underline{X} = (X_{1:m:n}, X_{2:m:n}, \dots, X_{m:m:n})$ and $C(n, m-1) = n(n - R_1 - 1)(n - R_1 - R_2 - 2) \dots (n - R_1 - \dots - R_{m-1} - m + 1)$, with $C(n, 0) = n$.

Then substituting (10) and (11) in (32) yields

$$L(\underline{\theta}; \underline{x}) = C(n, m-1) \prod_{i=1}^m \left[2\lambda\alpha\beta e^{\left[\beta x_{(i)} + 2\lambda(1 - e^{\beta x_{(i)}})\right]} \right]$$

$$\times \prod_{i=1}^m \left[1 - \exp\left(2\lambda(1 - e^{\beta x_{(i)}})\right) \right]^{\alpha-1}$$

$$\times \prod_{i=1}^m \left\{ 1 - \left[1 - \exp \left(2\lambda(1 - e^{\beta x_{(i)}}) \right) \right]^{\alpha} \right\}^{R_i}. \quad (33)$$

The natural logarithm of $L(\underline{\theta}; \underline{x})$ is given by

$$\begin{aligned} \ell &\propto m \ln \lambda + m \ln \alpha + m \ln \beta + \beta \sum_{i=1}^m x_{(i)} + \sum_{i=1}^m \ln z_{(i)} \\ &+ (\alpha - 1) \sum_{i=1}^m \ln[1 - z_{(i)}] + \sum_{i=1}^m R_i \ln[1 - (1 - z_{(i)})^{\alpha}], \end{aligned} \quad (34)$$

where $z_{(i)} = \exp[2\lambda(1 - e^{\beta x_{(i)}})]$.

4.1 Estimation of the parameters

The ML estimators of the parameters $\underline{\vartheta} = (\lambda, \alpha, \beta)'$ can be obtained by differentiating (34) with respect to λ, α and β and then setting to zero. Hence

$$\begin{aligned} \frac{\partial \ell}{\partial \lambda} &= \frac{m}{\lambda} + \sum_{i=1}^m \left(2 - 2e^{\beta x_{(i)}} \right) - (\hat{\alpha} - 1) \sum_{i=1}^m \frac{\dot{z}_{(i)}(\hat{\lambda})}{[1 - z_{(i)}]} \\ &+ \hat{\alpha} \sum_{i=1}^m R_i \frac{\dot{z}_{(i)}(\hat{\lambda}) (1 - z_{(i)})^{\hat{\alpha}-1}}{[1 - (1 - z_{(i)})^{\hat{\alpha}}]} = 0, \end{aligned} \quad (35)$$

where $\dot{z}_{(i)}(\lambda) = \frac{\partial z_{(i)}}{\partial \lambda} = z_{(i)} (2 - 2e^{\beta x_{(i)}})$.

$$\frac{\partial \ell}{\partial \alpha} = \frac{m}{\alpha} + \sum_{i=1}^m \ln[1 - z_{(i)}] - \sum_{i=1}^m R_i \frac{\ln[1 - z_{(i)}] (1 - z_{(i)})^{\hat{\alpha}}}{[1 - (1 - z_{(i)})^{\hat{\alpha}}]} = 0, \quad (36)$$

and

$$\begin{aligned} \frac{\partial \ell}{\partial \beta} &= \frac{m}{\beta} + \sum_{i=1}^m x_{(i)} - 2\hat{\lambda} \sum_{i=1}^m x_{(i)} e^{\beta x_{(i)}} - (\hat{\alpha} - 1) \sum_{i=1}^m \frac{\dot{z}_{(i)}(\hat{\beta})}{[1 - z_{(i)}]} \\ &+ \hat{\alpha} \sum_{i=1}^m R_i \frac{\dot{z}_{(i)}(\hat{\beta}) (1 - z_{(i)})^{\hat{\alpha}-1}}{[1 - (1 - z_{(i)})^{\hat{\alpha}}]} = 0, \end{aligned} \quad (37)$$

where $\dot{z}_{(i)}(\beta) = \frac{\partial z_{(i)}}{\partial \beta} = -2\lambda x_{(i)} z_{(i)} e^{\beta x_{(i)}}$.

A system of non-linear likelihood Equations (35)-(37) can be solved numerically using the Newton-Raphson method, to obtain the ML estimators of the parameters λ, α and β .

4.2 Estimation of the reliability, hazard and reversed hazard rate functions

The invariance property of the ML estimators enables us to obtain the ML estimators of the rf, hrf and rhrf by replacing the parameters λ, α and β by their ML estimators in (18), (21) and (22), respectively, as follows:

$$\hat{R}(x; \hat{\lambda}, \hat{\alpha}, \hat{\beta}) = 1 - \left[1 - \exp \left(2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right) \right]^{\hat{\alpha}}, \quad x > 0; (\hat{\lambda}, \hat{\alpha}, \hat{\beta} > 0), \quad (38)$$

$$\hat{h}(x; \hat{\lambda}, \hat{\alpha}, \hat{\beta}) = \frac{2\hat{\lambda}\hat{\alpha}\hat{\beta} e^{\left[\hat{\beta}x + 2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right]} \left[1 - \exp \left(2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right) \right]^{\hat{\alpha}-1}}{1 - \left[1 - \exp \left(2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right) \right]^{\hat{\alpha}}}, \quad x > 0; (\hat{\lambda}, \hat{\alpha}, \hat{\beta} > 0), \quad (39)$$

and

$$\widehat{rh}(x; \hat{\lambda}, \hat{\alpha}, \hat{\beta}) = \frac{2\hat{\lambda}\hat{\alpha}\hat{\beta} e^{\left[\hat{\beta}x + 2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right]}}{\left[1 - \exp \left(2\hat{\lambda}(1 - e^{\hat{\beta}x}) \right) \right]}, \quad x > 0; (\hat{\lambda}, \hat{\alpha}, \hat{\beta} > 0), \quad (40)$$

where $\hat{\lambda}$, $\hat{\alpha}$ and $\hat{\beta}$ are the ML estimators of λ , α and β .

Remarks:

- If $R_i = 0, i = 1, 2, \dots, m$ and n equals m , all the results obtained for progressive Type-II censored sample reduced to the complete sample case.
- If $R_i = 0, i = 1, 2, \dots, m - 1$ and $R_m = n - m$, then progressive Type-II censored sampling reduces to traditional Type-II censoring.
- Let $X_1, X_2, \dots, X_n$ be a random sample of size n drawn from a MTLCh (λ, α, β) distribution with pdf given by (10), where λ and β are known and α is unknown, one can obtain a complete minimal sufficient statistic for the parameter α , by using the exponential family. The complete minimal sufficient statistic for α is as follows:

$$s = \sum_{i=1}^n \ln[1 - \exp[2\lambda(1 - e^{\beta x})]]. \quad (41)$$

- Let $\hat{Y} = -\ln[1 - \exp[2\lambda(1 - e^{\beta x})]]$, then the pdf of $\hat{Y} \sim \text{exponential}(\alpha)$, if $\hat{s} = \sum \hat{Y}$, then the distribution of $\hat{s}$ is a gamma distribution and $\hat{s}$ is a complete minimal sufficient statistic, hence the confidence interval for the parameter α can be obtained as follows:

$$P \left[\frac{\chi^2_{(2n, \tau/2)}}{2\hat{s}} < \alpha < \frac{\chi^2_{(2n, 1-\tau/2)}}{2\hat{s}} \right] = 1 - \tau. \quad (42)$$

4.3 Asymptotic confidence intervals

The asymptotic Fisher information matrix $\hat{I}_{ij}(\vartheta)$ for the ML estimators of the parameters $\vartheta = (\lambda, \alpha, \beta)'$ is the 3×3 symmetric matrix of the negative second partial derivatives of log-LF, given by (34) with respect to the parameters $\vartheta = (\lambda, \alpha, \beta)'$. That

$$\text{is } \hat{I}_{ij}(\vartheta) \approx - \left[\frac{\partial^2 \ell}{\partial \vartheta_i \partial \vartheta_j} \right]_{\vartheta=\hat{\vartheta}}, \quad i, j = 1, 2, 3, \quad (43)$$

where $\vartheta_1 = \lambda, \vartheta_2 = \alpha$ and $\vartheta_3 = \beta$.

The asymptotic variance-covariance matrix $\widehat{V}$ for the ML estimators is the inverse of the asymptotic Fisher information matrix.

The asymptotic Fisher information matrix enables us to construct ACIs for the parameters based on the limiting normal distribution, where $V(\hat{\vartheta}_i) = \sigma_{ii}$, the i th diagonal element of the matrix $\widehat{V}$.

ACIs for the parameters, rf, hrf and rhrf can be obtained by using the asymptotic normality of the ML estimators. Hence, a two sided approximate $100(1 - \tau) \%$ CIs for ω is given by

$$L_{\omega} = \widehat{\omega} - Z_{(1-\tau/2)} \sqrt{\widehat{var}(\widehat{\omega})} \quad \text{and} \quad U_{\omega} = \widehat{\omega} + Z_{(1-\tau/2)} \sqrt{\widehat{var}(\widehat{\omega})}, \quad (48)$$

where L_{ω} and U_{ω} are the *lower limit* (LL) and *upper limit* (UL), $\widehat{\omega}$ is $\underline{\hat{\vartheta}} = \hat{\lambda}, \hat{\alpha}, \hat{\beta}, \hat{R}(x), \hat{h}(x)$ or $\hat{r}h(x)$, respectively, and $Z_{(1-\tau/2)}$ is a standard normal variate and τ is the confidence coefficient.

5. Numerical Illustration

This section aims to investigate the precision of the theoretical results of estimation on basis of simulated and real data.

5.1 Simulation study

In this subsection, a simulation study is conducted to illustrate the performance of the presented ML estimates on the basis of generated data from the MTLCh distribution. The ML averages of the parameters, rf and hrf based on progressive Type-II censoring scheme are computed. Moreover, CIs of the parameters, rf and hrf are calculated. Simulation studies are performed using Mathematica 9 for illustrating the obtained results.

Applying the algorithm in Balakrishnan and Sandhu (1995), the following steps are used to generate a progressive Type-II censored sample from the MTLCh distribution as follows:

Step 1: Generate m independent $U(0,1)$ random variables $U_1, U_2, \dots, U_m$.

Step 2: For given values of the progressive censoring scheme $R_1, R_2, \dots, R_m$, set $Y_i = U_i^{1/(i+\sum_{j=m-i+1}^m R_j)}$, for $i = 1, 2, \dots, m$.

Step 3: Set $UP_i = 1 - (Y_m Y_{m-1} Y_{m-2} \dots Y_{m-i+1})$, $i = 1, 2, \dots, m$. Then, $UP_1, UP_2, \dots, UP_m$ are progressive Type-II censored sample of size m from $U(0,1)$ distribution.

Step 4: For given values of the parameters λ , α and β , the inverse cdf method, can be used to generate m progressive Type-II censored sample from MTLCh whose cdf is given by (11). Thus, by solving the nonlinear equation

$$x_i = \frac{1}{\beta} \ln \left[1 - \frac{1}{2\lambda} \ln \left(1 - UP_i^{1/\alpha} \right) \right], \quad i = 1, 2, \dots, m,$$

the resulting set, $(x_1, x_2, \dots, x_m)$ is the required progressive Type-II censored sample of size m from MTLCh distribution and this obtained sample is ordered.

Step 5: Repeat all the previous steps N times where N represents a fixed number of simulated samples.

Evaluating the performance of the estimates is considered through some measurements of accuracy. In order to study the precision and variation of the estimates, it is convenient to use the *estimated risk* (ER) = $\frac{\sum_{i=1}^N (\text{estimated value} - \text{true value})^2}{N}$ and *relative error* (RE) = $\frac{\sqrt{\text{ER}}}{\text{true value}}$.

Simulation results of ML estimates are displayed in Tables 2-8, where $N = 2000$ is the number of repetitions and the samples of size ($n=30, 60, 100$). For each sample size, different ratio of effective sample sizes, $r = \frac{m}{n} = 0.2, 0.5$ and 0.8 , and set of different samples schemes, where

Scheme I: $R_1 = n - m$ and $R_2 = R_3 = \dots = R_m = 0$. It is Type-II censoring.

Scheme II: $R_1 = R_2 = \dots = R_{m-1} = 1$ and $R_m = n - 2m + 1$, for $r = 0.2$ and 0.5 . If $r = 0.8$, then $R_{i+1} = 1$ and $R_j = 0$; $i = 0, 4, 8, 12, \dots, m$, $i \neq j$.

Scheme III: $0, 0, \dots, 0, \frac{n-m}{3}, 0, 0, \dots, 0, \frac{n-m}{3}, 0, 0, \dots, 0, \frac{n-m}{3}$.

The best scheme is the scheme which minimizes the ERs, REs and length of CIs of the estimators.

Tables 2-4 display the ML averages, ERs, REs and CIs of the unknown parameters based on progressive Type-II censoring under different samples schemes. While Tables 5-7 present the ML averages, ERs, REs and CIs of the rf and hrf based on progressive Type-II censoring under different samples schemes.

Table 8 presents ML averages of rf and hrf under different time x_0 at different samples size based on progressive Type-II censoring under different samples schemes.

5.2 Some applications

The main aim of this subsection is to demonstrate how the proposed MTLCh distribution based on progressive Type-II censoring can be used in practice. Two real

lifetime data sets are used for this purpose. The MTLCh distribution is fitted to the two real data using Kolmogorov-Smirnov goodness of fit test through Mathematica 9.

Application 1:

The first application is given by Murthy *et al.* (2004). The data refers to the time between failures for a repairable item: 1.43, 0.11, 0.71, 0.77, 2.63, 1.49, 3.46, 2.46, 0.59, 0.74, 1.23, 0.94, 4.36, 0.40, 1.74, 4.73, 2.23, 0.45, 0.70, 1.06, 1.46, 0.30, 1.82, 2.37, 0.63, 1.23, 1.24, 1.97, 1.86 and 1.17.

Application 2:

The second data set is given by Nelson (1982). The data refers to the time to breakdown of an insulating fluid between electrodes at a voltage of 34 k.v. (minutes). The 19 times to breakdown are: 0.96, 4.15, 0.19, 0.78, 8.01, 31.75, 7.35, 6.50, 8.27, 33.91, 32.52, 3.16, 4.85, 2.78, 4.67, 1.31, 12.06, 36.71 and 72.89.

The Kolmogorov–Smirnov goodness of fit test is applied to check the validity of the introduced fitted model. The p values are given, respectively, 0.062 and 0.119. The p value given in each case showed that the proposed model fits the data very well.

Table 9 displays the ML estimates and REs of the unknown parameters for the real data sets based on progressive Type-II censoring under different samples schemes.

5.3 Concluding remarks

- It is noticed, from Tables 2-4, that the ML averages are very close to the initial values of the parameters as the sample size increases. Also, ERs and REs are decreasing when the sample size is increasing. This is indicative of the fact that the estimates are consistent and approaches the population parameter values as the sample size increases.
- From Tables 5-7, it is noticed that the ERs and REs of the rf and hrf are decreasing when the sample size is increasing.
- The lengths of the CIs of the parameters, rf and hrf become narrower as the sample size increases.
- It is noticed, from Tables 2-7, that the ERs and REs are decreasing as the ratio of effective samples size (r) increases.
- From Table 8 the estimated values of the rf decreases when the time x_0 increases, while the estimated values of the hrf are monotonically decreasing then it get approximately constant and it increases (taking a bathtub shape) when the time x_0 increases. The results get better when the sample size increases.
- Scheme I is the best censoring scheme where it has the lowest ERs and REs and the narrower lengths of the CIs.

6. General Conclusion

In this study, a class of distributions called the MTLCh (λ, α, β) distribution is introduced. This class of distributions include several special models discussed such as the exponentiated Chen, Chen and standard Chen. Some transformed distributions are derived such as Kumaraswamy Weibull, Kumaraswamy exponential, Kumaraswamy, exponentiated Weibull, exponentiated exponential, Weibull, Burr Type X, left truncated exponentiated exponential, TL-beta Type II, Kumaraswamy power function, exponentiated power function, power function, TL-Weibull, TL- exponential, TL-Rayleigh, TL-F, TL-Dagum, Dagum, TL-Burr Type III and Burr Type III distributions, among several others. From the graphical description of the MTLCh distribution, one can observe that it has several figures, therefore the MTLCh distribution is a very flexible reliability model. Some statistical properties of this distribution are derived. Also, the ML estimators for the parameters, rf and hrf of the MTLCh distribution based on progressive Type-II censored data are obtained. Monte Carlo simulation is performed. In general, ML estimators of the model parameters, rf and hrf when the sample size is n and the ratio of effective samples of size (r) increases, the ERs, REs and the lengths of the CIs decrease. Also, when the different values of time x_0 increases, the rf decreases while the hrf monotonically decreases then it approximately stay constant and it increases (taking a bathtub shape). Bayesian estimation and prediction bounds would be useful as a basis for further researches in distribution theory. The work is in progress.

Acknowledgments:

The authors are thankful to the Editor and Referees for their constructive comments that led to improvements to the original manuscript.

References

- AL-Hussaini, E. K. (2012).** Composition of cumulative distribution functions. *Journal of Statistical Theory and Applications*. Vol. 11, No. 4, pp. 323-336.
- Almetwally, E. M. and Almongy, H. M. (2018).** Bayesian estimation of the generalized power Weibull distribution parameters based on progressive censoring schemes. *International Journal of Mathematical Archive*. Vol. 9, pp. 1-8.
- Almetwally, E. M., Mubarak, A. E. and Almongy, H. M. (2018).** Bayesian and maximum likelihood estimation for the Weibull generalized exponential distribution parameters using progressive censoring schemes. *Pakistan Journal of Statistics and Operation Research*. Vol. 14, pp. 853-868. Doi: 10.18187/pjsor.v14i4.2600.
- Al-Shomrani, A., Arif, O., Shawky, A., Hanif, S. and Shahbaz, M. Q. (2016).** Topp–Leone family of distributions: some properties and application. *Pakistan Journal of Statistics and Operation Research*. Vol. 12, pp. 433-451.
- Arshad, M. and Jamal, Q. A. (2019).** Statistical inference for Topp–Leone-generated family of distributions based on records. *Journal of Statistical Theory and Applications*. Vol. 18, pp. 65–78. Doi: <https://doi.org/10.2991/jsta.d.190306.008>.

- 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*. Boston, Birkhauser.
Doi: 10.1007/978-0-8176-4807-7.
- Balakrishnan, N. and Sandhu, R. A. (1995).** A simple simulational algorithm for generating progressive Type-II censored samples. *The American Statistician*. Vol. 49, No. 2, pp. 229-230.
- Chaubey, Y. P. and Zhang, R. (2015).** An extension of Chen's family of survival distributions with bathtub shape or increasing hazard rate function. *Communications in Statistics-Theory and Methods*. Vol. 44, pp. 4049-4064.
Doi: 10.1080/03610926.2014.997357.
- Chen, Z. (2000).** A new two-parameter lifetime distribution with bathtub shape or increasing failure rate function. *Statistics and Probability Letters*. Vol. 49, pp. 155-161.
- Cordeiro, G. M., Hashimoto, E. M. and Ortega, E. M. (2014).** The mcdonald Weibull model. *A Journal of Theoretical and Applied Statistics*. Vol. 48, pp. 256-278.
- Dey, S., Kumar, D., Ramos, P. L. and Louzada, F. (2017).** Exponentiated Chen distribution: properties and estimation. *Communications in Statistics - Simulation and Computation*. Vol. 0, pp.1-22. Doi.org/10.1080/03610918.2016.1267752.
- EL-Helbawy, A. A. (2018 a).** Bayesian estimation for Topp-Leone Weibull distribution based on dual generalized order statistics. *Academy of Business Journal, AL-Azhar University, An Academic-Periodical Refereed Journal*.
- EL-Helbawy, A. A. (2018 b).** Estimation under a finite mixture of two-component Topp-Leone Rayleigh lifetime model. *Scientific Journal of Faculty of Commerce, Assuit University*.
- Eugene, N., Lee, C. and Famoye, F. (2002).** Beta-normal distribution and its applications. *Communications in Statistics - Theory and Methods*. Vol. 31, pp. 497–512.
- Karakoca, A. and Pekgör, A. (2019).** Maximum likelihood estimation of the parameters of progressively Type-II censored samples from Weibull distribution using Genetic Algorithm. *Academic Platform Journal of Engineering and Science*. Vol. 7. Doi: 10.21541/apjes.452564.
- Kayal, T., Tripathi, Y. M., Singh, D. P. and Rastogi, M. K. (2016).** Estimation and prediction for Chen distribution with bathtub shape under progressive censoring. *Journal of Statistical Computation and Simulation*. pp. 1-20.
Doi.org/10.1080/00949655.2016.1209199.
- Murthy, D. N. P., Xie, M. and Jiang, R. (2004).** *Weibull Models*. John Wiley and Sons, Inc., Hoboken, New Jersey.
- Nadarajah, S. and Kotz, S. (2003).** Moments of some J-shaped distributions. *Journal of Applied Statistics*. Vol. 30, pp. 311-317.

- Nassar, M., Dey, S. and Kumar, D. (2018).** A new generalization of the exponentiated Pareto distribution with an application. *American Journal of Mathematical and Management Sciences*. Vol. 37, pp. 1-29.
- Nelson, W. (1982).** *Applied Life Data Analysis*. John Wiley and Sons, Inc., New York.
- Oguntunde, P. E., Khaleel, M. A., Okagbue, H. I. and Odetunmbi, O. A. (2019).** The Topp–Leone Lomax distribution with applications to airborne communication transceiver dataset. *Wireless Personal Communications*. pp. 1-12. Doi.org/10.1007/s11277-019-06568-8.
- Prudnikov, A. P., Brychkov, Y. A. and Marichev, O. I. (1986).** *Integrals and Series*. New York: Gordon and Breach Science Publishers. Vol. 3.
- Reyad, H., Korkmaz, M. C., Afify, A. Z., Hamedani, G. G. and Othman, S. (2019).** The Fréchet Topp Leone-G family of distributions: properties, characterizations and applications. *Annals of Data Science*. pp. 1-22. Doi.org/10.1007/s40745-019-00212-9.
- Topp, C. W. and Leone, F. C. (1955).** A family of J-shaped frequency functions. *Journal of the American Statistical Association*. Vol. 50, pp. 209-219.

Table 1
Summary of some transformations applied to the MTLCh and resulting distributions

Transformations	Resulting Distribution	pdf	Range
$[2e^{\beta x} - 2]^{1/\theta_1}$	Exponentiated Weibull $(\lambda, \alpha, \theta_1)$	$\lambda \alpha \theta_1 y_1^{\theta_1-1} e^{-\lambda y_1^{\theta_1}} \times [1 - e^{-\lambda y_1^{\theta_1}}]^{\alpha-1}$	$y_1 > 0$
$[b - \theta_2(2 - 2e^{\beta x})]$	Left truncated exponentiated exponential $(\lambda, \alpha, \theta_2, b)$	$\frac{\lambda \alpha}{\theta_2} e^{-\lambda(\frac{y_2-b}{\theta_2})} \times [1 - e^{-\lambda(\frac{y_2-b}{\theta_2})}]^{\alpha-1}$	$b < y_2 < \infty$
$\left[\left(e^{2\lambda(1-e^{\beta x})} \right)^{\frac{-1}{2\theta_3}} - 1 \right]$	TL-beta Type II (α, θ_3)	$2\alpha \theta_3 (1 + y_3)^{-(2\theta_3+1)} \times [1 - (1 + y_3)^{-2\theta_3}]^{\alpha-1}$	$y_3 > 0$
$[e^{2\lambda(1-e^{\beta x})}]$	Beta Type I $(1, \alpha)$	$\alpha [1 - y_4]^{\alpha-1}$	$0 < y_4 < 1$
$[e^{2\lambda(1-e^{\beta x})}]^{\frac{1}{\theta_4 \theta_5}}$	Kumaraswamy power function $(\alpha, \theta_4, \theta_5)$	$\alpha \theta_4 \theta_5 y_5^{\theta_4 \theta_5 - 1} [1 - y_5^{\theta_4 \theta_5}]^{\alpha-1}$	$0 < y_5 < 1$
$[e^{\beta x} - 1]^{\frac{1}{\theta_6}}$	TL-Weibull $(\lambda, \alpha, \theta_6)$	$2\lambda \alpha \theta_6 y_6^{\theta_6-1} e^{-2\lambda y_6^{\theta_6}} \times [1 - e^{-2\lambda y_6^{\theta_6}}]^{\alpha-1}$	$y_6 > 0$
$\theta_7 \left[\left(e^{2\lambda(1-e^{\beta x})} \right)^{\frac{-1}{2\theta_7}} - 1 \right]$	TL-F (α, θ_7)	$2\alpha \left[1 + \frac{y_7}{\theta_7} \right]^{-(2\theta_7+1)} \times \left[1 - \left[1 + \frac{y_7}{\theta_7} \right]^{-2\theta_7} \right]^{\alpha-1}$	$y_7 > 0$
$\left[\frac{\left(e^{2\lambda(1-e^{\beta x})} \right)^{\frac{-1}{2\theta_{10}}} - 1}{\theta_9} \right]^{\frac{-1}{\theta_8}}$	TL-Dagum $(\alpha, \theta_8, \theta_9, \theta_{10})$	$2\alpha \theta_8 \theta_9 \theta_{10} y_8^{-(\theta_8+1)} \times [1 + \theta_9 y_8^{-\theta_8}]^{-(2\theta_{10}+1)} \times [1 - (1 + \theta_9 y_8^{-\theta_8})^{-2\theta_{10}}]^{\alpha-1}$	$y_8 > 0$

Table 2

ML averages, estimated risks, relative errors and 95% confidence intervals of the MTLCh parameters based on progressive Type II censoring under Scheme I
(N=2000, $r = 0.2, 0.5$ and 0.8 , $\lambda = 0.63, \alpha = 0.48$ and $\beta = 0.37$)

n	m	$\underline{\theta}$	Average	ER	RE	UL	LL	Length
30	6	λ	0.7335	0.0546	0.3710	1.1443	0.3227	0.8215
		α	0.5206	0.0129	0.2363	0.7282	0.3130	0.4152
		β	0.4349	0.0285	0.4566	0.7406	0.1292	0.6114
	15	λ	0.6766	0.0208	0.2287	0.9438	0.4094	0.5344
		α	0.5049	0.0103	0.2120	0.6982	0.3115	0.3867
		β	0.3960	0.0101	0.2715	0.5862	0.2059	0.3802
	24	λ	0.6611	0.0164	0.2033	0.9047	0.4177	0.4870
		α	0.4986	0.0093	0.2009	0.6841	0.3131	0.3709
		β	0.3905	0.0070	0.2269	0.5500	0.2309	0.3191
60	12	λ	0.6789	0.0248	0.2498	0.9721	0.3856	0.5865
		α	0.4964	0.0070	0.1750	0.6579	0.3348	0.3231
		β	0.4024	0.0126	0.3036	0.6132	0.1915	0.4217
	30	λ	0.6484	0.0149	0.1938	0.8850	0.4118	0.4732
		α	0.4889	0.0060	0.1621	0.6404	0.3375	0.3029
		β	0.3837	0.0055	0.2002	0.5264	0.2410	0.2854
	48	λ	0.6463	0.0125	0.1777	0.8634	0.4292	0.4342
		α	0.4866	0.0050	0.1481	0.6253	0.3478	0.2775
		β	0.3794	0.0036	0.1629	0.4961	0.2627	0.2334
100	20	λ	0.6678	0.0187	0.2172	0.9255	0.4101	0.5154
		α	0.4881	0.0055	0.1545	0.6326	0.3436	0.2890
		β	0.3883	0.0078	0.2391	0.5579	0.2186	0.3393
	50	λ	0.6458	0.0122	0.1750	0.8597	0.4319	0.4278
		α	0.4846	0.0038	0.1279	0.6046	0.3646	0.2399
		β	0.3776	0.0034	0.1575	0.4908	0.2644	0.2264
	80	λ	0.6461	0.0102	0.1606	0.8419	0.4504	0.3915
		α	0.4860	0.0029	0.1133	0.5919	0.3800	0.2119
		β	0.3727	0.0024	0.1325	0.4686	0.2767	0.1919

Table 3

ML averages, estimated risks, relative errors and 95% confidence intervals of the MTLCh parameters based on progressive Type II censoring under Scheme II
($N=2000, r = 0.2, 0.5$ and $0.8, \lambda = 0.63, \alpha = 0.48$ and $\beta = 0.37$)

n	m	$\underline{\theta}$	Average	ER	RE	UL	LL	Length
30	6	λ	0.7644	0.1116	0.5302	1.3638	0.1650	1.1988
		α	0.5320	0.0174	0.2751	0.7699	0.2941	0.4758
		β	0.4477	0.0402	0.5418	0.8099	0.0856	0.7243
	15	λ	0.6866	0.0292	0.2714	1.0028	0.3704	0.6324
		α	0.5136	0.0115	0.2238	0.7136	0.3136	0.4000
		β	0.4027	0.0138	0.3179	0.6241	0.1812	0.4429
	24	λ	0.6689	0.0174	0.2091	0.9157	0.4222	0.4934
		α	0.5031	0.0087	0.1942	0.6800	0.3261	0.3539
		β	0.3898	0.0075	0.2343	0.5552	0.2244	0.3308
60	12	λ	0.7093	0.0704	0.4211	1.2055	0.2132	0.9923
		α	0.5118	0.0113	0.2221	0.7113	0.3123	0.3990
		β	0.4154	0.0253	0.4301	0.7143	0.1164	0.5980
	30	λ	0.6603	0.0176	0.2107	0.9136	0.4069	0.5068
		α	0.4960	0.0063	0.1661	0.6491	0.3429	0.3062
		β	0.3870	0.0072	0.2291	0.5498	0.2241	0.3256
	48	λ	0.6482	0.0123	0.1759	0.8625	0.4339	0.4286
		α	0.4908	0.0046	0.1416	0.6223	0.3593	0.2630
		β	0.3822	0.0040	0.1718	0.5045	0.2600	0.2445
100	20	λ	0.6731	0.0449	0.3365	1.0800	0.2663	0.8136
		α	0.4947	0.0078	0.1836	0.6650	0.3244	0.3406
		β	0.3951	0.0166	0.3484	0.6429	0.1473	0.4956
	50	λ	0.6495	0.0126	0.1785	0.8666	0.4324	0.4342
		α	0.4891	0.0038	0.1287	0.6089	0.3693	0.2396
		β	0.3810	0.0048	0.1879	0.5155	0.2465	0.2690
	80	λ	0.6424	0.0093	0.1532	0.8301	0.4547	0.3753
		α	0.4859	0.0027	0.1081	0.5870	0.3849	0.2021
		β	0.3752	0.0026	0.1386	0.4752	0.2753	0.1999

Table 4

ML averages, estimated risks, relative errors and 95% confidence intervals of the MTLCh parameters based on progressive Type II censoring under Scheme III
($N=2000$, $r = 0.2, 0.5$ and 0.8 , $\lambda = 0.63$, $\alpha = 0.48$ and $\beta = 0.37$)

n	m	ϑ	Average	ER	RE	UL	LL	Length
30	6	λ	0.7713	0.1086	0.5230	1.3547	0.1878	1.1669
		α	0.5346	0.0180	0.2799	0.7752	0.2941	0.4812
		β	0.4507	0.0405	0.5436	0.8118	0.0895	0.7223
	15	λ	0.6971	0.0325	0.2863	1.0253	0.3688	0.6565
		α	0.5187	0.0123	0.2307	0.7220	0.3154	0.4066
		β	0.4097	0.0146	0.3271	0.6337	0.1856	0.4481
	24	λ	0.6646	0.0167	0.2054	0.9090	0.4201	0.4889
		α	0.5049	0.0089	0.1966	0.6833	0.3265	0.3568
		β	0.3890	0.0072	0.2302	0.5517	0.2263	0.3254
60	12	λ	0.6994	0.0616	0.3939	1.1664	0.2323	0.9341
		α	0.5051	0.0109	0.2176	0.7039	0.3063	0.3975
		β	0.4085	0.0233	0.4127	0.6981	0.1189	0.5792
	30	λ	0.6658	0.0168	0.2058	0.9100	0.4216	0.4884
		α	0.4990	0.0066	0.1692	0.6538	0.3442	0.3096
		β	0.3904	0.0074	0.2325	0.5542	0.2265	0.3277
	48	λ	0.6445	0.0113	0.1688	0.8510	0.4380	0.4130
		α	0.4917	0.0046	0.1412	0.6225	0.3609	0.2616
		β	0.3808	0.0038	0.1672	0.5002	0.2614	0.2388
100	20	λ	0.6726	0.0421	0.3257	1.0660	0.2792	0.7868
		α	0.4937	0.0075	0.1811	0.6619	0.3255	0.3365
		β	0.3950	0.0162	0.3443	0.6398	0.1502	0.4896
	50	λ	0.6495	0.0116	0.1707	0.8568	0.4421	0.4147
		α	0.4904	0.0042	0.1351	0.6159	0.3649	0.2509
		β	0.3794	0.0048	0.1877	0.5143	0.2446	0.2697
	80	λ	0.6436	0.0089	0.1496	0.8264	0.4609	0.3655
		α	0.4876	0.0030	0.1136	0.5935	0.3817	0.2117
		β	0.3757	0.0025	0.1355	0.4733	0.2781	0.1952

Table 5

ML averages, estimated risks, relative errors and 95% confidence intervals of the reliability and hazard rate functions based on progressive Type II censoring under Scheme I

(N=2000, $r = 0.2, 0.5$ and $0.8, \lambda = 0.63, \alpha = 0.48, \beta = 0.37$ and $x_0 = 0.1$)

n	m	rf and hrf	Average	ER	RE	UL	LL	Length
30	6	$R(x_0)$	0.7593	0.0043	0.0851	0.8858	0.6327	0.2531
		$h(x_0)$	1.6149	0.3059	0.3900	2.6280	0.6018	2.0261
	15	$R(x_0)$	0.7671	0.0040	0.0823	0.8912	0.6429	0.2483
		$h(x_0)$	1.4752	0.0984	0.2212	2.0796	0.8707	1.2089
	24	$R(x_0)$	0.7670	0.0039	0.0806	0.8886	0.6455	0.2431
		$h(x_0)$	1.4562	0.0720	0.1892	1.9768	0.9356	1.0412
60	12	$R(x_0)$	0.7627	0.0030	0.0714	0.8693	0.6560	0.2132
		$h(x_0)$	1.5072	0.1095	0.2334	2.1320	0.8824	1.2495
	30	$R(x_0)$	0.7671	0.0026	0.0664	0.8672	0.6670	0.2002
		$h(x_0)$	1.4432	0.0494	0.1567	1.8760	1.0103	0.8657
	48	$R(x_0)$	0.7674	0.0022	0.0611	0.8594	0.6753	0.1840
		$h(x_0)$	1.4372	0.0354	0.1327	1.8042	1.0702	0.7340
100	20	$R(x_0)$	0.7629	0.0026	0.0658	0.8612	0.6647	0.1964
		$h(x_0)$	1.4807	0.0700	0.1866	1.9846	0.9767	1.0079
	50	$R(x_0)$	0.7678	0.0018	0.0548	0.8504	0.6852	0.1653
		$h(x_0)$	1.4357	0.0322	0.1265	1.7855	1.0858	0.6997
	80	$R(x_0)$	0.7706	0.0013	0.0461	0.8403	0.7009	0.1394
		$h(x_0)$	1.4218	0.0215	0.1034	1.7092	1.1344	0.5748

Table 6

ML averages, estimated risks, relative errors and 95% confidence intervals of the reliability and hazard rate functions based on progressive Type II censoring under Scheme II

(N=2000, $r = 0.2, 0.5$ and 0.8 , $\lambda = 0.63$, $\alpha = 0.48$, $\beta = 0.37$ and $x_0 = 0.1$)

n	m	rf and hrf	Average	ER	RE	UL	LL	Length
30	6	$R(x_0)$	0.7567	0.0044	0.0864	0.8843	0.6291	0.2551
		$h(x_0)$	1.7112	0.5442	0.5202	3.0381	0.3843	2.6538
	15	$R(x_0)$	0.7699	0.0031	0.0723	0.8792	0.6606	0.2186
		$h(x_0)$	1.4886	0.1206	0.2449	2.1552	0.8219	1.3333
	24	$R(x_0)$	0.7701	0.0030	0.0709	0.8773	0.6630	0.2143
		$h(x_0)$	1.4518	0.0650	0.1798	1.9473	0.9563	0.9910
60	12	$R(x_0)$	0.7638	0.0026	0.0661	0.8626	0.6649	0.1977
		$h(x_0)$	1.5840	0.3035	0.3885	2.6137	0.5542	2.0595
	30	$R(x_0)$	0.7693	0.0020	0.0576	0.8563	0.6823	0.1740
		$h(x_0)$	1.4529	0.0578	0.1695	1.9191	0.9867	0.9324
	48	$R(x_0)$	0.7698	0.0018	0.0550	0.8528	0.6868	0.1661
		$h(x_0)$	1.4350	0.0338	0.1296	1.7937	1.0763	0.7174
100	20	$R(x_0)$	0.7646	0.0018	0.0554	0.8474	0.6818	0.1656
		$h(x_0)$	1.5186	0.1846	0.3030	2.3373	0.6999	1.6374
	50	$R(x_0)$	0.7696	0.0012	0.0459	0.8389	0.7003	0.1386
		$h(x_0)$	1.4395	0.0342	0.1305	1.7997	1.0794	0.7203
	80	$R(x_0)$	0.7707	0.0011	0.0436	0.8366	0.7047	0.1319
		$h(x_0)$	1.4235	0.0213	0.1029	1.7094	1.1377	0.5717

Table 7

ML averages, estimated risks, relative errors and 95% confidence intervals of the reliability and hazard rate functions based on progressive Type II censoring under Scheme III

($N=2000$, $r = 0.2, 0.5$ and 0.8 , $\lambda = 0.63$, $\alpha = 0.48$, $\beta = 0.37$ and $x_0 = 0.1$)

n	m	rf and hrf	Average	ER	RE	UL	LL	Length
30	6	$R(x_0)$	0.7564	0.0045	0.0871	0.8849	0.6279	0.2570
		$h(x_0)$	1.7171	0.5417	0.5190	3.0352	0.3989	2.6363
	15	$R(x_0)$	0.7691	0.0030	0.0708	0.8760	0.6622	0.2138
		$h(x_0)$	1.5123	0.1358	0.2599	2.2107	0.8139	1.3968
	24	$R(x_0)$	0.8080	0.0026	0.0637	0.9088	0.7074	0.2014
		$h(x_0)$	1.6500	0.0861	0.1795	2.2245	1.0754	1.1491
60	12	$R(x_0)$	0.7624	0.0027	0.0679	0.8636	0.6612	0.2025
		$h(x_0)$	1.5679	0.2723	0.3680	2.5476	0.5881	1.9594
	30	$R(x_0)$	0.7690	0.0019	0.0566	0.8545	0.6834	0.1710
		$h(x_0)$	1.4656	0.0625	0.1763	1.9466	0.9845	0.9621
	48	$R(x_0)$	0.8072	0.0014	0.0473	0.8820	0.7324	0.1496
		$h(x_0)$	1.6399	0.0472	0.1328	2.0655	1.2142	0.8513
100	20	$R(x_0)$	0.7640	0.0018	0.0551	0.8462	0.6818	0.1644
		$h(x_0)$	1.5183	0.1743	0.2944	2.3127	0.7239	1.5888
	50	$R(x_0)$	0.7704	0.0012	0.0457	0.8394	0.7013	0.1381
		$h(x_0)$	1.4366	0.0353	0.1324	1.8029	1.0702	0.7327
	80	$R(x_0)$	0.8070	0.0010	0.0390	0.8686	0.7454	0.1232
		$h(x_0)$	1.6387	0.0310	0.1077	1.9840	1.2934	0.6906

Table 8

ML averages of the reliability and hazard rate functions under different time x_0 , different samples size based on progressive Type II censoring under different samples schemes

n	m	x_0	Average					
			Scheme I		Scheme II		Scheme III	
			$R(x_0)$	$h(x_0)$	$R(x_0)$	$h(x_0)$	$R(x_0)$	$h(x_0)$
30	6	0.1	0.7593	1.6149	0.7578	1.7160	0.7587	1.7069
		0.6	0.4226	1.1960	0.4070	1.3962	0.4132	1.3667
		1	0.2828	1.3309	0.2700	1.6511	0.2858	1.5648
	15	0.1	0.7671	1.4752	0.7675	1.5056	0.7683	1.5051
		0.6	0.4530	0.9674	0.4451	1.0277	0.4469	1.0391
		1	0.3093	1.0270	0.3102	1.0776	0.3044	1.1176
	24	0.1	0.7670	1.4562	0.7695	1.4478	0.7698	1.4530
		0.6	0.4589	0.9235	0.4584	0.9320	0.4572	0.9396
		1	0.3227	0.9404	0.3219	0.9474	0.3222	0.9601
60	12	0.1	0.7627	1.5072	0.7617	1.5885	0.7612	1.5929
		0.6	0.4430	1.0138	0.4380	1.1391	0.4342	1.1384
		1	0.3051	1.0532	0.2921	1.3166	0.3056	1.2190
	30	0.1	0.7671	1.4432	0.7676	1.4612	0.7692	1.4586
		0.6	0.4598	0.9115	0.4557	0.9381	0.4564	0.9419
		1	0.3224	0.9223	0.3202	0.9597	0.3174	0.9747
	48	0.1	0.7674	1.4372	0.7699	1.4306	0.7697	1.4357
		0.6	0.4629	0.8910	0.4623	0.8967	0.4608	0.9045
		1	0.3265	0.8935	0.3257	0.9046	0.3266	0.9014
100	20	0.1	0.7629	1.4807	0.7642	1.5149	0.7641	1.5177
		0.6	0.4494	0.9447	0.4500	1.0337	0.4478	1.0331
		1	0.3146	0.9698	0.3138	1.1203	0.3138	1.0971
	50	0.1	0.7678	1.4357	0.7692	1.4438	0.7689	1.4409
		0.6	0.4625	0.8917	0.4613	0.9048	0.4595	0.9115
		1	0.3286	0.8831	0.3266	0.9067	0.3250	0.9169
	80	0.1	0.7706	1.4218	0.7706	1.4216	0.7700	1.4261
		0.6	0.4678	0.8750	0.4661	0.8795	0.4648	0.8818
		1	0.3312	0.8763	0.3293	0.8855	0.3297	0.8829

Table 9

ML estimates and relative errors of the parameters for the real data sets based on progressive Type II censoring under different samples schemes

	n	m	$\underline{\theta}$	Scheme I		Scheme II		Scheme III	
				Estimate	SE	Estimate	SE	Estimate	SE
Application 1	30	6	λ	5.612	0.593	0.732	0.794	2.089	0.737
			α	4.050	0.796	1.118	0.897	1.245	0.869
			β	0.081	0.728	0.041	0.862	0.053	0.825
		15	λ	2.726	0.436	2.391	0.374	3.248	0.652
			α	1.458	0.718	1.963	0.687	1.643	0.826
			β	0.113	0.624	0.125	0.583	0.070	0.768
		24	λ	3.419	0.419	2.481	0.352	3.28	0.611
			α	1.648	0.710	1.594	0.676	1.453	0.806
			β	0.116	0.613	0.13	0.568	0.078	0.741
Application 2	19	6	λ	0.747	0.858	0.072	0.934	0.747	0.858
			α	0.825	0.965	0.376	0.983	0.825	0.965
			β	0.014	0.929	0.007	0.967	0.014	0.929
		15	λ	0.840	0.835	1.588	0.925	0.737	0.813
			α	0.707	0.959	0.642	0.981	0.576	0.953
			β	0.017	0.917	0.007	0.963	0.019	0.906

Estimation and Prediction Based on Constant Stress-Partially Accelerated Life Testing for Topp Leone-Inverted Kumaraswamy Distribution

Behairy, S. M., Refaey, R. M., AL-Dayian, G. R. and EL-Helbawy, A. A.

Statistics Department, Faculty of Commerce,
AL-Azhar University (Girls' Branch), Cairo, Egypt

Abstract

In this paper constant stress-partially accelerated life tests is discussed based on Type II censored sampling from Topp Leone-inverted Kumaraswamy distribution. The maximum likelihood estimators for the unknown parameters and the acceleration factor are obtained. Maximum likelihood prediction (point and bounds) is considered for future order statistics under Type-II censored informative samples obtained from constant stress-partially accelerated life testing models. Numerical study are given and some interesting comparisons are presented to illustrate the theoretical results. Moreover, the results are applied on real data set.

Keywords: *Topp Leone-inverted Kumaraswamy distribution; censored samples; asymptotic information matrix; two-sample prediction; maximum likelihood prediction.*

1. Introduction

Rapid developments, improvements of the high technology, consumer's demands for highly reliable products and competitive markets have placed pressure on manufacturers to deliver products with high quality and reliability. In life testing, it is very difficult to estimate the time of failure for modern high reliability products such as electronics, power cables, metal fatigue, insulating materials, laser, airplane parts, aerospace vehicles, etc.; since these types of products are not likely to fail under usual operating conditions in the relatively short time available for test. For this reason, *accelerated life testing* (ALT) or *partially accelerated life testing* (PALT) are preferred to be used in manufacturing industries to obtain enough failure data in a short period of time and necessary to study its relationship with external stress variables. Such testing could save much time, man power, material sources and money. The stress can be applied in different ways like constant stress, step stress and progressive stress among others.

In constant stress where under one higher than normal stress level, each specimen is run at a constant stress level. In practical use, most products run at constant stress as a constant stress test mimics actual use, it is simple and has a lot of advantages over time-dependent stress loadings because most of real products are operated at a constant-stress condition, see Nelson (1990). For more details about ALT, [see Bai and Chung (1989), Balakrishnan and Han (2008), AL-Dayian *et al.* (2014) and Basak and Balakrishnan (2018)] among others.

In ALT the main assumption is that a life-stress relationship is known or can be assumed so that the data obtained from accelerated conditions can be extrapolated to usual conditions.

In some cases, such relationship cannot be known or assumed. So, PALT are often used in such cases.

In a *constant stress-PALT* (CS-PALT), each test item is run at a constant stress under either normal usual condition or accelerated condition only until the test is terminated, and the analysis of PALT has been extensively studied in recent years, [see Bai *et al.* (1993), Ismail (2009), Hassan (2013), Hassan *et al.* (2015), Hyun and Lee (2015) and EL-Sagheer (2018)]. From the prediction viewpoint, few studies have considered PALT such as Abushal and AL-Zaydi (2017) and Prakash and Singh (2018).

Behairy *et al.* (2019) introduced *Topp Leone-inverted Kumaraswamy* (TL-IK) distribution as a composite distribution of $TL(\theta)$ and $IK(a, b)$ distributions. It's denoted by $TL-IK(a, b, \theta)$, its *cumulative distribution function* (cdf) and *probability density function* (pdf) are given, respectively, by:

$$F(x; a, b, \theta) = [\varphi(a)]^{\theta b} [2 - (\varphi(a))^b]^\theta, \quad 0 < x < \infty, a, b, \theta > 0, \quad (1)$$

where

$$\varphi(a) = 1 - (1 + x)^{-a}, \quad (2)$$

the pdf corresponding to (1) is given by

$$f(x; a, b, \theta) = 2ab\theta(1 + x)^{-(a+1)}[\varphi(a)]^{\theta b-1}[1 - (\varphi(a))^b][2 - (\varphi(a))^b]^{\theta-1}, \quad 0 < x < \infty, a, b, \theta > 0, \quad (3)$$

where a, b and θ are shape parameters.

The *reliability function* (rf) and *hazard rate function* (hrf) are given, respectively, by:

$$R(x_0; a, b, \theta) = P(X \geq x_0) = 1 - [\varphi(a)]^{\theta b} [2 - (\varphi(a))^b]^\theta, \quad 0 < x_0 < \infty, a, b, \theta > 0, \quad (4)$$

and

$$h(x_0; a, b, \theta) = \frac{f(x_0)}{1 - F(x_0)} = \frac{2ab\theta(1+x_0)^{-(a+1)}[\varphi(a)]^{\theta b-1}[1 - (\varphi(a))^b][2 - (\varphi(a))^b]^{\theta-1}}{1 - [\varphi(a)]^{\theta b} [2 - (\varphi(a))^b]^\theta}, \quad 0 < x_0 < \infty, a, b, \theta > 0. \quad (5)$$

TL-IK (a, b, θ) distribution contains some special well-known distributions in lifetime, such as the TL-Lomax (Pareto Type II), the TL-log-logistic (Fisk), the Lomax and the log-logistic (Fisk) distributions. Behairy *et al.* (2019) derived some transformed distributions such as the TL-exponentiated Weibull, TL-exponentiated Burr Type XII, TL-Kumaraswamy Dagum and TL-Kumaraswamy-inverse Weibull, among several others. They studied the properties of this distribution; they obtained the stress-strength reliability, moments, moment generating and quantile functions of the TL-IK distribution. Also, they

derived the *maximum likelihood* (ML) estimators, asymptotic variances and covariance matrix of the ML estimators and *confidence intervals* (CIs) for the parameters, also they obtained the ML two-sample predictors for the future observation based on Type-II censored data.

This paper is organized as follows: in Section 2, a description of the model is presented, the ML estimators for the parameters, the acceleration factor, rf and hrf for CS-PALT based on Type II censoring are derived. Point and bound prediction using the ML prediction for a future observation based on two-sample prediction are presented in Section 3. Numerical illustration is presented in Section 4. Finally, some general conclusions are introduced in Section 5.

2. Model Description and Maximum Likelihood Estimation

In this section, the description of the model and basic assumptions are presented in Subsection 2.1. Estimation for the unknown parameters, the acceleration factor, rf and hrf of the TL-IK distribution for CS-PALT under Type II censored data are discussed in Subsection 2.2.

2.1 Model description and assumptions

Total items are divided randomly into two samples of size $n(1 - \pi)$ and $n\pi$, respectively, where π is the sample proportion. The first sample is allocated to usual conditions and the other is assigned to accelerated conditions. Each test item of every sample is run without changing the test condition until reaching the censoring number.

- **Assumptions**

1. The lifetimes $X_i, i = 1, \dots, n(1 - \pi)$ of items allocated to usual conditions, are *independent and identically distributed* (iid) random variables.
2. The lifetimes $Y_j, j = 1, \dots, n\pi$ of items allocated to accelerated conditions, are iid random variables.
3. The lifetimes X_i and Y_j are mutually statistically independent.

In this study, the lifetimes of test items are assumed to have a TL-IK distribution. The pdf of an item at usual conditions is given by (3).

The pdf and cdf for an item tested at accelerated conditions are given by:

$$f(y; a, b, \theta, \beta) = 2ab\theta\beta(1 + \beta y)^{-(a+1)}[\varphi(a\beta)]^{\theta b - 1}[1 - (\varphi(a\beta))^b][2 - (\varphi(a\beta))^b]^{\theta - 1},$$

$$y > 0; a, b, \theta > 0; \beta > 1, \quad (6)$$

where

$$\varphi(a\beta) = 1 - (1 + \beta y)^{-a}, \quad (7)$$

and

$$F(y; a, b, \theta, \beta) = [\varphi(a\beta)]^{\theta b} [2 - (\varphi(a\beta))^b]^\theta, \quad y > 0; a, b, \theta > 0; \beta > 1. \quad (8)$$

where

$Y = \beta^{-1}X$, β is the acceleration factor which is the ratio of mean life at usual condition to that at accelerated condition and $\beta > 1$.

The rf and hrf for an item tested at accelerated conditions are as follows:

$$R(y_0; a, b, \theta, \beta) = 1 - [\varphi(a\beta)]^{\theta b} [2 - (\varphi(a\beta))^b]^\theta, \quad y_0 > 0; a, b, \theta > 0; \beta > 1, \quad (9)$$

and

$$h(y_0; a, b, \theta, \beta) = \frac{2ab\theta\beta(1+\beta y_0)^{-(a+1)}[\varphi(a\beta)]^{\theta b-1}[1-(\varphi(a\beta))^b][2-(\varphi(a\beta))^b]^{\theta-1}}{1-[\varphi(a\beta)]^{\theta b}[2-(\varphi(a\beta))^b]^\theta}, \quad y_0 > 0; a, b, \theta > 0; \beta > 1. \quad (10)$$

2.2 Maximum likelihood estimation

Considering CS-PALT based on Type-II censoring, the failure times consist of r th smallest lifetimes $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(r)}$ out of a random sample of $n(1 - \pi)$ lifetimes $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n(1-\pi))}$ under usual conditions and $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(r)}$ out of a random sample of $n\pi$ lifetimes $Y_{(1)} \leq Y_{(2)} \leq \dots \leq Y_{(n\pi)}$ at accelerated conditions, respectively.

The *likelihood function* (LF) for $\{(x_{(i)}): i = 1, \dots, n(1 - \pi)\}$ at usual conditions is given by

$$L_u(\underline{y}; \underline{x}) = \frac{(n-n\pi)!}{(n-n\pi-r)!} \left\{ \prod_{i=1}^r f(x_{(i)}; \underline{y}) \right\} [R(x_{(r)}; \underline{y})]^{n-n\pi-r}, \quad (11)$$

where $\underline{y} = (a, b, \theta)'$ and $f(x_{(i)}; \underline{y})$, $R(x_{(r)}; \underline{y})$ are given by (3) and (4), respectively.

The LF for $\{(y_{(j)}): j = 1, \dots, n\pi\}$ at accelerated conditions is given by

$$L_g(\underline{\Phi}; \underline{y}) = \frac{(n\pi)!}{(n\pi-r)!} \left\{ \prod_{j=1}^r f(y_{(j)}; \underline{\Phi}) \right\} [R(y_{(r)}; \underline{\Phi})]^{n\pi-r}, \quad (12)$$

where $\underline{\Phi} = (a, b, \theta, \beta)'$, $f(y_{(j)}; \underline{\Phi})$ and $R(y_{(r)}; \underline{\Phi})$ are given by (6) and (9), respectively.

Let c_u and c_g be the number of items censored at usual and accelerated conditions, respectively, where

$$c_u = n(1 - \pi) - r \text{ and } c_g = n\pi - r. \quad (13)$$

Substituting (3) and (4) into (11), and substituting (6) and (9) into (12), hence the LF according to CS-PALT, for $\{(x_{(i)}), (y_{(j)}): i = 1, \dots, n(1 - \pi), j = 1, \dots, n\pi\}$, can be written as:

$$\begin{aligned} L(\underline{\Phi}; \underline{x}, \underline{y}) &= \frac{(n-n\pi)!}{(c_u)!} (2ab\theta)^r \left\{ \prod_{i=1}^r (1 + x_{(i)})^{-(a+1)} (\varphi_i(a))^{\theta b-1} (1 - (\varphi_i(a))^b) \right. \\ &\quad \times (2 - (\varphi_i(a))^b)^{\theta-1} \left. \right\} [1 - (\varphi_r(a))^{\theta b} (2 - (\varphi_r(a))^b)^{\theta}]^{c_u} \\ &\quad \times \frac{(n\pi)!}{(c_g)!} (2ab\theta\beta)^r \left\{ \prod_{j=1}^r (1 + \beta y_{(j)})^{-(a+1)} (\varphi_j(a\beta))^{\theta b-1} (1 - (\varphi_j(a\beta))^b) \right. \\ &\quad \times (2 - (\varphi_j(a\beta))^b)^{\theta-1} \left. \right\} [1 - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^{\theta}]^{c_g}, \end{aligned} \quad (14)$$

where

$$\begin{cases} \underline{\Phi} = (a, b, \theta, \beta)', \varphi_i(a) = 1 - (1 + x_{(i)})^{-a}, \varphi_r(a) = 1 - (1 + x_{(r)})^{-a}, \\ \varphi_j(a\beta) = 1 - (1 + \beta y_{(j)})^{-a}, \varphi_r(a\beta) = 1 - (1 + \beta y_{(r)})^{-a}. \end{cases} \quad (15)$$

I. Point estimation

The ML estimators of a, b, θ and β can be derived by maximizing the logarithm of (14), denoted by ℓ which can be written in the form:

$$\begin{aligned} \ell \equiv \ln L_1(\underline{\Phi}; \underline{x}, \underline{y}) &\propto 2r \ln(a) + 2r \ln(b) + 2r \ln(\theta) + r \ln(\beta) - a \sum_{i=1}^r \ln(1 + x_{(i)}) \\ &\quad + (\theta b - 1) \sum_{i=1}^r \ln \varphi_i(a) + \sum_{i=1}^r \ln(1 - (\varphi_i(a))^b) + (\theta - 1) \sum_{i=1}^r \ln(2 - (\varphi_i(a))^b) \\ &\quad + c_u \ln(1 - (\varphi_r(a))^{\theta b} (2 - (\varphi_r(a))^b)^{\theta}) - a \sum_{j=1}^r \ln(1 + \beta y_{(j)}) \\ &\quad + (\theta b - 1) \sum_{j=1}^r \ln \varphi_j(a\beta) + \sum_{j=1}^r \ln(1 - (\varphi_j(a\beta))^b) \\ &\quad + (\theta - 1) \sum_{j=1}^r \ln(2 - (\varphi_j(a\beta))^b) + c_g \ln(1 - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^{\theta}). \end{aligned} \quad (16)$$

The first partial derivatives of the logarithm of the LF with respect to a, b, θ and β are given below:

$$\begin{aligned} \frac{\partial \ell}{\partial a} &= \frac{2r}{a} - \sum_{i=1}^r \ln(1 + x_{(i)}) + (\theta b - 1) \sum_{i=1}^r \frac{(1+x_{(i)})^{-a} \ln(1+x_{(i)})}{\varphi_i(a)} \\ &\quad - \sum_{i=1}^r \frac{b(\varphi_i(a))^{b-1} (1+x_{(i)})^{-a} \ln(1+x_{(i)})}{(1-(\varphi_i(a))^b)} - (\theta - 1) \sum_{i=1}^r \frac{b(\varphi_i(a))^{b-1} (1+x_{(i)})^{-a} \ln(1+x_{(i)})}{(2-(\varphi_i(a))^b)} \\ &\quad - c_u \left[\frac{\theta b (1+x_{(r)})^{-a} \ln(1+x_{(r)}) \{ (2-(\varphi_r(a))^b)^{\theta} (\varphi_r(a))^{\theta b-1} - (\varphi_r(a))^{b(\theta+1)-1} (2-(\varphi_r(a))^b)^{\theta-1} \}}{1-(\varphi_r(a))^{\theta b} (2-(\varphi_r(a))^b)^{\theta}} \right] \end{aligned}$$

$$\begin{aligned}
 & - \sum_{j=1}^r \ln(1 + \beta y_{(j)}) + (\theta b - 1) \sum_{j=1}^r \frac{(1 + \beta y_{(j)})^{-a} \ln(1 + \beta y_{(j)})}{\varphi_j(a\beta)} \\
 & - \sum_{j=1}^r \frac{b(\varphi_j(a\beta))^{b-1} (1 + \beta y_{(j)})^{-a} \ln(1 + \beta y_{(j)})}{(1 - (\varphi_j(a\beta))^b)} - (\theta - 1) \sum_{j=1}^r \frac{b(\varphi_j(a\beta))^{b-1} (1 + \beta y_{(j)})^{-a} \ln(1 + \beta y_{(j)})}{(2 - (\varphi_j(a\beta))^b)} \\
 & - c_g \left[\frac{\theta b (1 + \beta y_{(r)})^{-a} \ln(1 + \beta y_{(r)}) \{ (2 - (\varphi_r(a\beta))^b)^\theta (\varphi_r(a\beta))^{\theta b - 1} - (\varphi_r(a\beta))^{b(\theta + 1) - 1} (2 - (\varphi_r(a\beta))^b)^{\theta - 1} \}}{1 - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^\theta} \right], \quad (17)
 \end{aligned}$$

$$\begin{aligned}
 \frac{\partial \ell}{\partial b} &= \frac{2r}{b} + \theta \sum_{i=1}^r \ln \varphi_i(a) - \sum_{i=1}^r \frac{(\varphi_i(a))^b \ln(\varphi_i(a))}{(1 - (\varphi_i(a))^b)} - (\theta - 1) \sum_{i=1}^r \frac{(\varphi_i(a))^b \ln(\varphi_i(a))}{(2 - (\varphi_i(a))^b)} \\
 & - c_u \left[\frac{(2 - (\varphi_r(a))^b)^\theta \theta (\varphi_r(a))^{\theta b} \ln(\varphi_r(a)) - (\varphi_r(a))^{b(\theta + 1)} \theta (2 - (\varphi_r(a))^b)^{\theta - 1} \ln(\varphi_r(a))}{1 - (\varphi_r(a))^{\theta b} (2 - (\varphi_r(a))^b)^\theta} \right] \\
 & + \theta \sum_{j=1}^r \ln \varphi_j(a\beta) - \sum_{j=1}^r \frac{(\varphi_j(a\beta))^b \ln(\varphi_j(a\beta))}{(1 - (\varphi_j(a\beta))^b)} - (\theta - 1) \sum_{j=1}^r \frac{(\varphi_j(a\beta))^b \ln(\varphi_j(a\beta))}{(2 - (\varphi_j(a\beta))^b)} \\
 & - c_g \left[\frac{(2 - (\varphi_r(a\beta))^b)^\theta \theta (\varphi_r(a\beta))^{\theta b} \ln(\varphi_r(a\beta)) - (\varphi_r(a\beta))^{b(\theta + 1)} \theta (2 - (\varphi_r(a\beta))^b)^{\theta - 1} \ln(\varphi_r(a\beta))}{1 - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^\theta} \right], \quad (18)
 \end{aligned}$$

$$\begin{aligned}
 \frac{\partial \ell}{\partial \theta} &= \frac{2r}{\theta} + b \sum_{i=1}^r \ln \varphi_i(a) + \sum_{i=1}^r (2 - (\varphi_i(a))^b) \\
 & - c_u \left[\frac{(2 - (\varphi_r(a))^b)^\theta b (\varphi_r(a))^{\theta b} \ln(\varphi_r(a)) - (\varphi_r(a))^{\theta b} (2 - (\varphi_r(a))^b)^\theta \ln(2 - (\varphi_r(a))^b)}{1 - (\varphi_r(a))^{\theta b} (2 - (\varphi_r(a))^b)^\theta} \right] \\
 & + b \sum_{j=1}^r \ln \varphi_j(a\beta) + \sum_{j=1}^r (2 - (\varphi_j(a\beta))^b) \\
 & - c_g \left[\frac{(2 - (\varphi_r(a\beta))^b)^\theta b (\varphi_r(a\beta))^{\theta b} \ln(\varphi_r(a\beta)) - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^\theta \ln(2 - (\varphi_r(a\beta))^b)}{1 - (\varphi_r(a\beta))^{\theta b} (2 - (\varphi_r(a\beta))^b)^\theta} \right], \quad (19)
 \end{aligned}$$

and the first partial derivative of the acceleration factor is

$$\begin{aligned}
 \frac{\partial \ell}{\partial \beta} &= \frac{r}{\beta} - a \sum_{j=1}^r \frac{y_{(j)}}{1 + \beta y_{(j)}} + (\theta b - 1) \sum_{j=1}^r \frac{a y_{(j)} (1 + \beta y_{(j)})^{-(a+1)}}{\varphi_j(a\beta)} \\
 & - \sum_{j=1}^r \frac{a b y_{(j)} (\varphi_j(a\beta))^{b-1} (1 + \beta y_{(j)})^{-(a+1)}}{(1 - (\varphi_j(a\beta))^b)} - (\theta - 1) \sum_{j=1}^r \frac{a b y_{(j)} (\varphi_j(a\beta))^{b-1} (1 + \beta y_{(j)})^{-(a+1)}}{(2 - (\varphi_j(a\beta))^b)}
 \end{aligned}$$

$$-c_g \left[\frac{ab\theta y_{(r)}(1+\beta y_{(r)})^{-(a+1)} \{(\varphi_r(a\beta))^{b(\theta+1)-1} (2-(\varphi_r(a\beta))^b)^{\theta-1} - (2-(\varphi_r(a\beta))^b)^\theta (\varphi_r(a\beta))^{\theta b-1}\}}{1-(\varphi_r(a\beta))^{\theta b} (2-(\varphi_r(a\beta))^b)^\theta} \right]. \quad (20)$$

The ML estimators are obtained by setting Equations (17)-(20) to zeros. The system of non-linear equations can be solved numerically using Newton-Raphson method, to obtain the ML estimates of $\hat{a}$, $\hat{b}$, $\hat{\theta}$ and $\hat{\beta}$.

Depending on the invariance property of the ML estimators, then, the ML estimators of the rf and hrf are obtained by replacing the parameters a, b, θ and β in (9) and (10) respectively by their ML estimates.

Hence, for a given value of y , the ML estimators of $R(y)$ and $h(y)$ are as follows:

$$\hat{R}(y) = 1 - [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}} \left[2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}} \right]^{\hat{\theta}}, \quad y > 0, \quad (21)$$

and

$$\hat{h}(y) = \frac{2\hat{a}\hat{b}\hat{\theta}\hat{\beta}(1+\hat{\beta}y)^{-(\hat{a}+1)} [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}-1} \left[1 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}} \right] \left[2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}} \right]^{\hat{\theta}-1}}{1 - [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}} \left[2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}} \right]^{\hat{\theta}}}, \quad y > 0. \quad (22)$$

II. Confidence intervals

The *asymptotic variance covariance matrix* (AVCM) of the estimators $\hat{a}$, $\hat{b}$, $\hat{\theta}$ and $\hat{\beta}$ are obtained depending on the *inverse asymptotic Fisher information matrix* (AFIM) using the partial second derivatives of the logarithm of the LF.

The AFIM can be written as follows:

$$\tilde{I} = - \left[\frac{\partial^2 l}{\partial \psi_i \partial \psi_j} \right], \quad \text{where } i, j = 1, 2, 3, 4, \quad (23)$$

where $\psi_1 = a, \psi_2 = b, \psi_3 = \theta$ and $\psi_4 = \beta$.

For large sample size, the ML estimators under regularity conditions are consistent and asymptotically unbiased as well as asymptotically normally distributed. Therefore, the *asymptotic confidence interval* (ACI) for the parameters; ψ , can be obtained by

$$P \left[-Z < \frac{\hat{\psi}_{iML} - \psi_i}{\sigma_{\hat{\psi}_{iML}}} < Z \right] = 1 - \tau, \text{ where } Z \text{ is the } 100 \left(1 - \frac{\tau}{2} \right) \text{th standard normal percentile.}$$

The two-sided approximate $100(1 - \tau)\%$ the confidence intervals are

$$L_{\psi_i} = \hat{\psi}_{iML} - Z_{\frac{\tau}{2}} \hat{\sigma}_{\hat{\psi}_{iML}}, \text{ and } U_{\psi_i} = \hat{\psi}_{iML} + Z_{\frac{\tau}{2}} \hat{\sigma}_{\hat{\psi}_{iML}}, \quad (24)$$

where $\hat{\sigma}_{\hat{\psi}_{iML}}$ is the standard deviation and $\hat{\psi}_{iML}$ is $\hat{a}$, $\hat{b}$, $\hat{\theta}$, $\hat{\beta}$, $\hat{R}(y)$ or $\hat{h}(y)$, respectively.

3. Maximum Likelihood Prediction Based on Two-Sample Prediction

In this section two-sample prediction is considered. ML point and interval prediction for a future observation from the CS-PALT based on Type II censoring for TL-IK distribution are discussed.

Assuming that $Y_{(1)} \leq Y_{(2)} \leq \dots \leq Y_{(r)}$ are the first r ordered life times in a random sample of n components (Type II censoring) whose failure times are identically distributed as a random variable Y having the pdf for an item tested at accelerated conditions is given by (6) which is an informative sample, and that $T_{(1)}, T_{(2)}, \dots, T_{(m)}$ is a future independent random sample (of size m) from the same distribution. Our aim is to predict a statistic in the future sample based on the informative sample [see Kaminsky and Rhodin (1985), Ateya and Mohammed (2018) and Raqab *et al.* (2019)].

For the future sample of size m , let $T_{(s)}$ denotes the s^{th} order statistic, $1 \leq s \leq m$, then he pdf of $T_{(s)}$ is given by

$$f_{s:m}(t_{(s)}) = D(s)f(t_{(s)}|\underline{\Phi})[F(t_{(s)}|\underline{\Phi})]^{s-1}[1 - F(t_{(s)}|\underline{\Phi})]^{m-s}, \quad t_{(s)} > 0,$$

$$\underline{\Phi} = (a, b, \theta, \beta)', D(s) = \frac{1}{B(s, m-s+1)}, s = 1, 2, 3, \dots, m, \quad (25)$$

using the binomial expansion theorem for $[1 - F(t_{(s)}|\underline{\Phi})]^{m-s}$, yields

$$f_{s:m}(t_{(s)}) = D(s)f(t_{(s)}) \sum_{l_7=0}^{m-s} (-1)^{l_7} \binom{m-s}{l_7} [F(t_{(s)})]^{s+l_7-1}, \quad (26)$$

and substituting (6) and (8) into (26), then one can obtain the pdf of s^{th} order statistic for an item tested at accelerated conditions:

$$h(t_{(s)}; \underline{\Phi}) = 2ab\theta\beta D(s) \sum_{l_7=0}^{m-s} \delta (1 + \beta t_{(s)})^{-(a+1)} [\varphi(a\beta)]^{\theta b(s+l_7)-1} [1 - (\varphi(a\beta))^b] \\ \times [2 - (\varphi(a\beta))^b]^{\theta(s+l_7)-1}, t_{(s)} > 0; \beta > 1; (\underline{\vartheta} > \underline{0}), \quad (27)$$

where

$$\sum_{l_7=0}^{m-s} \delta = \sum_{l_7=0}^{m-s} (-1)^{l_7} \binom{m-s}{l_7}, \underline{\vartheta} = (a, b, \theta)' \text{ and } \varphi_t(a\beta) = 1 - (1 + \beta t_{(s)})^{-a}. \quad (28)$$

Assuming that the parameters $\underline{\Phi}$ are unknown and independent, then the ML prediction density (MLPD) of $T_{(s)}$ given $\hat{\underline{\Phi}}_{ML}$ can be obtained using the conditional pdf of the s^{th} order statistic which is given by (27) after replacing the vector of parameters $\underline{\Phi}$ by their ML estimates $\hat{\underline{\Phi}}_{ML}$, as follows:

$$h(t_{(s)}; \hat{\Phi}_{ML}) = 2\hat{a}\hat{b}\hat{\theta}\hat{\beta}D(s) \sum_{l_7=0}^{m-s} \delta(1 + \hat{\beta}t_{(s)})^{-(\hat{a}+1)} [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}(s+l_7)-1} [1 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}}] \times [2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}}]^{\hat{\theta}(s+l_7)-1}, t_{(s)} > 0; \hat{\beta} > 1; (\hat{\vartheta}_{ML} > 0), \quad (29)$$

where

$$\hat{\Phi}_{ML} = (\hat{a}, \hat{b}, \hat{\theta}, \hat{\beta})', \hat{\vartheta} = (\hat{a}, \hat{b}, \hat{\theta})', \text{ and } \varphi(\hat{a}\hat{\beta}) = 1 - (1 + \hat{\beta}t_{(s)})^{-\hat{a}}, \quad (30)$$

where $D(s)$ and $\sum_{l_7=0}^{m-s} \delta$ are given by (25) and (28) respectively.

I. Point prediction

The *ML predictor* (MLP) for the future observation $T_{(s)}$, based on Type II censoring can be derived using (29) as follows:

$$\begin{aligned} \hat{t}_{(s)(ML)} &= E(t_{(s)} | \hat{\Phi}_{ML}) = \int_{t_{(s)}} t_{(s)} h(t_{(s)}; \hat{\Phi}_{ML}) dt_{(s)} \\ &= 2D(s) \sum_{l_7=0}^{m-s} \delta \int_{t_{(s)}} \hat{a}\hat{b}\hat{\theta}\hat{\beta} t_{(s)} (1 + \hat{\beta}t_{(s)})^{-(\hat{a}+1)} [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}(s+l_7)-1} \\ &\quad \times [1 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}}] [2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}}]^{\hat{\theta}(s+l_7)-1} dt_{(s)}. \end{aligned} \quad (31)$$

II. Interval prediction

A $100(1 - \tau)\%$ *ML predictive bound* (MLPB) for the future observation $T_{(s)}$, such that $P(L_{(s)}(\underline{y}) < T_{(s)} < U_{(s)}(\underline{y}) | \underline{y}) = 1 - \tau$, are given by:

$$P(T_{(s)} > L_{(s)}(\underline{y}) | \underline{y}) = \int_{L_{(s)}(\underline{y})}^{\infty} h(t_{(s)}; \hat{\Phi}_{ML}) dt_{(s)} = 1 - \frac{\tau}{2}, \quad (32)$$

and

$$P(T_{(s)} > U_{(s)}(\underline{y}) | \underline{y}) = \int_{U_{(s)}(\underline{y})}^{\infty} h(t_{(s)}; \hat{\Phi}_{ML}) dt_{(s)} = \frac{\tau}{2}. \quad (33)$$

Substituting (29) in (32) and (33), the MLPB are obtained as follows:

$$\begin{aligned} P(T_{(s)} > L_{(s)}(\underline{y}) | \underline{y}) &= 2D(s) \sum_{l_7=0}^{m-s} \delta \int_{L_{(s)}(\underline{y})}^{\infty} \hat{a}\hat{b}\hat{\theta}\hat{\beta} (1 + \hat{\beta}t_{(s)})^{-(\hat{a}+1)} \\ &\quad \times [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}(s+l_7)-1} [2 - (\varphi(\hat{a}\hat{\beta}))^{\hat{b}}]^{\hat{\theta}(s+l_7)-1} dt_{(s)} \end{aligned}$$

$$\times \left[1 - \left(\varphi(\hat{a}\hat{\beta}) \right)^{\hat{b}} \right] dt_{(s)} = 1 - \frac{\tau}{2}, \quad (34)$$

and

$$\begin{aligned} P(T_{(s)} > U_{(s)}(\underline{y}) | \underline{y}) &= 2D(s) \sum_{l_7=0}^{m-s} \delta \int_{U_{(s)}(\underline{y})}^{\infty} \hat{a}\hat{b}\hat{\theta}\hat{\beta}(1 + \hat{\beta}t_{(s)})^{-(\hat{a}+1)} \\ &\times [\varphi(\hat{a}\hat{\beta})]^{\hat{\theta}\hat{b}(s+l_7)-1} \left[2 - \left(\varphi(\hat{a}\hat{\beta}) \right)^{\hat{b}} \right]^{\hat{\theta}(s+l_7)-1} \\ &\times \left[1 - \left(\varphi(\hat{a}\hat{\beta}) \right)^{\hat{b}} \right] dt_{(s)} = \frac{\tau}{2}. \end{aligned} \quad (35)$$

Remark:

If $s = 1$, $s = m$ and $s = \frac{m+1}{2}$ in (31), one can predict the minimum observable, $T_{(1)}$, the maximum observable, $T_{(m)}$, and the median observable when m is odd, $T_{(\frac{m+1}{2})}$.

4. Numerical Illustration

This section aims to investigate the precision of the theoretical results of estimation and prediction on the basis of simulated and real data.

4.1 Simulation algorithm

In this subsection, a simulation study is conducted to illustrate the performance of the presented ML estimates on the basis of generated data from the TL-IK (a, b, θ) distribution considering the CS-PALT. The ML averages of the parameters, rf and hrf based on Type II censoring are computed. Moreover, CIs of the parameters, rf and hrf are calculated. Also, ML two-sample predictors (point and interval) are calculated. Simulation studies are performed using Mathematica 9 for illustrating the obtained results.

The steps of the simulation procedure based on Type II censored data are as follows:

- a) For given values of $\underline{\vartheta}_i$, random samples of size n are generated from the TL-IK (a, b, θ) distribution.
- b) For given values of $\underline{\vartheta}_j$, random samples of size n are generated from the TL-IK (a, b, θ) distribution.
- The transformation between the uniform distribution and TL-IK distribution is obtained as follows:

$$x = \left[\left(1 - (u)^{\frac{1}{b}} \right)^{-\frac{1}{a}} - 1 \right] \left(1 - \sqrt{1 - (u)^{\frac{1}{\theta}}} \right), \quad 0 < u < 1.$$

- c) For each sample size n sort x_i 's, such that $x_1 \leq x_2 \leq \dots \leq x_{n(1-\pi)}$.

- d) For each sample size n sort y_j 's, such that $y_1 \leq y_2 \leq \dots \leq y_{n\pi}$.
 - e) Choose the number of failures r to be less than $n(1 - \pi)$ and $n\pi$.
 - f) Repeat all the previous steps N times, where N represents a fixed number of simulated samples, where $N = 3000$ is the number of repetitions.
- For the number of the units r and the population parameter values and using the Newton-Raphson method, the ML averages and the CIs of the parameters are obtained. Also, the rf, hrf and their CIs are calculated using the ML averages of the parameters.
 - Evaluating the performance of the estimates is considered through some measurements of accuracy. In order to study the precision and variation of the estimates, it is convenient to use the *estimated risk* (ER) = $\frac{\sum_{i=1}^N (\text{estimate} - \text{true value})^2}{N}$.
 - The ML predictors (point and interval) for a future observation from the TL-IK distribution based on Type II censored data are computed for the two-sample case.
 - Simulation results of the ML estimates are displayed in Tables 1 and 2, where the samples of size ($n=30, 60, 100$), are used. For each sample size, the censoring level is 10% and the chosen parameters value are selected as (Case 1, $a = 1.1, b = 1.2, \theta = 2.1, \beta = 1.2, y_0 = 0.1, \pi = 0.20$ and $\pi = 0.30$) and (Case 2, $a = 1.1, b = 1.2, \theta = 2.1, \beta = 1.5, y_0 = 0.1, \pi = 0.30$).
 - Table 1 and 2 display the ML averages, ERs and CIs of the unknown parameters, rf and hrf based on Type II censoring in Case 1 and Case 2, respectively. While Table 3 presents the ML averages, of the rf and hrf for different values of time x_0 based on Type II censoring.
 - The ML two-sample predictors are presented in Table 5.

4.2 Some applications

This subsection aims to demonstrate how the proposed method can be used in practice. Two real lifetime data sets are used for this purpose. The TL-IK (a, b, θ) distribution is fitted to two real data using Kolmogorov-Smirnov goodness of fit test through the R programming language.

Application 1

The data given by Murthy *et al.* (2004), refers to the time between failures for a repairable item: 1.43, 0.11, 0.71, 0.77, 2.63, 1.49, 3.46, 2.46, 0.59, 0.74, 1.23, 0.94, 4.36, 0.40, 1.74, 4.73, 2.23, 0.45, 0.70, 1.06, 1.46, 0.30, 1.82, 2.37, 0.63, 1.23, 1.24, 1.97, 1.86 and 1.17.

Application 2

The second application is given by Dumonceaux and Antle (1973). With respect to the maximum flood level (in millions of cubic feet per second) for the Susquehanna River at Harrisburg, Pennsylvania. Each number is the maximum flood level for a four year period, the first, 0.654, being for the period 1890-1893, and the last, 0.265, being for the period 1966-1969. The data are

0.654, 0.613, 0.315, 0.449, 0.297, 0.402, 0.379, 0.423, 0.379, 0.3235 0.269, 0.740, 0.418, 0.412 0.494, 0.416, 0.338, 0.392, 0.484, 0.265.

The Kolmogorov–Smirnov goodness of fit test is applied to check the validity of the fitted model. The p values are given, respectively, 0.5860 and 0.6465. The p value given in each case showed that the model fits the data very well.

- Table 4 gives the ML averages and standard error (SE) of the unknown parameters, rf and hrf for the real data sets based on Type II censoring. Table 6 presents the ML two-sample predictors for the future observation based on Type II censoring of the real data sets.

4.3 Concluding remarks

- It is noticed, from Tables 1 and 2 that the ML averages are very close to the population parameter values as the sample size increases. Also, ERs are decreasing when the sample size is increasing. This is indicative of the fact that the estimates are consistent and approaches the true parameter values as the sample size increases.
- The lengths of the CIs of the parameters become narrower as the sample size increases.
- From Table 3 the estimated values of the rf decreases when the time y_0 increases, while the estimated values of the hrf increases when the time y_0 increases then decreases. The results get better when the sample size increases.
- The length of the first future order statistic is smaller than the length of the last future order statistic. [Tables 5 and 6].
- The ML intervals include the predictive values (between the *lower limit* (LL) and *upper limit* (UL)).

5. General Conclusion

For products having a high reliability, the test of product life under usual conditions often requires a long period of time. So ALT or PALT is used to facilitate estimating the reliability of the unit in a short period of time. In ALT test items are run only at accelerated conditions, in some cases, such relationship cannot be known or assumed. So, PALT are often used in such cases, in PALT they are run at both usual and accelerated conditions. This paper deals with the CS-PALT in the case of Type II censoring. It is assumed that the lifetime of test

units has the TL-IK distribution. ML estimators of the acceleration factor and the parameters are derived. Point and bounds prediction using the ML prediction method for a future observation based on two-sample prediction are studied. From the results, one can clearly see that the acceleration factor decreases the estimates have smaller ER. As the sample size increases the ER and the interval of the parameters, rf and hrf decrease. This indicates that the ML estimates provide asymptotically normally distributed and consistent estimators for the parameters and the acceleration factor.

Acknowledgements:

The authors would like to thank the Referees and the Editor for their constructive comments which led to improvements of an earlier version of this article.

6. References

- Abushal, T., A. and AL-Zaydi, A., M. (2017).** Inference on constant-partially accelerated life tests for mixture of Pareto distributions under progressive Type-II censoring. *Open Journal of Statistics*, Vol. 7, pp. 323-346.
- Ateya, S., F. and Mohammed, H., S. (2018).** Prediction under Burr-XII distribution based on generalized Type-II progressive hybrid censoring. *Journal of the Egyptian Mathematical Society*, Vol. 26, pp. 491-508.
- AL-Dayian G. R., EL-Helbawy A. A. and Rezk H. R. (2014).** Statistical inference for a simple constant stress model based on censored sampling data from the Kumaraswamy Weibull distribution. *International Journal of Statistics and Probability*, Vol. 3, No. 3, pp. 80-95.
- Bai, D. S., and Chung, S. W. (1989).** Accelerated life test model with the inverse power law. *Reliability Engineering and System Safety*, Vol. 24, pp. 223-230.
- Bai, D., S., Chung, S., W. and Chun Y., R. (1993).** Optimal design of partially accelerated life tests for lognormal distribution under Type-I censoring. *Reliability Engineering and System Safety*, Vol. 40, pp. 85 – 92.
- Balakrishnan, N. and Han, D. (2008).** Exact inference for a simple step-stress model with competing risks for failure from exponential distribution under Type II censoring. *Journal of Statistical Planning and Inference*, Vol.138, pp. 4172-4186.
- Basak, I. and Balakrishnan, N. (2018).** A note on the prediction of censored exponential lifetimes in a simple step-stress model with Type II censoring. *Calcutta Statistical Association Bulletin*, Vol. 70, pp. 57-73.
- Behariy, S. M., EL-Helbawy, A. A., Refaey, R. M. and AL-Dayian, G. R. (2019).** *Topp Leone-inverted Kumaraswamy distribution: properties, estimation and prediction*. This paper proceedings of the reliability session of 31th Annual International Conference on Statistics and Modeling in Human and Social Sciences, Faculty of Economics and Political Science, Cairo, Egypt.
- Dumonceaux, R. and Antle, C. E. (1973).** Discrimination between the log-normal and the Weibull distributions. *Technometrics*, Vol. 15, No. 4, pp. 923-926.

- EL-Sagheer, R. M. (2018).** Inferences in constant-partially accelerated life tests based on progressive Type-II censoring. *Bulletin of the Malaysian Mathematical Sciences Society*, Vol. 41, pp. 609–626.
- Hassan, A. S. (2013).** On the optimal design of failure step stress-partially accelerated life tests for exponentiated inverted Weibull with censoring. *Australian Journal of Basic and Applied Sciences*, Vol. 1, pp. 97-104.
- Hassan, A. S., Assar, S. M. and Zaky, A. N. (2015).** Constant-stress partially accelerated life tests for inverted Weibull distribution with multiple censored data. *International Journal of Advanced Statistics and Probability*, Vol. 3, pp. 72-82.
- Hyun, S. and Lee, J. (2015).** Constant stress-partially accelerated life testing for log-logistic distribution with censored data. *Journal of Statistics Applications and Probability*, Vol. 4, pp. 193-201.
- Ismail, A. (2009).** Optimum constant stress-partially accelerated life test plans with Type-II censoring: the case of Weibull failure distribution. *International Journal of Statistics and Economics*, Vol. 3, No. S09, pp. 39-46.
- Kaminsky, K. S. and Rhodin, L. S. (1985).** Maximum likelihood prediction. *Annals of the Institute of Statistical Mathematics*, Vol. 37, pp. 507-517.
- Murthy, D. N. P., Xie, M. and Jiang, R. (2004).** *Weibull Models*. John Wiley and Sons, Inc., Hoboken, New Jersey.
- Nelson, W. (1990).** *Accelerated testing: Statistical Models, Test Plans and Data Analysis*. John Wiley, New York.
- Prakash, G. and Singh, P. (2018).** Derived the bound lengths based on constant stress-partially accelerated life test under different censoring patterns. *International Journal of Scientific World*, Vol. 6, pp. 19-26.
- Raqab, M. Z., Alkhalfan, L. A., Bdair, O. M. and Balakrishnan, N. (2019).** Maximum likelihood prediction of records from 3-parameter Weibull distribution and some approximations. *Journal of Computational and Applied Mathematics*, Vol. 356, pp.118-132.

Appendix

Table 1: ML averages, estimated risks and 95% CIs for the parameters a, b, θ, β , rf and hrf based on Type II censoring, ($N=3000, \pi = 0.20, \pi = 0.30$ and $r=0.1n$)
(Case 1, $a = 1.1, b = 1.2, \theta = 2.1, \beta = 1.2, y_0 = 0.1$)

n	r	π	Parameters, rf and hrf	average	ER	UL	LL	length
30	3	0.20	a	0.95499	0.06071	1.00615	0.90383	0.10231
			b	1.07131	0.05263	1.10699	1.03562	0.07137
			θ	1.85000	0.30411	1.92863	1.77137	0.15726
			β	1.05474	0.00339	1.12652	0.98296	0.14356
			$R(y_0)$	0.97041	0.00065	0.99027	0.95056	0.03971
			$h(y_0)$	0.52186	0.13771	0.73632	0.30738	0.42894
30	3	0.30	a	0.95351	0.06169	1.01358	0.89344	0.12013
			b	1.07091	0.05279	1.10523	1.03660	0.06863
			θ	1.84956	0.30422	1.91860	1.78052	0.13809
			β	1.05031	0.00335	1.10838	0.99225	0.11613
			$R(y_0)$	0.97079	0.00059	0.98542	0.95616	0.02926
			$h(y_0)$	0.51654	0.12832	0.67245	0.36062	0.31184
60	6	0.20	a	0.95441	0.06037	0.96810	0.94070	0.02739
			b	1.07182	0.05213	1.08750	1.05613	0.03137
			θ	1.85197	0.30065	1.88653	1.81741	0.06912
			β	1.05119	0.00255	1.07671	1.02567	0.05104
			$R(y_0)$	0.97113	0.00053	0.97885	0.96340	0.01545
			$h(y_0)$	0.51327	0.12099	0.58306	0.44348	0.13958
60	6	0.30	a	0.95323	0.06094	0.96673	0.93974	0.02698
			b	1.07120	0.05238	1.08266	1.05975	0.02291
			θ	1.85096	0.30158	1.87448	1.82745	0.04704
			β	1.04782	0.00283	1.06782	1.02781	0.04001
			$R(y_0)$	0.97126	0.00051	0.97279	0.96973	0.00306
			$h(y_0)$	0.51099	0.11824	0.52889	0.49310	0.03579
100	10	0.20	a	0.95461	0.06024	0.96407	0.94515	0.18924
			b	1.07175	0.05211	1.08009	1.06342	0.01666
			θ	1.85206	0.30031	1.86894	1.83518	0.03377
			β	1.05072	0.00250	1.06769	1.03374	0.03396
			$R(y_0)$	0.97119	0.00051	0.97202	0.97035	0.00166
			$h(y_0)$	0.51284	0.11945	0.52170	0.50399	0.01772
100	10	0.30	a	0.95342	0.60826	0.96304	0.94379	0.01925
			b	1.07123	0.05235	1.07916	1.06329	0.01588
			θ	1.85109	0.30137	1.86746	1.83472	0.03273
			β	1.04757	0.00280	1.06158	1.03356	0.02802
			$R(y_0)$	0.97128	0.00051	0.97225	0.97031	0.00194
			$h(y_0)$	0.51081	0.11807	0.52242	0.49921	0.02321

Table 2: ML averages, estimated risks and 95% CIs of the parameters a, b, θ, β , rf and hrf based on Type II censoring, (N=3000, $\pi = 0.30$ and $r=0.1n$) (Case 2, $a = 1.1$, $b = 1.2$, $\theta = 2.1$, $\beta = 1.5$, $y_0 = 0.1$)

n	r	π	Parameters, rf and hrf	Average	ER	UL	LL	length
30	3	0.30	a	0.95445	0.06203	1.03598	0.87292	0.16307
			b	1.07104	0.05303	1.11918	1.02290	0.09628
			θ	1.84897	0.30584	1.94112	1.75682	0.18431
			β	1.31230	0.00903	1.38406	1.24054	0.14351
			$R(y_0)$	0.95721	0.00109	0.97843	0.93599	0.04243
			$h(y_0)$	0.74105	0.20179	0.95533	0.52677	0.42857
60	6	0.30	a	0.95299	0.06106	0.96711	0.93889	0.02821
			b	1.07127	0.05239	1.08722	1.05531	0.03190
			θ	1.85077	0.30198	1.88587	1.81567	0.07019
			β	1.30879	0.00848	1.33380	1.28377	0.05003
			$R(y_0)$	0.95804	0.00094	0.96625	0.94982	0.01644
			$h(y_0)$	0.73131	0.18261	0.79822	0.66439	0.13382
100	10	0.30	a	0.95305	0.06101	0.96265	0.94345	0.01920
			b	1.07132	0.05231	1.07970	1.06295	0.01675
			θ	1.85103	0.30144	1.86818	1.83389	0.03428
			β	1.30840	0.00847	1.32590	1.29091	0.03499
			$R(y_0)$	0.95813	0.00091	0.95942	0.95683	0.00259
			$h(y_0)$	0.73045	0.18077	0.74522	0.715672	0.02955

**Table 3: ML average of rf and hrf under different time y_0
at different sample sizes based on Type II censoring
($N=3000$, $\pi = 0.30$ and $r=0.1n$)**

n	y_0	Average	
		$\hat{R}(y_0)$	$\hat{h}(y_0)$
30	0.1	0.97079	0.51654
	0.3	0.83930	0.85348
	0.6	0.64195	0.90342
	0.9	0.49322	0.84503
	1.0	0.45394	0.81859
60	0.1	0.97126	0.51099
	0.3	0.84094	0.84742
	0.6	0.64328	0.90059
	0.9	0.49486	0.84093
	1.0	0.45569	0.81569
100	0.1	0.97128	0.51081
	0.3	0.84090	0.84749
	0.6	0.64318	0.90070
	0.9	0.49495	0.84079
	1.0	0.45557	0.81594

**Table 4: ML estimates for the parameters, rf and hrf and their standard
error for the real data sets based on Type II censoring
($r=0.1n$, $\pi = 0.30$, $y_0 = 0.1$)**

	n	r	Parameters	Estimates	SE
Application I	30	3	a	1.27325	0.00537
			b	1.22876	0.00507
			θ	2.46650	0.00442
			β	1.55154	0.20389
			$R(y_0)$	0.97855	0.00023
			$h(y_0)$	0.53019	0.13172
Application II	20	2	a	0.92518	0.07553
			b	1.42167	0.01480
			θ	1.3972	7.8186×10^{-6}
			β	2.08405	0.96836
			$R(y_0)$	0.93389	0.00022
			$h(y_0)$	1.12692	0.07592

Table 5: ML predictive and bounds of the future observation based on Type II censoring under two-sample prediction (N=3000, n=30, r=0.1n, $a = 1.1$, $b = 1.2$, $\theta = 2.1$, $\beta = 1.2$, $y_0 = 0.1$)

$\pi = 0.20$					$\pi = 0.30$				
s	$\hat{t}_{(s)}(ML)$	UL	LL	length	s	$\hat{t}_{(s)}(ML)$	UL	LL	length
1	0.09659	0.01483	0.23585	0.22103	1	0.09719	0.01491	0.23735	0.22244
15	0.86561	0.52170	1.35255	0.83085	15	0.87158	0.52505	1.36251	0.83751
25	2.44757	1.33847	4.34017	3.0017	25	2.46488	1.34756	4.37215	3.02459

Table 6: ML predictive and bounds of the future observation for real data sets based on Type II censoring under two-sample prediction (r=0.1n, $\pi = 0.30$, $y_0 = 0.1$)

Application I					Application II				
s	$\hat{t}_{(s)}(ML)$	UL	LL	length	s	$\hat{t}_{(s)}(ML)$	UL	LL	length
1	0.10803	0.02968	0.21743	0.18775	1	0.26590	0.19116	0.34402	0.15286
15	1.23907	0.81267	1.82910	1.01644	10	0.40791	0.33062	0.50349	0.17287
25	2.30459	1.19327	4.30294	3.10966	15	0.53358	0.41813	0.68931	0.27118

Estimation for Exponentiated Generalized Xgamma Distribution from New General Class Based on Dual Generalized Order Statistics

A. M. Abd AL-Fattah, R. E. Abd EL-Kader, G. R. AL-Dayian and A. A. EL-Helbawy

Statistics Department, Faculty of Commerce,
AL-Azhar University (Girls' Branch), Cairo, Egypt

Abstract

In this paper, the exponentiated generalized general class of distributions is introduced. Bayes estimators for the parameters, reliability and hazard rate functions are derived based on dual generalized order statistics. As a special model from exponentiated generalized general class distributions, Bayesian estimation for the unknown parameters, reliability and hazard rate functions of exponentiated generalized xgamma distribution based on dual generalized order statistics are considered. All results are specialized to lower records, also a numerical study is given to illustrate the theoretical procedures.

Keywords: *Exponentiated generalized distributions; Bayesian estimation; Dual generalized order statistics; Exponentiated generalized xgamma distribution.*

1. Introduction

The statistical literature contains many classes of distributions that have been constructed by extending the common families of continuous distributions providing more flexibility as far as applications is concerned. These new families have been used for modeling data in many applied areas such as engineering, economics, biological studies, environmental sciences lifetime analysis, finance and insurance. These *generalized* (G) distributions give more flexibility by adding new parameters to the baseline model and are useful in obtaining general results that could be applied to special cases to obtain new results.[Alzaatreh *et al.* (2013) and Alizadeh *et al.* (2017)].

There are several ways to generate new distributions, such as adding a positive parameter to a general survival function by Marshall and Olkin (1997). Also, Gupta *et al.* (1998) introduced the *exponentiated* (E) exponential distribution; by exponentiating a cumulative distribution function with a positive real number, beta generalized-G family by Eugene *et al.* (2002), G gamma generated distributions by Zografos and Balakrishnan (2009). Cordeiro and de Castro (2011) presented a class of Kumaraswamy G distribution and Alzaatreh *et al.* (2013) introduced a new method for generating families of continuous distributions.

Many authors focused on the E distributions and its applications; for example, Nadarajah and Kotz (2006), Ali *et al.* (2007), Silva *et al.* (2010), Lemonte *et al.* (2013), Elgarhy and Shawki (2017) and Rather and Subramanian (2018).

Cordeiro *et al.* (2013) proposed a class of distributions as an extension of the E type distribution which has greater flexibility of its tails and can be widely applied in many areas of biology and engineering. Given a non-negative continuous random variable X , the *cumulative distribution function* (cdf) of the *Exponentiated Generalized* (EG) general class of distribution is defined by

$$F(x; \alpha, \beta) = [1 - (1 - G(x))^\alpha]^\beta, \quad \alpha, \beta > 0, \quad (1)$$

where α and β are additional shape parameters, the corresponding *probability density function* (pdf) for (1) is given by

$$f(x; \alpha, \beta) = \alpha\beta g(x)(1 - G(x))^{\alpha-1}[1 - (1 - G(x))^\alpha]^{\beta-1}, \quad \alpha, \beta > 0. \quad (2)$$

By setting $\alpha = 1$ in (1) the E type distributions is derived by Gupta *et al.* (1998); further the E exponential and E gamma distributions can be obtained if $G(x)$ is the exponential or gamma cumulative distributions, respectively. For $\beta = 1$ in (1) and if $G(x)$ is the Gumbel or Fréchet cumulative distributions, then, one can get the E Gumbel and E Fréchet distributions, respectively, as defined by Nadarajah and Kotz (2006). Thus, the class of distributions (1) extends both E type distributions.

General class of distributions is important to obtain a general result that could be applied to special cases in obtaining new results. It is more flexible in dealing with statistical problems. Many authors focused on the G and EG distributions and its applications; for example, Oguntunde *et al.* (2014), Yousof *et al.* (2015), De Andrade *et al.* (2016), Mustafa *et al.* (2016) and Sindi *et al.* (2017).

AL-Hussaini (1999) introduced a general class of distributions with positive domain to be the underlying population model. This class of distributions includes, the Weibull, Pareto Type I, beta Type I, Gompertz and Compound Gompertz distribution, among others.

$$G_1(x|\underline{\theta}) = 1 - \exp[-\lambda(x)], \quad x > 0, \quad (3)$$

where $\lambda(x) \equiv \lambda(x, \underline{\theta})$ is a non-negative continuous function of x such that $\lambda(x, \underline{\theta}) \rightarrow 0$ as $x \rightarrow 0^+$ and $\lambda(x, \underline{\theta}) \rightarrow \infty$ as $x \rightarrow \infty$, and $\underline{\theta} = (\theta_1, \theta_2, \dots, \theta_s)$.

In this paper the new *exponentiated generalized general class* (EGGC) of distributions will be introduced; using (1) and (3), the cdf and pdf can be derived as follows:

$$F(x; \alpha, \beta) = [1 - \exp[-\alpha\lambda(x)]]^\beta, \quad x > 0, \quad \alpha, \beta > 0, \quad (4)$$

and

$$f(x; \alpha, \beta) = \alpha\beta\lambda(x)\exp[-\alpha\lambda(x)][1 - \exp[-\alpha\lambda(x)]]^{\beta-1}, x > 0, \alpha, \beta > 0, \quad (5)$$

where $\lambda(x)$ is the derivative of $\lambda(x)$ with respect to x .

The *reliability function* (rf), *hazard rate function* (hrf) and *reversed hazard rate function* (rhrf) are given, respectively, by

$$R(x) = 1 - [1 - \exp[-\alpha\lambda(x)]]^{\beta}, \alpha, \beta, x > 0, \quad (6)$$

$$h_1(x) = \frac{f(x)}{R(x)} = \frac{\alpha\beta\lambda(x)\exp[-\alpha\lambda(x)][1 - \exp[-\alpha\lambda(x)]]^{\beta-1}}{1 - [1 - \exp[-\alpha\lambda(x)]]^{\beta}}, x > 0; \alpha, \beta > 0, \quad (7)$$

and

$$h_2(x) = \frac{f(x)}{F(x)} = \alpha\beta\lambda(x)\exp[-\alpha\lambda(x)][1 - \exp[-\alpha\lambda(x)]]^{-1}, x > 0; \alpha, \beta > 0. \quad (8)$$

Burkschat *et al.* (2003) studied the *dual generalized order statistics* (dgos) that enables a common approach to descending ordered random variables as reversed ordered order statistics, lower record models and lower Pfeifer records.

Let $X_{(1,n,m,k)}, X_{(2,n,m,k)}, \dots, X_{(n,n,m,k)}$ be n dgos from an absolutely cdf with corresponding pdf. Then, the joint pdf has the form

$$f_{X_{(1,n,m,k)}, X_{(2,n,m,k)}, \dots, X_{(n,n,m,k)}}(x_{(1)}, \dots, x_{(n)}) = k \left(\prod_{j=1}^{n-1} \gamma_j \right) \left[\prod_{i=1}^{n-1} \left(F(x_{(i)}) \right)^m f(x_{(i)}) \right] \left(F(x_{(n)}) \right)^{k-1} f(x_{(n)}), \quad (9)$$

where $F^{-1}(1) \geq x_{(1)} \dots \geq x_{(n)} \geq F^{-1}(0), n \in N, k \geq 1, m_1, \dots, m_{n-1} = m,$

$m \in R$ is the parameters such that $\gamma_r = k + (n - r)(m + 1) \geq 1, for all 1 \leq r \leq n.$

This paper is organized as follows: in Section 2, Bayes estimators for the parameters, rf and hrf of EGGC of distributions based on dgos under *squared error* (SE) and *linear exponential* (LINEX) loss functions are derived. The *exponentiated generalized xgamma* (EG-Xg) distribution is studied in details as an application for this general class in Section 3. Finally a numerical study is presented in Section 4, to illustrate the application procedures of the various results developed in this paper.

2. Bayesian Estimation Based on Dual Generalized Order Statistics

This section is devoted to estimate the parameters, rf and hrf based on dgos using Bayesian approach, under SE and LINEX loss functions. Also the credible intervals are obtained.

Suppose that $X_{(1,n,m,k)}, X_{(2,n,m,k)}, \dots, X_{(n,n,m,k)}$ are n dgos from EGGC distribution, the likelihood function can be obtained by substituting (4) and (5) in (9) as follows:

$$L(\alpha, \beta; \underline{x}) = k \left(\prod_{j=1}^{n-1} \gamma_j \right) \alpha^n \beta^n \prod_{i=1}^n \lambda(x_i) \exp \left[-\alpha \sum_{i=1}^n \lambda(x_i) \right] \prod_{i=1}^{n-1} [1 - \exp[-\alpha \lambda(x_i)]]^{\beta(m+1)-1} \\ \times [1 - \exp[-\alpha \lambda(x_n)]]^{\beta k - 1}. \quad (10)$$

Assuming that the parameters α and β of EGGC distributions are random variables with a joint bivariate prior density function that was used by AL-Hussaini and Jaheen (1992) as

$$g(\alpha, \beta) = g_1(\beta|\alpha)g_2(\alpha), \quad (11)$$

$$\text{where } g_1(\beta|\alpha) = \frac{\alpha^{\tau+1}}{\Gamma(\tau+1)w^{\tau+1}} \beta^\tau e^{-\frac{\alpha\beta}{w}}, \quad \tau > -1, w, \alpha, \beta > 0, \quad (12)$$

and the prior of α is

$$g_2(\alpha) = \frac{\alpha^{c-1}}{\Gamma(c)b^c} e^{-\frac{\alpha}{b}}, \quad \alpha, b, c > 0, \quad (13)$$

which is the gamma (c, b) density.

The joint prior pdf of α and β , will be obtained by substituting (12) and (13) in (11) and it's given by;

$$g(\alpha, \beta) \propto \alpha^{c+\tau} \beta^\tau e^{-\alpha(\frac{1}{b} + \frac{\beta}{w})}, \quad c, b, w > 0, \tau > -1. \quad (14)$$

The joint posterior of α and β can be derived by using (10) and (14) as follows:

$$\pi(\alpha, \beta|\underline{x}) \propto L(\alpha, \beta|\underline{x})g(\alpha, \beta)$$

$$\pi(\alpha, \beta|\underline{x}) \propto \alpha^{n+\tau+c} \beta^{n+\tau} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} e^{-\beta[\frac{\alpha}{w} - (m+1)\sum_{i=1}^{n-1} \ln u_i - k \ln u_n]} \prod_{i=1}^n (u_i)^{-1},$$

$$\text{where } u_i = [1 - \exp[-\alpha \lambda(x_i)]] \text{ and } u_n = [1 - \exp[-\alpha \lambda(x_n)]], \quad (15)$$

Hence, the joint posterior distribution of α and β is given by;

$$\pi(\alpha, \beta|\underline{x}) = \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} e^{-\beta[\frac{\alpha}{w} - (m+1)\sum_{i=1}^{n-1} \ln u_i - k \ln u_n]} \prod_{i=1}^n (u_i)^{-1}}{\Gamma(n+\tau+1) \varphi}, \quad (16)$$

$$\text{where } \varphi = \int_0^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1}}{\left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n\right]^{n+\tau+1}} d\alpha. \quad (17)$$

2.1 Point estimation

In this subsection the Bayes estimators for the parameters, rf and hrf based on dgos under SE and LINEX loss functions of the EGGC distribution are obtained.

a. Bayesian estimation of the parameters, rf and hrf under squared error loss function

Under SE loss function the Bayes estimators of the parameters α and β are given by their marginal posterior expectations using (16) as shown below

$$\begin{aligned} \alpha_{(SE)}^* &= E(\alpha|\underline{x}) \\ &= \int_0^\infty \frac{1}{\varphi\left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n\right]^{n+\tau+1}} \alpha^{n+\tau+c+1} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1} d\alpha, \end{aligned} \quad (18)$$

and

$$\begin{aligned} \beta_{(SE)}^* &= E(\beta|\underline{x}) \\ &= \int_0^\infty \frac{(n+\tau+1) \alpha^{n+\tau+c} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1}}{\varphi\left(\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n\right)^{n+\tau+2}} d\alpha. \end{aligned} \quad (19)$$

The Bayes estimators of the rf and hrf under SE loss function are as follows:

$$\begin{aligned} R_{(SE)}^*(x) &= E(R(x)|\underline{x}) \\ &= 1 - \int_0^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1}}{\varphi\left(\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n - \ln u\right)^{n+\tau+1}} d\alpha, \end{aligned} \quad (20)$$

and

$$\begin{aligned} h_{1(SE)}^*(x) &= E(h(x)|\underline{x}) = \\ &= \int_0^\infty \int_0^\infty \frac{\alpha^{n+\tau+c+1} \beta^{n+\tau+1} \lambda(x) \exp[-\alpha \lambda(x)] u^{\beta-1} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} e^{-\beta[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n]} \prod_{i=1}^n (u_i)^{-1}}{(1-u^\beta) \Gamma(n+\tau+1) \varphi} d\alpha d\beta. \end{aligned} \quad (21)$$

$$\begin{aligned} \text{where } u &= [1 - \exp[-\alpha \lambda(x)]] \\ \text{and } u_i \text{ and } u_n &\text{ are given by (15).} \end{aligned} \quad (22)$$

To obtain the Bayes estimates of the parameters, rf and hrf, the Equations (18) - (21) should be solved numerically.

b. Bayesian estimation of the parameters, rf and hrf under linear exponential loss function

Under the LINEX loss function, the Bayes estimators for the shape parameters α and β are given, respectively, by

$$\alpha_{(LNX)}^* = \frac{-1}{\vartheta} \ln E(e^{-\vartheta\alpha} | \underline{x}), \quad (23)$$

where

$$E(e^{-\vartheta\alpha} | \underline{x}) = \int_0^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha[\vartheta + \frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1}}{\varphi[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n]^{n+\tau+1}} d\alpha, \quad (24)$$

and

$$\beta_{(LNX)}^* = \frac{-1}{\vartheta} \ln E(e^{-\vartheta\beta} | \underline{x}), \quad (25)$$

where

$$E(e^{-\vartheta\beta} | \underline{x}) = \int_0^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} \prod_{i=1}^n (u_i)^{-1}}{\varphi(\vartheta + \frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n)^{n+\tau+1}} d\alpha. \quad (26)$$

Also, the Bayes estimator for rf based on dgos can be obtained as follows:

$$R_{(LNX)}^*(x) = \frac{-1}{\vartheta} \ln E(e^{-\vartheta R(x)} | \underline{x}), \quad (27)$$

where

$$E(e^{-\vartheta R(x)} | \underline{x}) = 1 - \int_0^\infty \int_0^\infty \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\vartheta(1-u)^\beta} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} e^{-\beta[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n]} \prod_{i=1}^n (u_i)^{-1}}{\Gamma(n+\tau+1) \varphi} d\alpha d\beta, \quad (28)$$

and the Bayes estimators for hrf and rhf based on dgos can be derived as given below:

$$h_{1(LNX)}^*(x) = \frac{-1}{\vartheta} \ln E(e^{-\vartheta h_1(x)} | \underline{x}). \quad (29)$$

where

$$E(e^{-\vartheta h_1(\underline{x})}|\underline{x}) = \int_0^\infty \int_0^\infty \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\vartheta \left[\frac{\alpha \beta \lambda(\underline{x}) \exp[-\alpha \lambda(\underline{x})] [u] \beta^{-1}}{1-[u]^\beta} \right]} e^{-\alpha \left[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i) \right]} e^{-\beta \left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n \right]} \prod_{i=1}^n (u_i)^{-1}}{\Gamma(n+\tau+1) \varphi} d\alpha d\beta, \quad (30)$$

and

$$h_{2(LNX)}^*(x) = \frac{-1}{\vartheta} \ln E(e^{-\vartheta h_2(x)}|\underline{x}). \quad (31)$$

where

$$E(e^{-\vartheta h_2(\underline{x})}|\underline{x}) = \int_0^\infty \int_0^\infty \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\vartheta [\alpha \beta \lambda(\underline{x}) \exp[-\alpha \lambda(\underline{x})] u^{-1}]}}{\Gamma(n+\tau+1) \varphi} e^{-\alpha \left[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i) \right]} e^{-\beta \left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n \right]} \prod_{i=1}^n (u_i)^{-1}} d\alpha d\beta. \quad (32)$$

To obtain the Bayes estimators of the parameters, rf and hrf, the Equations (23) - (32) should be solved numerically.

2.2 Credible intervals

In this subsection the Bayesian analog to the confidence interval which is called a credible intervals are introduced. In general, $(L(\underline{x}), U(\underline{x}))$ is $100(1 - \omega)\%$ credible interval of $\underline{\theta}$ if

$$P[L(\underline{x}) < \underline{\theta} < U(\underline{x})|\underline{x}] = \int_{L(\underline{x})}^{U(\underline{x})} \pi^*(\underline{\theta}|\underline{x}) d\underline{\theta} = 1 - \omega. \quad (33)$$

Since, the posterior distribution is given by (16), then a 100 (1- ω) % credible interval for α is $(L(\underline{x}), U(\underline{x}))$,

where

$$P[\alpha > L(\underline{x})|\underline{x}] = \int_{L(\underline{x})}^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha \left[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i) \right]}}{\varphi \prod_{i=1}^n u_i \left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n \right]^{n+\tau+1}} d\alpha = 1 - \frac{\omega}{2}, \quad (34)$$

and

$$P[\alpha > U(\underline{x})|\underline{x}] = \int_{U(\underline{x})}^\infty \frac{\alpha^{n+\tau+c} e^{-\alpha \left[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i) \right]}}{\varphi \prod_{i=1}^n u_i \left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n \right]^{n+\tau+1}} d\alpha = \frac{\omega}{2}. \quad (35)$$

Also, 100 (1- ω) % credible interval for β is $(L(\underline{x}), U(\underline{x}))$,

where

$$\begin{aligned} P[\beta > L(\underline{x})|\underline{x}] &= \int_{L(\underline{x})}^\infty \int_0^\infty \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\alpha \left[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i) \right]} e^{-\beta \left[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n \right]} \prod_{i=1}^n (u_i)^{-1}}{\varphi \Gamma(n+\tau+1)} d\alpha d\beta \\ &= 1 - \frac{\omega}{2}, \end{aligned} \quad (36)$$

and

$$P[\beta > U(\underline{x})|\underline{x}] = \int_{U(\underline{x})}^{\infty} \int_0^{\infty} \frac{\alpha^{n+\tau+c} \beta^{n+\tau} e^{-\alpha[\frac{1}{b} + \sum_{i=1}^n \lambda(x_i)]} e^{-\beta[\frac{\alpha}{w} - (m+1) \sum_{i=1}^{n-1} \ln u_i - k \ln u_n]} \prod_{i=1}^n (u_i)^{-1}}{\varphi \Gamma(n+\tau+1)} d\alpha d\beta = \frac{\omega}{2}. \quad (37)$$

3. Exponentiated Generalized Xgamma Distribution Based on Dual Generalized Order Statistics

Sen *et al.* (2016) introduced the *xgamma* (xg) distribution which is generated as a special finite mixture of exponential(θ) and gamma(3, θ) distributions with mixing proportion $\pi_1 = \frac{\theta}{1+\theta}$ and $\pi_2 = 1 - \pi_1 = \frac{1}{1+\theta}$. The cdf and pdf of the xgamma distribution are, respectively

$$F_{xg}(x; \theta) = 1 - \frac{1+\theta+\theta x + \frac{\theta^2 x^2}{2}}{1+\theta} e^{-\theta x}, \quad x > 0, \theta > 0, \quad (38)$$

and

$$f_{xg}(x; \theta) = \frac{\theta^2}{1+\theta} \left(1 + \frac{\theta x^2}{2}\right) e^{-\theta x}, \quad x > 0, \theta > 0. \quad (39)$$

Yadav *et al.* (2018) studied the *generalized xgamma* (G-xg) distribution by adding power shape parameter to the cdf; some statistical properties of this G-xg are discussed. They used many methods of estimation to estimate the rf and hrf of the G-xg distribution.

Assuming that X is a random variable distribution with EG-xg distribution with shape parameters, $\alpha, \beta > 0$ and scale parameter $\theta > 0$ denoted by $X \sim \text{EGxg}(\alpha, \beta, \theta)$, hence the pdf and cdf are given, respectively, by

$$F_{EGxg}(x; \alpha, \beta, \theta) = \left[1 - \frac{(1+\theta+\theta x + \frac{\theta^2 x^2}{2})^\alpha}{(1+\theta)^\alpha} e^{-\alpha \theta x}\right]^\beta, \quad x > 0, \alpha, \beta, \theta > 0, \quad (40)$$

and

$$f_{EGxg}(x; \alpha, \beta, \theta) = \frac{\alpha \beta \theta^2 e^{-\alpha \theta x}}{(1+\theta)^\alpha} \left(1 + \frac{\theta x^2}{2}\right) \left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^{\alpha-1} \times \left[1 - \frac{(1+\theta+\theta x + \frac{\theta^2 x^2}{2})^\alpha}{(1+\theta)^\alpha} e^{-\alpha \theta x}\right]^{\beta-1}, \quad x > 0, \alpha, \beta, \theta > 0. \quad (41)$$

The rf, hrf and rhf are, respectively, given by

$$R_{EGxg}(x; \alpha, \beta, \theta) = 1 - \left[1 - \frac{\left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^\alpha}{(1 + \theta)^\alpha} e^{-\alpha \theta x} \right]^\beta, \quad x > 0, \alpha, \beta, \theta > 0, \quad (42)$$

$$\begin{aligned} h_{1EGxg}(x; \alpha, \beta, \theta) &= \frac{\frac{\alpha \beta \theta^2 e^{-\alpha \theta x}}{(1 + \theta)^\alpha} \left(1 + \frac{\theta x^2}{2}\right) \left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^{\alpha-1} \left[1 - \frac{\left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^\alpha}{(1 + \theta)^\alpha} e^{-\alpha \theta x} \right]^{\beta-1}}{1 - \left[1 - \frac{\left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^\alpha}{(1 + \theta)^\alpha} e^{-\alpha \theta x} \right]^\beta}, \\ & \quad x > 0, \alpha, \beta, \theta > 0, \quad (43) \end{aligned}$$

and

$$\begin{aligned} h_{2EGxg}(x; \alpha, \beta, \theta) &= \frac{\alpha \beta \theta^2 e^{-\alpha \theta x}}{(1 + \theta)^\alpha} \left(1 + \frac{\theta x^2}{2}\right) \left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^{\alpha-1} \\ & \quad \times \left[1 - \frac{\left(1 + \theta + \theta x + \frac{\theta^2 x^2}{2}\right)^\alpha}{(1 + \theta)^\alpha} e^{-\alpha \theta x} \right]^{-1}, \quad x, \alpha, \beta, \theta > 0. \end{aligned} \quad (44)$$

Plots of the pdf, hrf and rhrf of EG-xg are, respectively, given in Figures 1-3. The plots, in Figure 1 indicate the behavior of the density function and explain the flexibility of the model graphically with its sub-families. From Figure 1, one can observe that the curves of the pdf are monotone decreasing, increasing and bathtub, with different values of shape and scale parameters. In Figure 2, the curves of the hrf are increasing, decreasing and monotone decreasing with different values of shape and scale parameters. The curves of the rhrf at the all values are decreasing and then constant in Figure 3.

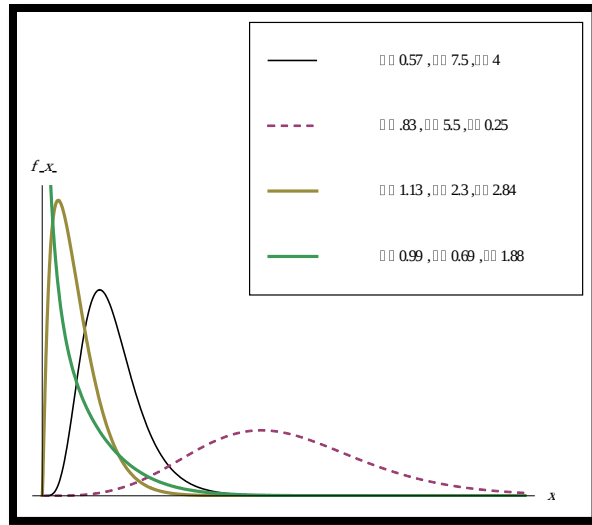


Figure 1

Plots of the probability density function

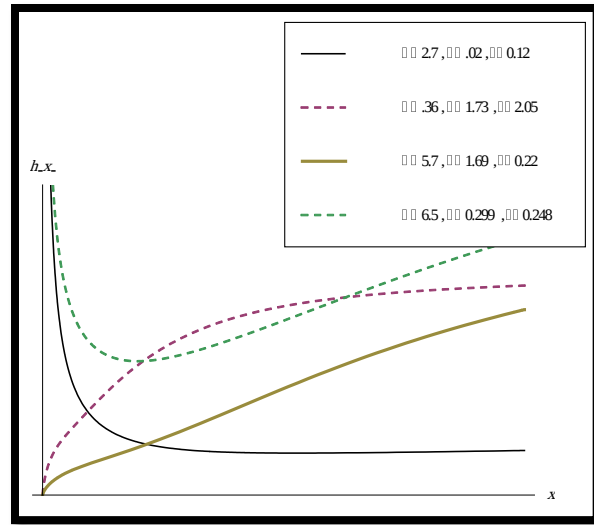


Figure 2

Plots of the hazard rate function

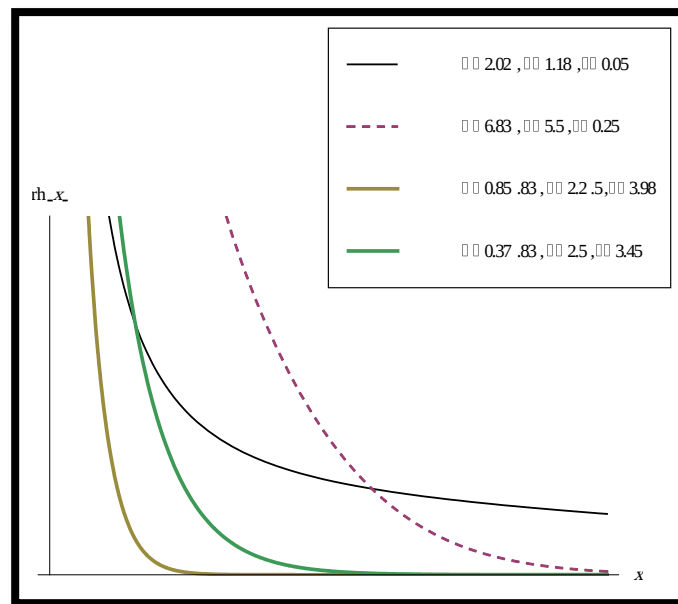


Figure 3

Plots of the reversed hazard rate function

3.1 Bayesian estimation for exponentiated generalized xgamma distribution

This section derived the estimation of the parameters, rf and hrf based on dgos using Bayesian approach. Also the credible intervals of the parameters, rf and hrf are obtained.

Suppose that $X_{(1,n,m,k)}, X_{(2,n,m,k)}, \dots, X_{(n,n,m,k)}$ are n dgos from EG-Xg distribution, the likelihood function can be derived by substituting (40) and (41) in (9) as follows:

$$L_{EGXg}(\alpha, \beta, \theta; \underline{x}) = k \left(\prod_{j=1}^{n-1} \gamma_j \right) \frac{\alpha^n \beta^n \theta^{2n} e^{-\alpha \theta \sum_{i=1}^n x_i}}{(1+\theta)^{n\alpha}} \prod_{i=1}^n \left(1 + \frac{\theta x_i^2}{2} \right) \delta_i^{\alpha-1} \\ \times \prod_{i=1}^{n-1} \left[1 - \left(\frac{\delta_i}{(1+\theta)} \right)^\alpha e^{-\alpha \theta x_i} \right]^{\beta(m+1)-1} \left[1 - \left(\frac{\delta_n}{(1+\theta)} \right)^\alpha e^{-\alpha \theta x_n} \right]^{\beta k-1}, \quad (45)$$

$$\text{where } \delta_i = \left(1 + \theta + \theta x_i + \frac{\theta^2 x_i^2}{2} \right) \text{ and } \delta_n = \left(1 + \theta + \theta x_n + \frac{\theta^2 x_n^2}{2} \right). \quad (46)$$

3.1.1 Point estimation

In this subsection, the Bayesian approach is considered under SE and LINEX loss functions to estimate the parameters, rf and hrf of the EG-xg distribution based on dgos. Also credible intervals for the parameters are obtained.

Let $\vartheta_1 = \alpha, \vartheta_2 = \beta$ and $\vartheta_3 = \theta$ are independent random variables with gamma prior distribution with the pdf as follows:

$$\pi(\vartheta_j) = \frac{d_j^{c_j}}{\Gamma(c_j)} \vartheta_j^{c_j-1} e^{-d_j \vartheta_j}, \quad \vartheta_j, d_j, c_j > 0, \quad j = 1, 2, 3, \text{ where } c_j, d_j \text{ are the hyper parameters.}$$

A joint prior density function of $\underline{\vartheta} = (\vartheta_1, \vartheta_2, \vartheta_3)'$ is then given by

$$\pi(\underline{\vartheta}) \propto \prod_{j=1}^3 \vartheta_j^{c_j-1} e^{-d_j \vartheta_j} \\ \propto \alpha^{c_1-1} \beta^{c_2-1} \theta^{c_3-1} e^{-[d_1 \alpha + d_2 \beta + d_3 \theta]}. \quad (47)$$

The likelihood function given by (45) can be rewritten as

$$L_{EGXg}(\alpha, \beta, \theta; \underline{x}) = k \left(\prod_{j=1}^{n-1} \gamma_j \right) \frac{\alpha^n \beta^n \theta^{2n} e^{-\alpha \theta \sum_{i=1}^n x_i}}{(1+\theta)^{n\alpha}} \prod_{i=1}^n \rho_i (u_i^*)^{-1} \prod_{i=1}^{n-1} u_i^{*[\beta(m+1)]} u_n^{*[\beta k]}, \quad (48)$$

$$\text{where } u_i^* = \left[1 - \left(\frac{\delta_i}{(1+\theta)} \right)^\alpha e^{-\alpha \theta x_i} \right], \quad u_n^* = \left[1 - \left(\frac{\delta_n}{(1+\theta)} \right)^\alpha e^{-\alpha \theta x_n} \right] \text{ and } \rho_i = \left(1 + \frac{\theta x_i^2}{2} \right) \delta_i^{\alpha-1}.$$

The joint posterior density can be derived by using (47) and (48) as follows:

$$\begin{aligned}
 & \pi_{EGXg}^*(\underline{y}|\underline{x}) \\
 &= T(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} \\
 & \times e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1}, \quad (49)
 \end{aligned}$$

where T is the normalizing constant defined by

$$\begin{aligned}
 & T^{-1} \\
 &= \int_0^\infty \int_0^\infty \left(\alpha^{c_1+n-1} \theta^{c_3+2n-1} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \right) \\
 & \times \left[\int_0^\infty \beta^{c_2+n-1} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta \right] d\alpha d\theta. \\
 &= \Gamma(c_2 + n) \int_0^\infty \int_0^\infty \left(\alpha^{c_1+n-1} \theta^{c_3+2n-1} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \right) \\
 & \times \left[\frac{1}{[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \right] d\alpha d\theta
 \end{aligned}$$

Let

$$\begin{aligned}
 \varphi^* &= \int_0^\infty \int_0^\infty \left(\alpha^{c_1+n-1} \theta^{c_3+2n-1} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \right) \\
 & \times \left[\frac{1}{[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \right] d\alpha d\theta
 \end{aligned}$$

Hence, the joint posterior distribution of α, β and θ given $\underline{x}$ can be written as;

$$\pi_{EGXg}^*(\alpha, \beta, \theta | \underline{x}) = \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]}}{\varphi^* \Gamma(c_2+n)}, \quad (50)$$

a. Bayesian estimation of the parameters, rf and hrf under squared error loss function

Under SE loss function the Bayes estimators of the parameters α, β and θ are given by their marginal posterior expectations using (50) as shown below

$$\begin{aligned}
 \alpha_{(SE)EGXg}^* &= E(\alpha | \underline{x}) \\
 &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{1}{\varphi^* \Gamma(c_2 + n)} (\alpha^{c_1+n} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} \\
 & \times e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\theta d\alpha,
 \end{aligned}$$

$$= \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n} \theta^{c_3+2n-1})}{\varphi^*[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\theta d\alpha, \quad (51)$$

$$\begin{aligned} \beta_{(SE)EGXg}^* &= E(\beta|\underline{x}) \\ &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{1}{\varphi^*\Gamma(c_2+n)} (\alpha^{c_1+n-1} \beta^{c_2+n} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \\ &\times e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\alpha d\theta, \\ &= \int_0^\infty \int_0^\infty \frac{(c_2+n)(\alpha^{c_1+n-1} \theta^{c_3+2n}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1}}{\varphi^*[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n+1}} d\alpha d\theta, \end{aligned} \quad (52)$$

and

$$\begin{aligned} \theta_{(SE)EGXg}^* &= E(\theta|\underline{x}) \\ &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n})}{\varphi^*\Gamma(c_2+n)} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \\ &\times e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\alpha d\theta, \\ &= \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1}}{\varphi^*[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} d\alpha d\theta, \end{aligned} \quad (53)$$

The Bayes estimators of the rf, hrf and rhrf under SE loss function can be obtained using (42), (43) and (50) as follows:

$$\begin{aligned} R_{(SE)EGXg}^*(x) &= E(R_{EGXg}(x)|\underline{x}) = 1 - \int_{\underline{g}} (u^*)^\beta \pi_{EGXg}^*(\underline{g}|\underline{x}) d\underline{g} \\ &= 1 \\ &- \int_0^\infty \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1})}{\varphi^*\Gamma(c_2+n)} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \\ &\times e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*) - \ln u^*]} d\beta d\theta d\alpha, \\ &= 1 - \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1}}{\varphi^*[d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*) - \ln u^*]^{c_2+n}} d\theta d\alpha, \end{aligned}$$

(54)

$$\text{where } u^* = \left[1 - \left(\frac{\delta}{(1+\theta)} \right)^\alpha e^{-\theta \alpha x} \right]$$

and

$$\begin{aligned} h_{1(SE)EGXg}^*(x) &= E(h_{1EGXg}(x)|\underline{x}) \\ &= \int_{\underline{\vartheta}} h_{1EGXg}(x) \pi_{EGXg}^*(\underline{\vartheta}|\underline{x}) d\underline{\vartheta}, \end{aligned} \quad (55)$$

To obtain the Bayes estimators of the parameters, rf, hrf and rhrf, the Equations (51) - (55) should be solved numerically.

b. Bayesian estimation of the parameters, rf and hrf under linear exponential loss function

Under the LINEX loss function, the Bayes estimators for the shape parameters α , β and θ are given, respectively, by

$$\alpha_{(LNX)EGXg}^* = \frac{-1}{v} \ln E(e^{-v\alpha}|\underline{x}), \quad (56)$$

where

$$\begin{aligned} &E(e^{-v\alpha}|\underline{x}) \\ &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{1}{\varphi^* \Gamma(c_2 + n)} (\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta)+v)} e^{-\theta d_3} \\ &\times \prod_{i=1}^n \rho_i (u_i^*)^{-1} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\theta d\alpha, \\ &= \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta)+v)} e^{-\theta d_3}}{\varphi^* [d_2 - (m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i (u_i^*)^{-1} d\theta d\alpha, \end{aligned}$$

and

$$\beta_{(LNX)EGXg}^* = \frac{-1}{v} \ln E(e^{-v\beta}|\underline{x}), \quad (57)$$

where

$$\begin{aligned}
 & E(e^{-\nu\beta}|\underline{x}) \\
 &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1})}{\varphi^* \Gamma(c_2+n)} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \\
 &\times e^{-\beta[v+d_2-(m+1)\sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\alpha d\theta, \\
 &= \int_0^\infty \int_0^\infty \frac{(c_2+n)(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3}}{\varphi^*[v+d_2-(m+1)\sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta,
 \end{aligned}$$

Also

$$\theta_{(LN X)EGXg}^* = \frac{-1}{\nu} \ln E(e^{-\nu\theta}|\underline{x}), \quad (58)$$

where

$$\begin{aligned}
 & E(e^{-\nu\theta}|\underline{x}) \\
 &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta(d_3+\nu)}}{\varphi^* \Gamma(c_2+n)} \prod_{i=1}^n \rho_i(u_i^*)^{-1} \\
 &\times e^{-\beta[d_2-(m+1)\sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} d\beta d\alpha d\theta, \\
 &= \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta(d_3+\nu)}}{\varphi^*[d_2-(m+1)\sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta.
 \end{aligned}$$

Also, the Bayes estimator for rf based on dgos can be obtained as follows:

$$R_{(LN X)EGXg}^*(x) = \frac{-1}{\nu} \ln E(e^{-\nu R_{EGXg}(x)}|\underline{x}), \quad (59)$$

where

$$\begin{aligned}
 & E(e^{-\nu R(x)}|\underline{x}) = \\
 &= \\
 &\int_0^\infty \int_0^\infty \int_0^\infty \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3 - \nu(1-(u^*)^\beta)}}{\varphi^* \Gamma(c_2+n)} e^{-\beta[d_2-(m+1)\sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*) - \ln u^*]} \times \\
 &\prod_{i=1}^n \rho_i(u_i^*)^{-1} d\beta d\theta d\alpha,
 \end{aligned}$$

and the Bayes estimator for hrf and rhrf based on dgos can be derived as follows:

$$h_{1(LN X)EGXg}^*(x) = \frac{-1}{\nu} \ln E(e^{-\nu h_1(x)}|\underline{x}). \quad (60)$$

where

$$E(e^{-\nu h_1(x)}|\underline{x}) = \int_0^\infty \int_0^\infty \int_0^\infty e^{-\nu h_1(x)} \pi_{EGXg}^*(\alpha, \beta, \theta|\underline{x}) d\beta d\theta d\alpha.$$

3.1.2 Credible interval

Since, the posterior distribution is given by (51), then a 100 (1- ω) % credible interval for α is $(L(\underline{x}), U(\underline{x}))$,

where

$$\begin{aligned} P[\alpha_{EGXg} > L(\underline{x})|\underline{x}] &= \int_{L(\underline{x})}^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3}}{\varphi^*[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\theta d\alpha \\ &= 1 - \frac{\omega}{2}. \end{aligned} \quad (61)$$

and

$$\begin{aligned} P[\alpha_{EGXg} > U(\underline{x})|\underline{x}] &= \int_{U(\underline{x})}^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3}}{\varphi^*[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\theta d\alpha \\ &= \frac{\omega}{2}. \end{aligned} \quad (62)$$

Also, 100 (1- ω) % credible interval for β is $(L(\underline{x}), U(\underline{x}))$,

where

$$\begin{aligned} P[\beta_{EGXg} > L(\underline{x})|\underline{x}] &= \int_{L(\underline{x})}^{\infty} \int_0^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1})}{\Gamma(c_2+n) \varphi^*} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} \\ &\times \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta d\beta \\ &= 1 - \frac{\omega}{2}, \end{aligned} \quad (63)$$

and

$$\begin{aligned} P[\beta_{EGXg} > U(\underline{x})|\underline{x}] &= \int_{U(\underline{x})}^{\infty} \int_0^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1} \beta^{c_2+n-1} \theta^{c_3+2n-1})}{\Gamma(c_2+n) \varphi^*} e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3} e^{-\beta[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]} \\ &\times \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta d\beta = \frac{\omega}{2}. \end{aligned} \quad (64)$$

Also, 100 (1- ω) % credible interval for θ is $(L(\underline{x}), U(\underline{x}))$,

$$\begin{aligned} P[\theta_{EGXg} > L(\underline{x})|\underline{x}] &= \int_{L(\underline{x})}^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1} \theta^{c_3+2n-1}) e^{-\alpha(d_1+\theta \sum_{i=1}^n x_i + n \ln(1+\theta))} e^{-\theta d_3}}{\varphi^*[d_2-(m+1) \sum_{i=1}^{n-1} \ln(u_i^*) - k \ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta \\ &= 1 - \frac{\omega}{2}, \end{aligned} \quad (65)$$

and

$$P[\theta_{EGxg} > U(\underline{x})|\underline{x}] = \int_{U(\underline{x})}^{\infty} \int_0^{\infty} \frac{(\alpha^{c_1+n-1}\theta^{c_3+2n-1})e^{-\alpha(d_1+\theta\sum_{i=1}^n x_i+n\ln(1+\theta))}e^{-\theta d_3}}{\varphi^*[d_2-(m+1)\sum_{i=1}^{n-1}\ln(u_i^*)-k\ln(u_n^*)]^{c_2+n}} \prod_{i=1}^n \rho_i(u_i^*)^{-1} d\alpha d\theta$$

$$= \frac{\omega}{2}. \quad (66)$$

4. Numerical Results

This section aims to illustrate the theoretical results of Bayesian estimation under SE and LINEX loss functions. Numerical results are presented for EG-xg distribution based on lower record values through a simulation study and some applications.

4.1 Simulation study

The lower record values can be obtained as a special case from dogs by setting $m = -1$, $k = 1$; the estimation results obtained in Sections 2 and 3 can be specialized to lower records. The Bayes estimates of α, β, θ , rf and hrf and their average estimates, *estimated risks* (ERs) are computed based on lower record values through Monte Carlo simulation study according to the following steps:

- a. To generate random number from EG-xg with shape parameters α, β and scale parameter θ , the following steps may be used.
 - Specified the values α, β, θ and n .
 - Generate U_i from uniform (0, 1) distribution ($i = 1, 2, 3, \dots, n$).
 - Generate V_i from gamma (1, θ) distribution ($i = 1, 2, 3, \dots, n$).
 - Generate W_i from gamma (3, θ) distribution ($i = 1, 2, 3, \dots, n$).
 - If $U_i \leq \frac{\theta}{1+\theta}$, set $X_i = V_i$, otherwise set $X_i = W_i$.
- b. For each sample size n , consider the first observation is the first lower record value x_1 denote it by R_1 and the second observation x_2 denote it by R_2 ; which is smaller than the maximum ($x_1 > x_2$) record and if $x_1 \leq x_2$ ignore it and repeat until you get sample of *record values* (Rv) records.
- c. At the number of the surviving units, the population parameter values α, β, θ and the hyper parameters of the prior distribution, the Bayes estimates of the parameters, rf and hrf under SE and LINEX loss functions are computed. The computations are performed using R programming language.
- d. Tables 1 and 2 present the Bayes estimates under SE and LINEX loss functions of the parameters and their ERs, *relative absolute biases* (RAB) and credible intervals based on lower record values for different population parameter values for $\alpha = (0.2, 2)$, $\beta =$

(1.3, 3) and $\theta = (0.9, 0.3)$, based on size of records $R_v=5, 7, 9$ and replications $NR = 10000$.

- e. Table 3 displays the Bayes averages and 95% confidence intervals of the rf and hrf at $t_0 = 0.5, 1, 1.5$ from EG-xg distribution based on lower record values for different sample size of records $R_v= 5, 9$, and replications $NR = 10000$.

4.2 Applications

In this subsection, the application to real data set is provided to illustrate the importance of the EG-xg distribution based on lower records. Table 4 displays Bayes estimates of rf, hrf and *standard deviation* (sd) from EG-xg distribution for the real data based on lower records. Bayes estimates of the parameters, sds and ERs for the real data based on lower records in Table 5.

To check the validity of the fitted model, Kolmogorov-Smirnov goodness of fit test is performed for each data set. The p values are given, respectively, 0.35 and 0.204. The p values in each case indicate that the model fits the data very well.

- I. The application is the vinyl chloride data obtained from clean upgrading, monitoring wells in mg/L; this data set was used for Bhaumik *et al.* (2009). The data are: 5.1, 1.2, 1.3, 0.6, 0.5, 2.4, 0.5, 1.1, 8.0, 0.8, 0.4, 0.6, 0.9, 0.4, 2.0, 0.5, 5.3, 3.2, 2.7, 2.9, 2.5, 2.3, 1.0, 0.2, 0.1, 0.1, 1.8, 0.9, 2.0, 4.0, 6.8, 1.2, 0.4, 0.2.
- II. The second data set are service times of 63 aircraft windshield from Tahir *et al.* (2015). The data are: 0.046, 1.436, 2.592, 0.140, 1.492, 2.600, 0.150, 1.580, 2.670, 0.248, 1.719, 2.717, 0.280, 1.794, 2.819, 0.313, 1.915, 2.820, 0.389, 1.920, 2.878, 0.487, 1.963, 2.950, 0.622, 1.978, 3.003, 0.900, 2.053, 3.102, 0.952, 2.065, 3.304, 0.996, 2.117, 3.483, 1.003, 2.137, 3.500, 1.010, 2.141, 3.622, 1.085, 2.163, 3.665, 1.092, 2.183, 3.695, 1.152, 2.240, 4.015, 1.183, 2.341, 4.628, 1.244, 2.435, 4.806, 1.249, 2.464, 4.881, 1.262, 2.543, 5.140.

4.3 Concluding remarks

- It is clear from Tables 1-3 that the ERs of the Bayes estimates of the parameters, rf and hrf performs better and the length of the credible intervals get shorter when the sample size of R_v increases.
- One can notice that the ERs for the estimates of the parameters, rf and hrf under LINEX loss function have the less values than the corresponding ERs of the estimates under SE loss function.
- The results obtained in this paper can be modified to obtain special results for sub-models of EG-xg distribution as follows:

- i. The exponentiated xgamma distribution, if $\alpha = 1$.
- ii. The xgamma distribution, if $\alpha = 1$ and $\beta = 1$.
- iii.

Acknowledgements

The authors would like to thank Referees and the Editor for their constructive comments and corrections which led to improvements of an earlier version of this article.

Table 1

Bayes averages of the parameters and their estimated risks, relative absolute biases and credible intervals based on lower records ($\alpha = 2, \beta = 1.3, \theta = 0.9, \nu = 0.5$ and $NR = 10000$)

Rv	Loss functions	Estimators	Average	ER	RAB	LL	UL	Length
5	SE	α^*	1.9986	2.85e-06	0.0006	1.9968	2.0003	0.0035
		β^*	1.2988	2.25e-06	0.0024	1.2967	1.3003	0.0035
		θ^*	0.9021	5.38e-06	0.0009	0.8999	0.9037	0.0037
	LINEX	α^*	2.0010	2.57e-06	0.0005	1.9994	2.0032	0.0037
		β^*	1.3007	1.08e-06	0.0017	1.2992	1.3019	0.0027
		θ^*	0.9016	3.31e-06	0.0006	0.8995	0.9028	0.0033
7	SE	α^*	1.9996	1.23e-06	0.0002	1.9978	2.0008	0.0029
		β^*	1.3010	1.00e-06	0.0008	1.2996	1.3019	0.0023
		θ^*	0.9004	1.08e-06	0.0005	0.8990	0.9011	0.0021
	LINEX	α^*	2.0006	7.56e-07	0.0003	1.9992	2.0015	0.0023
		β^*	1.2996	5.33e-07	0.0003	1.2982	1.3006	0.0025
		θ^*	0.8995	5.17e-07	0.0005	0.8984	0.9004	0.0020
9	SE	α^*	2.0003	6.50e-07	0.0001	1.9990	2.0014	0.0023
		β^*	1.2997	1.66e-07	0.0002	1.2989	1.3004	0.0014
		θ^*	0.8997	3.30e-07	0.0003	0.8987	0.9007	0.0019
	LINEX	α^*	1.9999	3.03e-07	4.52e-05	1.9986	2.0009	0.0022
		β^*	1.3001	2.56e-07	5.48e-05	1.2988	1.3007	0.0019
		θ^*	0.8999	1.54e-07	9.09e-05	0.8988	0.9005	0.0017

Table 2

Bayes averages for the parameters, their estimated risks, relative absolute biases and credible intervals based on lower records ($\alpha = 0.2, \beta = 3, \theta = 0.3, \nu = 0.5, NR = 10000$)

Rv	Loss functions	Estimators	Average	ER	RAB	LL	UL	Length
5	SE	α^*	0.1985	3.78e-06	0.0071	0.1962	0.2004	0.0042
		β^*	2.9992	1.41e-06	0.0002	2.9971	3.0002	0.0031
		θ^*	0.3011	3.79e-06	0.0036	0.2985	0.3036	0.0051
	LINEX	α^*	0.1979	4.75e-06	0.0101	0.1962	0.1990	0.0027
		β^*	3.0021	5.90e-06	0.0007	2.9996	3.0040	0.0044
		θ^*	0.2979	5.17e-06	0.0068	0.2964	0.2998	0.0034
7	SE	α^*	0.2008	1.13e-06	4.41e-03	0.1996	0.2018	0.0021
		β^*	2.9997	3.09e-07	8.23e-05	2.9985	3.0004	0.0019
		θ^*	0.2993	7.67e-07	2.18e-03	0.2981	0.3002	0.0021
	LINEX	α^*	0.1985	2.42e-06	0.0073	0.1975	0.1996	0.0021
		β^*	2.9990	1.10e-06	0.0003	2.9975	2.9998	0.0022
		θ^*	0.2995	4.52e-07	0.0015	0.2980	0.3001	0.0021
9	SE	α^*	0.1992	7.99e-07	3.84e-03	0.1980	0.1999	0.0019
		β^*	2.9998	1.27e-07	6.11e-05	2.9988	3.0003	0.0014
		θ^*	0.3006	5.59e-07	2.06e-03	0.2996	0.3013	0.0015
	LINEX	α^*	0.2002	3.76e-07	0.0013	0.1992	0.2011	0.0019
		β^*	2.9995	4.15e-07	0.0001	2.9983	3.0005	0.0021
		θ^*	0.30008	2.35e-07	0.0002	0.2988	0.30077	0.0019

Table 3

Bayes averages, their estimated risks, relative absolute biases and credible intervals of the rf and hrf at $t_0 = 0.5, 1, 1.5$ from EG-xg distribution based on SE and LENIX loss functions for different sample size of records Rv, and repetitions NR = 10000

Rv	t_0	Loss functions	Estimators	Average	ER	RAB	LL	UL	Length
5	0.5	SE	$R_{EGxg}^*(t_0)$	0.6235	1.63e-06	0.0010	0.6216	0.6248	0.0032
			$h_{1EGxg}^*(t_0)$	0.3315	4.58e-06	0.0057	0.3297	0.3331	0.0033
		LINEX	$R_{EGxg}^*(t_0)$	0.6254	1.83e-06	0.0019	0.6239	0.6262	0.0023
			$h_{1EGxg}^*(t_0)$	0.3327	2.07e-06	0.0022	0.3304	0.3342	0.0038
	1	SE	$R_{EGxg}^*(t_0)$	0.5425	6.52e-06	0.0039	0.5399	0.5446	0.0046
			$h_{1EGxg}^*(t_0)$	0.2351	8.73e-07	0.0021	0.2334	0.2368	0.0033
		LINEX	$R_{EGxg}^*(t_0)$	0.5467	5.57e-06	0.0039	0.5439	0.5478	0.0039
			$h_{1EGxg}^*(t_0)$	0.2343	4.05e-07	0.0014	0.2334	0.2352	0.0018
	1.5	SE	$R_{EGxg}^*(t_0)$	0.4894	1.24e-06	0.0020	0.4883	0.4903	0.0020
			$h_{1EGxg}^*(t_0)$	0.2054	4.03e-07	0.0001	0.2036	0.2063	0.0026
		LINEX	$R_{EGxg}^*(t_0)$	0.4881	8.47e-07	0.0007	0.4859	0.4891	0.0032
			$h_{1EGxg}^*(t_0)$	0.2040	3.84e-06	0.0063	0.2013	0.2061	0.0049
9	0.5	SE	$R_{EGxg}^*(t_0)$	0.6242	1.54e-07	7.39e-05	0.6234	0.6249	0.0015
			$h_{1EGxg}^*(t_0)$	0.3326	1.00e-06	2.59e-03	0.3316	0.3332	0.0016
		LINEX	$R_{EGxg}^*(t_0)$	0.6235	5.09e-07	0.0010	0.6228	0.6240	0.0011
			$h_{1EGxg}^*(t_0)$	0.3337	3.63e-07	0.0009	0.3326	0.3345	0.0019
	1	SE	$R_{EGxg}^*(t_0)$	0.5442	3.40e-07	0.0006	0.5430	0.5451	0.0021
			$h_{1EGxg}^*(t_0)$	0.2351	8.02e-07	0.0019	0.2339	0.2367	0.0027
		LINEX	$R_{EGxg}^*(t_0)$	0.5449	3.96e-07	0.0006	0.5436	0.5458	0.0022
			$h_{1EGxg}^*(t_0)$	0.2347	3.07e-07	0.0002	0.2334	0.2355	0.0021
	1.5	SE	$R_{EGxg}^*(t_0)$	0.4886	2.87e-07	4.35e-04	0.4874	0.4894	0.0019
			$h_{1EGxg}^*(t_0)$	0.2053	8.17e-08	3.29e-05	0.2045	0.2057	0.0012
		LINEX	$R_{EGxg}^*(t_0)$	0.4891	8.92e-07	8.92e-07	0.4878	0.4899	0.0020
			$h_{1EGxg}^*(t_0)$	0.2056	2.64e-07	2.64e-07	0.2048	0.2064	0.0015

Table 5

Bayes estimates of the parameters, rf, hrf and standard error from EG-xg distribution for the real data based on lower records

Rv		Estimators	Estimates	s.e
I	3	SE	$R_{EGxg}^*(t_0)$	0.0017
			$h_{1EGxg}^*(t_0)$	10.282
		LINEX	$R_{EGxg}^*(t_0)$	0.0001
			$h_{1EGxg}^*(t_0)$	10.285
II	7	SE	$R_{EGxg}^*(t_0)$	0.0035
			$h_{1EGxg}^*(t_0)$	6.9606
		LINEX	$R_{EGxg}^*(t_0)$	0.0034
			$h_{1EGxg}^*(t_0)$	6.9607

Table 5

Bayes estimates of the parameters, standard error and estimated risks for the real data based on lower records

Rv		Loss functions	Estimators	Estimates	s.e
I	3	SE	α^*	2.0017	0.0007
			β^*	1.0024	0.0008
			θ^*	6.0042	0.0013
		LINEX	α^*	2.0001	0.0012
			β^*	0.9994	0.0007
			θ^*	5.9986	0.0010
II	7	SE	α^*	1.5015	0.0009
			β^*	3.0014	0.0014
			θ^*	6.0008	0.0005
		LINEX	α^*	1.5011	0.0005
			β^*	3.0024	0.0008
			θ^*	5.9992	0.0004

REFERENCES

- AL-Hussaini, E. K., and Jaheen, Z. F. (1992).** Bayesian prediction bounds for the Burr Type XII failure model. *Communication in Statistics-Theory and methods*, Vol. 24, pp. 1829– 1842.
- AL-Hussaini, E. K. (1999).** Predicting observables from a general class of distributions. *Journal of Statistical Planning and Inference*, pp. 79-91.
- Ali, M. M., Pal, M. and Woo, J. (2007).** Some exponentiated distributions. *The Korean Communications in Statistics*. Vol. 14 No. 1, pp. 93-109.
- Alzaatreh. A, Lee, C. and Famoye, F. (2013).** A new method for generating families of continuous distributions. *Metron*, Vol. 71, pp. 63-79.
- Alizadeh, M., Cordeiro, G. M., Nascimento. D, C. A., Lima, M. S., and Ortega E. M. M. (2017).** Odd-Burr generalized family of distributions with some applications. *Journal of Statistical Computation and Simulation*, Vol. 87, No. 2, pp. 367-389.
- Bhaumik, D. K, Kapur, K. and Gibbons, R. D. (2009).** Testing parameters of a gamma distribution for small samples. *Technometrics*, Vol. 51, pp. 326-334.
- Burkschat, M., Cramer, E. and Kamps, U. (2003).** Dual generalized order statistics. *International Journal of Statistics*, Vol. LX1, No. 1, pp. 13-26.
- Cordeiro, G. M. and de Castro, M. (2011).** New family of generalized distributions. *Journal of Statistical Computation and Simulation*. Vol. 81, No. 7, pp. 883-898.
- Cordeiro, G. M., Ortega, E. M., and da Cunha D. C. (2013).** The exponentiated generalized class of distributions. *Journal of Data Science*. Vol. 11, pp. 1-27.
- De Andrade, T. A. N., Bourguignon, M. and Cordeiro, G. M. (2016).** The exponentiated generalized extended exponential distribution. *Journal of Data Science*. Vol. 14, pp. 393-414.
- Elgarhy, M. and Shawki, A. W. (2017).** Exponentiated SUSHILA distribution. *International Journal of Scientific Engineering and Science*. Vol. 1, No. 7, pp. 9-12.
- Elgarhy, M., Ul Haq, M. A and Ul Ain, Q. (2017).** Exponentiated generalized Kumaraswamy distribution with applications. *Annals of Data Science*. DOI 10.1007/s40745-017-0128-x.
- Eugene, N., Lee, C. and Famoye, F. (2002).** Beta-normal distribution and its applications. *Communications in Statistics Theory and Methods*. Vol. 31, pp. 497-512.
- Gupta, R. C., Gupta, P. L. and Gupta, R. D. (1998).** Modeling failure time data by Lehman alternatives. *Comm. Statist. A. Theory Methods* .Vol. 27, pp. 887–904.

- Lemonte A. J., Barreto-Souza W., Cordeiro G. M. (2013).** The exponentiated Kumaraswamy distribution and its log-transform. *Brazilian Journal of Probability and Statistics*. Vol. 27, No 1, pp. 31-53
- Marshall, A. W. and Olkin, I. (1997):** A new method for adding a parameter to a family of distributions with applications to the exponential and Weibull families. *Biometrika* .Vol. 84, No. 3, pp. 641–652.
- Mustafa, A., EL-Desouky, B. S. and AL-Garash, S. (2016).** The exponentiated generalized flexible Weibull extension distribution. *Fundamental Journal of Mathematics and Mathematical Sciences*. Vol. 6No. 2, pp. 75-98. <http://www.frdint.com/>.
- Nadarajah, S. and Kotz, S. (2006).** The exponentiated type distributions. *Acta Applicandae Mathematicae*. Vol. 92, pp. 97-111.
- Oguntunde, P. E., A. de jumo, A.O and Balogun, O. S. (2014).** Statistical properties of the exponentiated generalized inverted exponential distribution. *Applied Mathematics*. Vol. 4, No. 2, pp. 47-55. <http://journal.sapub.org/am>.
- Rather, A. A. and Subramanian, C. (2018).** Exponentiated Mukherjee-Islam distribution. *Journal of Statistics Applications and Probability*. Vol. 7, No. 2, pp. 357-361.
- Sen, S., Maiti, S. S. and Chandra, S. (2016).** The xgamma Distribution: Statistical Properties and Application. *Journal of Modern Applied Statistical Methods*. Vol. 15, No. 1, pp. 774-788. <http://digitalcommons.wayne.edu/jmasm>.
- Silva, R. B., Barreto-Souza, W., Cordeiro, G. M. (2010).** A new distribution with decreasing, increasing and upside-down bathtub failure rate. *Computational Statistics and Data analysis*. Vol. 54, pp. 935-944.
- Sindi, S. N., AL-Dayian, G. R. and Shahbaz, S. H. (2017).** Inference for exponentiated general class of distributions based on record values. *Pakistan Journal of Statistics and Operation Research*. Vol. XIII, No.3, pp. 575-587.
- Yadav, A. S., Saha, M., Tripathi, H and Kumar, S. (2018).** The exponentiated xgamma distribution: Estimation and its application. <https://www.researchgate.net/publication/328418395>
- Yousof, H. M, Afify. A. Z., Alizadeh, M, Butt. N. S., Hamedan, G.G and Ali, M. M. (2015).** The transmuted exponentiated generalized-G family of distributions. *Pakistan Journal of Statistics and Operation Research*. Vol. 11, No. 4. pp. 441-464. [<http://dx.doi.org/10.18187/pjsor.v11i4.1164>].
- Zografos, K. and Balakrishnan, N. (2009).** On families of beta- and generalized gamma generated distributions and associated inference. *Statistical Methodology*. Vol. 6, pp. 344-362

Discrete Inverted Kumaraswamy Distribution: Properties and Estimation

Hegazy, M. A. , Abd EL-Kader, R. E., AL-Dayian, G. R. and EL-Helbawy, A. A.

*Department of Statistics, Faculty of Commerce, AL-Azhar University
(Girls' Branch), Cairo, Egypt*

Abstract

In this paper, a discrete inverted Kumaraswamy distribution; which is a discrete version of the continuous inverted Kumaraswamy variable, is derived using the general approach of discretization of a continuous distribution. Some important distributional and reliability properties of the discrete inverted Kumaraswamy distribution are obtained. Maximum likelihood and Bayesian approaches are applied to estimate the model parameters. A simulation study is carried out to illustrate the theoretical results. Finally, a real data set is applied.

Keywords

Inverted Kumaraswamy distribution; Discrete lifetime models; Survival and hazard rate functions; Order statistics; Maximum likelihood and Bayesian estimation.

1. Introduction

It is well known that the life length in the real world may be associated with continuous non-negative lifetime distributions, however, it is sometimes difficult to get samples from a continuous distribution in real life. The observed values are discrete because they are usually measured to only a finite number of decimal places and can't really constitute all points in a continuum. Even if the measures are taken on a continuous (ratio or interval) scale, the observations may be recorded in a way making discrete model more appropriate. Therefore, it is reasonable to consider the observations as coming from a discretized distribution generated from the continuous model.

In many practical situations, the reliability data are measured in terms of the numbers of runs, cycles or shocks the device sustains before it fails. For example, the number of times the devices are switched on/off, the lifetime of the switch is a *discrete random variable* (**drv**). Also, the number of voltage fluctuations; which an electrical or electronic item can withstand before its failure, is a **drv**, the life of equipment is measured by the number of completed cycles or the number of times it operated before failure, or the life of weapon is measured by the number of rounds fired prior to failure. Similarly, in survival analysis the *survival function* (sf) may be a function of **drv** that is considered as a discrete version of the analogue *continuous random variable* (**crv**). Such as the length of stay in observation ward; when it is measured by

the number of days, or the survival time that the leukemia patients survived since therapy may be counted by number of days or weeks.

Many discrete distributions are available to model such mentioned situations, for example, the geometric and negative binomial distributions which are the discrete versions for the exponential and gamma distributions respectively, but it is well known that they have monotonic hazard rate functions and thus they are unsuitable for some situations.

On the other hand, there are few discrete distributions which can provide accurate models for both count and times. As Poisson distribution, which is used to model counts but not times. Also the binomial distribution is not considered to be popular model for reliability, failure times, count, etc. Beside that it can be approximated to Poisson distribution under suitable conditions. In addition to that, these discrete distributions only cater to positive integers along with zero, but in some analysis the variable of interest can take either zero, positive or negative values. In many situations the interest may be in the difference of two **drvs** each having integer support $(0, \infty)$. The resulting difference will be another **drv** with integer support $(-\infty, \infty)$, see Chakraborty and Chakravorty (2016). Thus, there is a need to derive appropriate discrete distributions by discretizing the continuous distributions to fit various types of data. In the light of the above contexts, the study of the discretization of continuous is meaningful.

There are several methods to construct discrete distributions from the continuous ones, for example discrete analogue of the Pearson system of continuous distributions, discretizing using the *probability density function* (pdf). The distribution generated using this method may not always have a compact form due to the normalizing constant. Also, discretizing can be by shifting the *cumulative distribution function* (cdf), discretizing using *hazard rate function* (hrf), discretizing using **sf** and two composite method. For a comprehensive review on this topic, see Bracquemond and Gaudoin (2003) and Chakraborty (2015).

Many researchers studied the general approach of discretization of some known continuous distributions for use as lifetime distributions. For example, Nakagawa and Osaki (1975) proposed a discrete Weibull distribution.

Khan *et al.* (1989) discussed two discrete Weibull distributions and they presented a simple method to estimate the parameters for one of them. They compared this method with the method of moments, and they concluded that the estimates appear to have almost similar properties.

Roy (2003) derived a discrete normal distribution and elaborated its application for evaluating the reliability of complex systems as an alternative to simulation method. Roy (2004) proposed a discrete Rayleigh distribution as a particular case of the discrete Weibull. Inusah and Kozubowski (2006) introduced a discrete version of the

Laplace (double exponential) distribution and discussed some of its statistical properties and statistical issues of estimation under the discrete Laplace model. Krishna and Pundir (2009) derived the discrete Burr XII distribution and applied it to fit the reliability in series system and a set of real data. Also, they derived the discrete Pareto distribution as a special case of the discrete Burr distribution. Jazi *et al.* (2010) introduced inverse Weibull distribution and they studied four methods of estimation (the heuristic algorithm, the inverse Weibull probability paper plot, the method of moments and the method of proportions).

Gomez-Deniz and Calderin-Ojeda (2011) obtained the discrete version of Lindley distribution and used it as an alternative to Poisson distribution to model automobile claim frequency data. Nekoukhou and *et al.* (2012) presented a new version of the discrete generalized exponential distribution, which can be viewed as different generalization of the geometric distribution, some of its distributional and moment properties were discussed. Al-Huniti and AL-Dayian (2012) introduced the discrete Burr Type III distribution, they discussed some important properties and estimated the parameters based on the maximum likelihood and Bayesian approaches.

Lekshmi and Sebastian (2014) introduced a skewed generalized discrete Laplace distribution which arises as the difference of two independently distributed count variables, they discussed some properties of the distribution and also illustrated a real data set. Para and Jan (2014) presented a discrete generalized Burr Type XII distribution. Hussain and Ahmad (2014) introduced discrete inverse Rayleigh distribution. Hussain *et al.* (2016) introduced a two parameter discrete Lindley distribution. Alamatsaz *et al.* (2016) derived the discrete generalized Rayleigh distribution. Para and Jan (2016) introduced a discrete three parameter of Burr Type XII and Discrete Lomax distributions. Chakraborty and Chakravorty (2016) proposed the discrete logistic distribution and applied it to model real life count data.

Sarhan (2017) introduced a two parameter discrete distribution with bathtub hazard shape, he discussed some statistical properties of the distribution. Also, he used three different methods to estimate the parameters and used the bootstrap method to estimate the distributions of these point estimators. Borah and Hazarika (2017) presented discrete Shanker distribution. Hegazy *et al.* (2018) introduced discrete Gompertz distribution.

Migdadi (2014) used Bayesian inference to estimate the scale parameter of discrete Rayleigh distribution based on *squared error* (SE) and general entropy loss functions. This study also involved prediction for the future ordered observation. Kamari *et al.* (2015) studied Bayesian analysis of discrete Burr distribution, they used the Metropolis-Hastings method for numerical parameters estimate with two loss functions, SE and absolute error loss functions.

The rest of the paper is organized as follows: the *discrete inverted Kumaraswamy* (DIKum) distribution is introduced and some statistical properties are given in Section 2. Some relationships between the DIKum distribution and other well-known distributions are provided in Section 3. While, in Section 4, *maximum likelihood* (ML) and Bayesian estimation are derived. Simulation study and results are presented, in Section 5, also, a real data set is analyzed showing applicability of the DIKum distribution.

2. Discretizing a Continuous Distribution

The general approach of discretizing a continuous variable can be used to construct a discrete model by introducing a grouping on the time axis see Roy (2003, 2004). If the **crv** X has the sf, $S(x) = P(X \geq x)$ and times are grouped into unit intervals so that the **drv** of X denoted by $dX = \lfloor X \rfloor$; which is the largest integer less than or equal to x , will have the *probability mass function* (pmf)

$$\begin{aligned} P(dX = x) &= P(x) = P[x \leq X < x + 1] \\ &= S(x) - S(x + 1), \quad x = 0, 1, 2, \dots \end{aligned} \quad (1)$$

The pmf of the **drv**, dX , can be viewed as discrete concentration of pdf of X . So, given any continuous distribution, it is possible to construct corresponding discrete distribution using (1).

One of the advantages of applying this approach of discretizing is that the sf for discrete distributions has the same functional form of the sf for the continuous distributions; as a result, many reliability characteristics and properties remain unchanged. Thus, discretization of a continuous lifetime model according to this approach is an interesting and simple approach to derive a discrete lifetime model corresponding to the continuous one.

2.1 Construction of discrete inverted Kumaraswamy distribution

Gupta *et al.* (1998) introduced two-parameter distribution as a generalization of the standard Pareto of second kind, called the *exponentiated Pareto* (EP) distribution. Also, Abd AL-Fattah *et al.* (2017) derived the *inverted Kumaraswamy* (IKum) distribution using special transformation, which has the same pdf of EP distribution. This distribution is important in a wide range of applications; for example engineering, medical research and lifetime problems. Gupta *et al.* (1998) proved that distribution is effective in analyzing many lifetime data. The EP distribution has failure rates that take decreasing and upside-down bathtub shapes depending on the value of the shape parameter as was done for the *exponentiated Weibull* (EW) distribution by Mudholkar *et al.* (1995). They observed that exponential distribution, generalized exponential distribution, Weibull distribution, beta distribution, Gamma distribution, uniform distribution, exponentiated exponential distribution,

exponentiated Gamma distribution and other distributions can be obtained as special cases of the EP distribution. The pdf of IKum distribution is given by

$$g(x; \alpha, \beta) = \alpha\beta(1+x)^{-(\alpha+1)}(1-(1+x)^{-\alpha})^{\beta-1}, \quad x > 0; \alpha, \beta > 0, \quad (2)$$

where α and β are shape parameters and should be positive.

The corresponding cdf and sf are, respectively, given by

$$G(x; \alpha, \beta) = (1 - (1+x)^{-\alpha})^{\beta}, \quad x > 0; \alpha, \beta > 0, \quad (3)$$

and

$$S(x; \alpha, \beta) = 1 - (1 - (1+x)^{-\alpha})^{\beta}, \quad x > 0; \alpha, \beta > 0. \quad (4)$$

The IKum distribution have a long right tail; compared with other commonly used distributions. Thus it will affect long term reliability predictions, producing optimistic predictions of rare events occurring in the right tail of the distribution compared with other distributions. Also the IKum distribution provides a good fit to several data in literature.

Using (1) dX can be viewed as the discrete analogue to the continuous IKum variable X , and is commonly said to follow DIKum distribution with two parameters α and β , denoted by DIKum(α, β) distribution, where the corresponding pmf of dX can be written as

$$\begin{aligned} p(x) &\equiv p(x; \alpha, \beta) \\ &= (1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta}, \quad x = 0, 1, 2, \dots, \quad \alpha, \beta > 0, \end{aligned} \quad (5)$$

and the cdf, sf and hrf are as follows:

$$F(x; \alpha, \beta) = P(X \leq x) = 1 - S(x) + P(X = x) = (1 - (2+x)^{-\alpha})^{\beta}, \quad x = 0, 1, 2, \dots, \quad (6)$$

$$\begin{aligned} S(x) &\equiv S(x; \alpha, \beta) \\ &= P(X \geq x) = 1 - F(x; \alpha, \beta) + P(X = x) = 1 - (1 - (1+x)^{-\alpha})^{\beta}, \\ &\quad x = 0, 1, 2, \dots, \end{aligned} \quad (7)$$

and

$$h(x; \alpha, \beta) = \frac{P(x)}{S(x)} = \frac{(1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta}}{1 - (1 - (1+x)^{-\alpha})^{\beta}}, \quad x = 0, 1, 2, \dots, \quad \alpha, \beta > 0. \quad (8)$$

There are some problems associated with the definition of $h(x)$, three of the more notable ones are given below:

- $h(x)$ is not additive for series system.
- The cumulative hrf, $H(x) = \sum h(x) \neq -\ln S(x)$.
- $h(x) \leq 1$ and it has the interpretation of a probability.[For more details, see Xie *et al.* (2002) and Lai (2013) and (2014)].

Therefore, it was necessary to find an alternative definition that is consistent with its continuous counterpart. Roy and Gupta (1992) provide an excellent alternative definition of a discrete hrf denoted by $h_1(x)$:

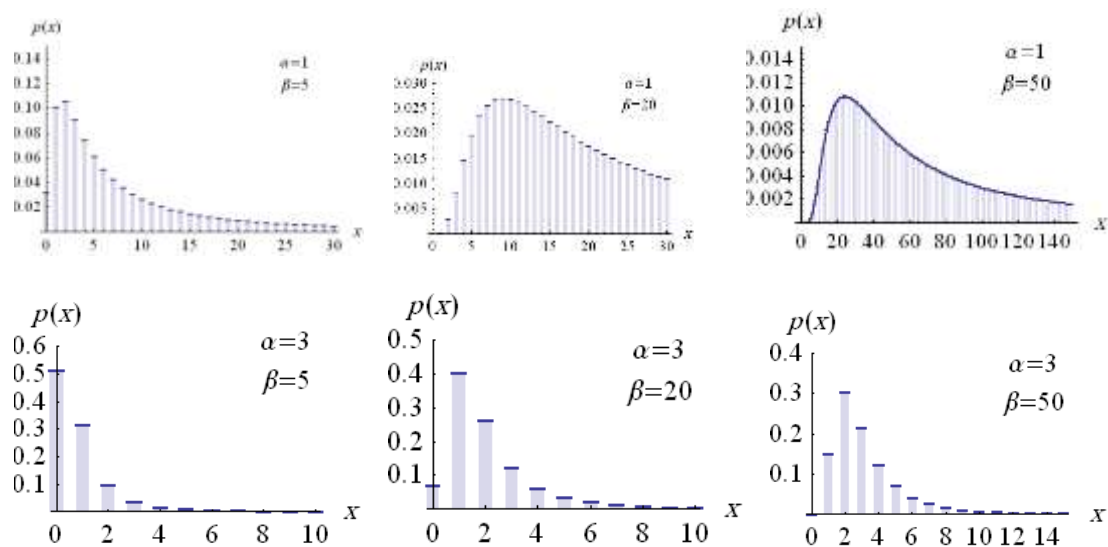
$$h_1(x) = \ln \left[\frac{S(x)}{S(x+1)} \right] = \ln \left[\frac{1 - (1 - (1+x)^{-\alpha})^\beta}{1 - (1 - (2+x)^{-\alpha})^\beta} \right], \quad x = 0, 1, 2, \dots, \alpha, \beta > 0. \quad (9)$$

There is a relationship between $h_1(x)$ and $h(x)$, given by:

$$h(x) = 1 - e^{-h_1(x)}. \quad (10)$$

The two concepts $h(x)$ and $h_1(x)$ have the same monotonic property, i.e., $h_1(x)$ is increasing (decreasing) if and only if $h(x)$ is increasing (decreasing).

Plots of pmf, hrf and *alternative hrf* (ahrf) of DIKum distribution are presented, respectively, in Figures 1 to 3, for some selected values of the parameters.



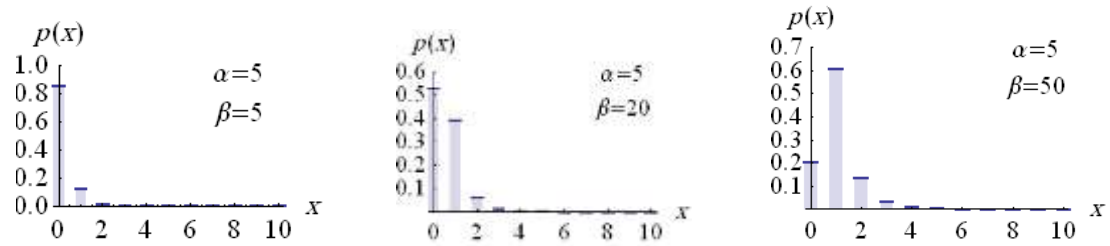


Figure: 1
The plots of the probability mass function

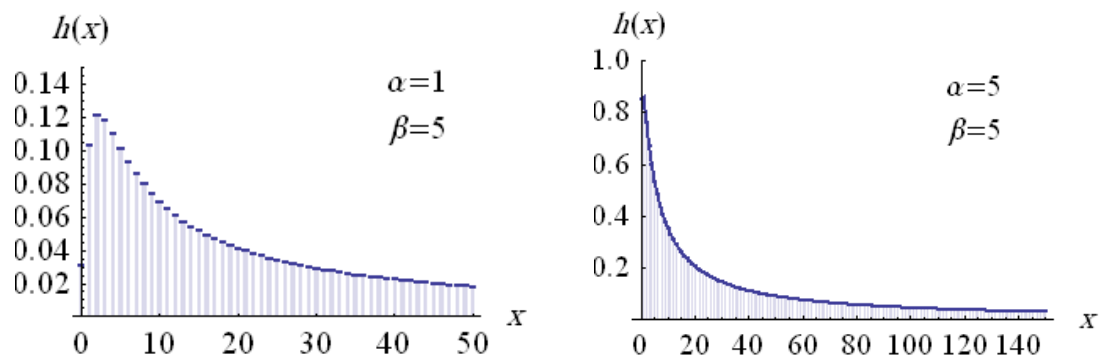


Figure: 2
The plots of the hazard rate function

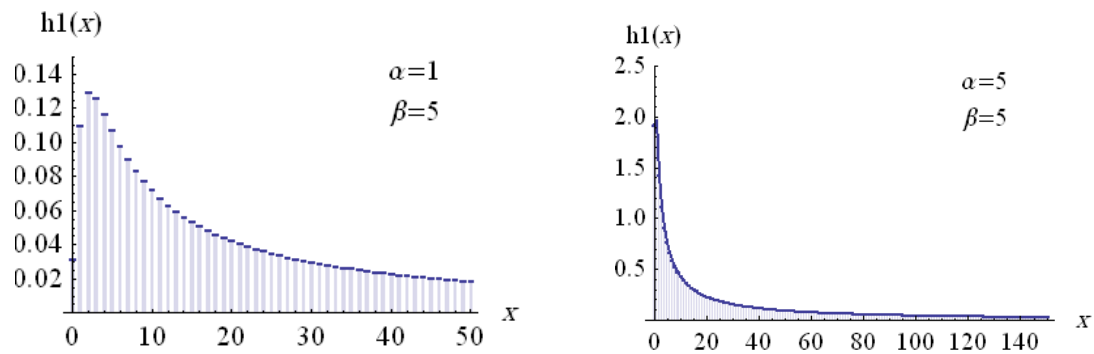


Figure: 3
The plots of the alternative hazard rate function

Figure 1, shows that the pmf of DIKum distribution is unimodal and right skewed according to the selected values of the parameters. For some values of parameters, the pmf is decreasing over $(0, \infty)$ and the mode is at zero. While for other values of the parameters, it indicates that the pmf is increasing on $(0, x_{mode})$ and reaches the maximum at x_{mode} , then decreases to the zero on (x_{mode}, ∞) ; in this case, $x_{mode} = \left\lfloor \left(\frac{\alpha+1}{\alpha\beta+1} \right)^{-\frac{1}{\alpha}} - 1 \right\rfloor$. Plots of pmf, hrf and ahrf show that the DIKum distribution exhibits a long right tail compared with other commonly used distributions. Thus, it will affect long term reliability predictions, producing optimistic predictions of rare events occurring in the right tail of the distribution compared with other distributions. Also the IKum distribution provides a good fit to several data in literature.

Figures 2 and 3 indicate that although the hrf and ahrf of DIKum distribution are decreasing and upside-down bathtub shapes depending on the value of the shape parameters, the hrf is less than 1.

2.2 Some Properties of Discrete Inverted Kumaraswamy Distribution

This subsection is devoted to obtain some distributional properties of DIKum (α, β) distribution, such as the mode, quantiles, r^{th} moments and order statistics.

2.2.1 Mode of discrete inverted Kumaraswamy distribution

The mode of DIKum distribution is at $x_{mode} = \left\lfloor \left(\frac{\alpha+1}{\alpha\beta+1} \right)^{-\frac{1}{\alpha}} - 1 \right\rfloor$, $\alpha, \beta > 0$.

This can be easily verified with pmf plots given in Figure 1.

2.2.2 Quantiles of discrete inverted Kumaraswamy distribution

The u^{th} quantile of a **drv** X , x_u , satisfies

$P(X \leq x_u) \geq u$ and $P(X \geq x_u) \geq 1 - u$, i.e., $F(x_u - 1) < u \leq F(x_u)$. [For more details see Rohatgi and Saleh (2001)].

The u^{th} quantile x_u , of the DIKum (α, β) distribution is given by

$$x_u = \left\lfloor \left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 2 \right\rfloor, \quad 0 < u < 1. \quad (11)$$

where $\lceil x \rceil$ denotes the smallest integer greater than or equal to x and $0 < u < 1$.

Proof

$P(X \leq x_u) \geq u$, from (6) $(1 - (2 + x)^{-\alpha})^\beta \geq u$, hence

$$x_u \geq \left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 2. \quad (12)$$

Similarly, if $P(X \geq x_u) \geq 1 - u$, one obtains

$$x_u \leq \left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 1. \quad (13)$$

Combining (12) and (13), one gets

$$\left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 2 \leq x_u \leq \left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 1.$$

Hence, x_u is an integer value given by

$$x_u = \left\lfloor \left\{ \left(1 - (u)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 2 \right\rfloor. \quad (14)$$

Thus, the median of DIKum (α, β) distribution can be computed from (14) as follows:

$$x_{0.5} = \left\lfloor \left\{ \left(1 - (0.5)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right\} - 2 \right\rfloor. \quad (15)$$

2.2.3 The moments of the inverted Kumaraswamy distribution

This subsection is devoted to obtain the main types of moments: non-central, central and standard moments.

a. The non-central moments of the discrete inverted Kumaraswamy distribution

The non-central moments of DIKum distribution are obtained using (5) as follows:

$$\mu'_r = E(X^r) = \sum_{x=0}^{\infty} x^r p(x) \quad (16)$$

$$= \sum_{x=0}^{\infty} x^r \left[(1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta} \right], \quad r = 1, 2, \dots \quad (17)$$

In particular, the mean (μ) of DIKum distribution is given by

$$\mu_1 = \mu = \sum_{x=0}^{\infty} x \left[(1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta} \right]. \quad (18)$$

b. The central moments of the discrete inverted Kumaraswamy distribution

The central moments can be derived using the relation between the central and non-central moments as given below

$$\mu_r = \sum_{j=0}^r \binom{r}{j} (-1)^j \mu^j \mu_{r-j}, \quad r = 1, 2, \dots, \quad (19)$$

thus, the *variance* (var) of DIKum distribution is

$$\mu_2 = \sum_{x=0}^{\infty} x^2 \left[(1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta} \right]$$

$$-\left[\sum_{x=0}^{\infty} X \left[(1 - (2+x)^{-\alpha})^{\beta} - (1 - (1+x)^{-\alpha})^{\beta} \right]\right]^2. \quad (20)$$

d. The standard moments of the discrete inverted Kumaraswamy distribution

The r^{th} standard moments can be obtained as follows:

$$\alpha_r = E \left(\frac{X-\mu}{\sigma} \right)^r. \quad (21)$$

The skewness and kurtosis of the DIKum distribution are given, respectively, by

$$\alpha_3 = \mu_3/\mu_2^{\frac{3}{2}} \text{ and } \alpha_4 = \mu_4/\mu_2^2, \text{ where } \mu_r = E(X - \mu)^r, \quad r = 1, 2, \dots$$

2.2.4 The order statistic of the discrete inverted Kumaraswamy distribution

Let $F_i(x; \alpha, \beta)$; the cdf of the i^{th} order statistic for a random sample $X_1, X_2, \dots, X_n$, from the DIKum (α, β) , is given by

$$F_i(x; \alpha, \beta) = \sum_{r=i}^n \binom{n}{r} [F(x; \alpha, \beta)]^r [1 - F(x; \alpha, \beta)]^{n-r}. \quad (22)$$

Using the binomial expansion for $[1 - F_i(x; \alpha, \beta)]^{n-r}$ and substituting (6) in (22),

where

$$\begin{aligned} F_i(x; \alpha, \beta) &= \sum_{r=i}^n \binom{n}{r} [F(x; \alpha, \beta)]^r \sum_{j=0}^{n-r} \binom{n-r}{j} (-1)^j [F(x; \alpha, \beta)]^j \\ &= \sum_{r=i}^n \binom{n}{r} \sum_{j=0}^{n-r} \binom{n-r}{j} (-1)^j [(1 - (2+x)^{-\alpha})^\beta]^{r+j}. \end{aligned} \quad (23)$$

Special cases

Case I: If $i=1$ in (23) one can obtain the distribution function of the first order statistic, as given below

$$F_1(x; \alpha, \beta) = 1 - [1 - F(x; \alpha, \beta)]^n = 1 - [1 - (1 - (2+x)^{-\alpha})^\beta]^n, \quad (24)$$

Case II: If $i=n$ in (23) the distribution function of the largest order statistic, as follows:

$$F_n(x; \alpha, \beta) = [F(x; \alpha, \beta)]^n = [(1 - (2+x)^{-\alpha})^\beta]^n. \quad (25)$$

which is the cdf of DIKum $(\alpha, n\beta)$, and the sf of DIKum $(\alpha, n\beta)$ is

$$S(x) = 1 - (1 - (1+x)^{-\alpha})^{n\beta}. \quad (26)$$

Suppose that $X_1, X_2, \dots, X_n$ is a random sample from the DIKum (α, β) distribution. Let $X_{1:n}, X_{2:n}, \dots, X_{n:n}$ denote the corresponding order statistics. (see Arnold *et al.* (2008)). Then, the pmf of $X_{i:n}$, is defined by

$$P(X_{i:n} = x) = \frac{n!}{(i-1)!(n-i)!} \int_{F(x-1)}^{F(x)} v^{i-1} (1-v)^{n-i} dv. \quad (27)$$

Using the binomial expansion for $(1-v)^{n-i}$, then, the pmf in (27) is

$$P(X_{i:n} = x) = \frac{n!}{(i-1)!(n-i)!} \sum_{j=0}^{n-i} \binom{n-i}{j} (-1)^j \left(\frac{1}{1+j}\right) \times \left[\left[(1 - (2+x)^{-\alpha})^\beta \right]^{1+j} - \left[(1 - (1+x)^{-\alpha})^\beta \right]^{1+j} \right]. \quad (28)$$

The pmf of the smallest order statistic is obtained by substituting $i=1$ in (28) as follows:

$$P(X_{1:n} = x) = n \sum_{j=0}^{n-1} \binom{n-1}{j} (-1)^j \left(\frac{1}{1+j}\right) \times \left[\left[(1 - (2+x)^{-\alpha})^\beta \right]^{1+j} - \left[(1 - (1+x)^{-\alpha})^\beta \right]^{1+j} \right], \quad (29)$$

and, the pmf of largest order statistic is obtained by substituting $i=n$ in (28) as follows:

$$P(X_{n:n} = x) = \frac{n}{n+1} \left[\left[(1 - (2+x)^{-\alpha})^\beta \right]^{n+1} - \left[(1 - (1+x)^{-\alpha})^\beta \right]^{n+1} \right]. \quad (30)$$

Also, (23) can be used to obtain the pmf of the DIKum (α, β) distribution.

3. Some Transformations Applied to the Discrete Inverted Kumaraswamy Distribution

In this section relationships between the DIKum distribution and other well-known distributions are provided through using appropriate transformations, which are displayed in Table 1.

Table 1.

Summary of some transformations applied to the inverted Kumaraswamy distribution, discrete inverted Kumaraswamy distribution and the resulting distributions

Distribution	Transformation	Result	Pmf
Inverted Kumaraswamy $X \sim IKum(\alpha, \beta)$	$Y_1 = [X]$	$Y_1 \sim DIKum(\alpha, \beta)$	$p(y_1) = (1 - (2 + y_1)^{-\alpha})^\beta - (1 - (1 + y_1)^{-\alpha})^\beta, \\ y_1 = 0, 1, 2, \dots, \quad \alpha, \beta > 0.$
Discrete inverted Kumaraswamy $X_i' s \sim DIKum(\alpha, \beta)$ $X_i' s (i = 1, 2, 3, \dots, n)$ be iid	$Y_2 = \max_{1 \leq i \leq n} X_i$	$Y_2 \sim DIK \sim (\alpha, n\beta)$	$p(y_2) = (1 - (2 + y_2)^{-\alpha})^{n\beta} - (1 - (1 + y_2)^{-\alpha})^{n\beta}, \\ y_2 = 0, 1, 2, \dots, \quad \alpha, \beta > 0$
Discrete inverted Kumaraswamy $X_i' s \sim DIKum(\alpha, \beta_i)$ $X_i' s$ are independent	$Y_3 = \max_{1 \leq i \leq n} X_i$	$Y_3 \sim DIKum(\alpha, \beta),$ $\beta = \prod_{i=1}^n \beta_i.$	$p(y_3) = (1 - (2 + y_3)^{-\alpha})^\beta - (1 - (1 + y_3)^{-\alpha})^\beta, \\ y_3 = 0, 1, 2, \dots, \quad \alpha, \beta > 0.$

Discrete inverted Kumaraswamy $X \sim DIKum(\theta, 1)$, where $\theta = e^{-\alpha}$	$Y_4 = \ln(2 + x)$	$Y_4 \sim$ geometric distribution with $\theta = e^{-\alpha}$.	$p(y_4) = \theta^{y_4}(1 - \theta)$, $y_4 = 0, 1, 2, \dots$, $0 < \theta < 1$.
Discrete inverted Kumaraswamy $X \sim DIKum(\theta, \beta)$, where $\theta = e^{-\alpha}$	$Y_5 = (\ln(2 + x))^{\frac{1}{\beta}} - 1$	$Y_5 \sim$ discrete generalized Raleigh (θ, β) , $\theta = e^{-\alpha}$.	$p(y_5) = (1 - \theta^{(1+y_5)^2})^\beta - (1 - \theta^{y_5^2})^\beta$, $y_5 = 0, 1, 2, \dots$, $0 < \theta < 1, \beta > 0$.
Discrete inverted Kumaraswamy $X \sim DIKum(\theta, 1)$, where $\theta = e^{-\alpha}$	$Y_6 = (\ln(2 + x))^{\frac{1}{2}} - 1$	$Y_6 \sim$ discrete Raleigh $h(\theta)$, where $\theta = e^{-\alpha}$.	$p(y_6) = (1 - \theta^{(1+y_6)^2}) - (1 - \theta^{y_6^2})$, $y_6 = 0, 1, 2, \dots$, $0 < \theta < 1$.
Discrete inverted Kumaraswamy $X \sim DIKum(\theta, 1)$, where $\theta = e^{-\alpha}$	$Y_7 = x$	$Y_7 \sim$ discrete Pareto (θ)	$p(y_7) = (\theta^{(1+y_7)}) - (\theta^{(2+y_7)})$, $y_7 = 0, 1, 2, \dots$, $0 < \theta < 1$.
Exponential distribution $X \sim exp(\alpha)$	$Y_8 = e^x - 2$	$Y_8 \sim DIKum(\theta, 1)$, $\theta = e^{-\alpha}$	$p(y_8) = (1 - (2 + y_8)^{-\alpha}) - (1 - (1 + y_8)^{-\alpha})$, $y_8 = 0, 1, 2, \dots$, $\alpha > 0$.

4. Estimation of the Parameters of Discrete Inverted Kumaraswamy Distribution

In this section, methods of ML and Bayesian are used to derive the estimators of the parameters for the DIKum distribution.

4.1 Method of maximum likelihood

This subsection is devoted to estimate the vector of two parameters, $\underline{\varphi} = (\alpha, \beta)$, sf, hrf and ahrf of the DIKum (α, β) distribution, based on Type II censored samples, also confidence intervals of the parameters α, β , sf, hrf and ahrf are derived.

Suppose that $X_1, X_2, \dots, X_r$ is a Type II censored sample of size r obtained from a life-test on n items whose lifetimes have a DIKum (α, β) distribution. Then, the likelihood function is

$$L(\underline{\varphi}; \underline{x}) \propto \{\prod_{i=1}^r p(x_{(i)})\} [S(x_{(r)})]^{n-r}, \quad (31)$$

where $p(x)$ and $S(x)$ are given, respectively, by (5) and (7). The $x_{(i)}$'s are ordered times for $i = 1, 2, \dots, r$.

$$L(\underline{\varphi}; \underline{x}) \propto \{\prod_{i=1}^r (1 - (2 + x_i)^{-\alpha})^\beta - (1 - (1 + x_i)^{-\alpha})^\beta\} \times [1 - (1 - (1 + x_r)^{-\alpha})^\beta]^{n-r}. \quad (32)$$

The natural logarithm of the likelihood function is given by

$$\ell \equiv \ln L(\underline{\varphi}; \underline{x}) \propto \ln \prod_{i=1}^r [(1 - (2 + x_i)^{-\alpha})^\beta - (1 - (1 + x_i)^{-\alpha})^\beta] + (n - r) \ln [1 - (1 - (1 + x_r)^{-\alpha})^\beta], \quad (33)$$

$$\ell = \ln L(\underline{\varphi}; \underline{x}) \propto \sum_{i=1}^r \ln [(1 - (2 + x_i)^{-\alpha})^\beta - (1 - (1 + x_i)^{-\alpha})^\beta] + (n - r) \ln [1 - (1 - (1 + x_r)^{-\alpha})^\beta]. \quad (34)$$

Considering the two parameters, α and β are unknown and differentiating the log likelihood function in (34), with respect to α and β , one obtains

$$\frac{\partial \ell}{\partial \alpha} = \sum_{i=1}^r \left\{ \frac{[\beta(1 - (2 + x_i)^{-\alpha})^{\beta-1} (2 + x_i)^{-\alpha} \ln(2 + x_i)] + [\beta(1 - (1 + x_i)^{-\alpha})^{\beta-1} (1 + x_i)^{-\alpha} \ln(1 + x_i)]}{[(1 - (2 + x_i)^{-\alpha})^\beta - (1 - (1 + x_i)^{-\alpha})^\beta]} \right\} + (n - r) \frac{\beta(1 - (1 + x_r)^{-\alpha})^{\beta-1} (1 + x_r)^{-\alpha} \ln(1 + x_r)}{[1 - (1 - (1 + x_r)^{-\alpha})^\beta]}, \quad (35)$$

and

$$\frac{\partial \ell}{\partial \beta} = \sum_{i=1}^r \left\{ \frac{[1 - (2 + x_i)^{-\alpha}]^\beta \ln[1 - (2 + x_i)^{-\alpha}] - [1 - (1 + x_i)^{-\alpha}]^\beta \ln[1 - (1 + x_i)^{-\alpha}]}{[(1 - (2 + x_i)^{-\alpha})^\beta - (1 - (1 + x_i)^{-\alpha})^\beta]} \right\} - (n - r) \frac{\{1 - (1 + x_r)^{-\alpha}\}^\beta \ln\{1 - (1 + x_r)^{-\alpha}\}}{[1 - \{1 - (1 + x_r)^{-\alpha}\}^\beta]}. \quad (36)$$

Then the ML estimators of the parameters, denoted by $\hat{\alpha}$ and $\hat{\beta}$ are derived by equating the two nonlinear likelihood (35) and (36) to zeros and solving numerically.

Depending on the invariance property, the ML estimators of $S(x)$, $h(x)$ and $h_1(x)$ can be obtained by replacing α and β with their corresponding ML estimators $\hat{\alpha}$ and $\hat{\beta}$, respectively, in (7), (8) and (9), as given below

$$\hat{S}_{ML}(x) = 1 - (1 - (1 + x)^{-\hat{\alpha}})^{\hat{\beta}}, \quad x = 1, 2, 3, \dots, \quad (37)$$

$$\hat{h}_{ML}(x) = \frac{(1 - (2 + x)^{-\hat{\alpha}})^{\hat{\beta}} - (1 - (1 + x)^{-\hat{\alpha}})^{\hat{\beta}}}{1 - (1 - (1 + x)^{-\hat{\alpha}})^{\hat{\beta}}}, \quad x = 0, 1, 2, \dots, \quad (38)$$

and

$$\hat{h}_{1ML}(x) = \ln \left[\frac{1 - (1 - (1 + x)^{-\hat{\alpha}})^{\hat{\beta}}}{1 - (1 - (2 + x)^{-\hat{\alpha}})^{\hat{\beta}}} \right], \quad x = 0, 1, 2, \dots. \quad (39)$$

When the sample size is large and the regularity conditions are satisfied, the asymptotic distribution of the ML estimators is

$$\underline{\hat{\varphi}} \sim \text{Bivariate Normal} \left(\underline{\varphi}, I_x^{-1}(\underline{\varphi}) \right), \text{ where } \underline{\varphi} = (\alpha, \beta), \quad \hat{\varphi} = (\hat{\alpha}, \hat{\beta}),$$

and

$I^{-1}(\varphi)$ is the asymptotic variance covariance matrix of the ML estimators α and β , which is the inverse of the observed Fisher information matrix. The asymptotic observed Fisher information matrix can be obtained as follows:

$$I_{\underline{x}}(\hat{\varphi}) \approx \left[\begin{array}{cc} -\left(\frac{\partial^2 \ell}{d\alpha^2}\right) & -\left(\frac{\partial^2 \ell}{\partial \alpha \partial \beta}\right) \\ -\left(\frac{\partial^2 \ell}{\partial \alpha \partial \beta}\right) & -\left(\frac{\partial^2 \ell}{d\beta^2}\right) \end{array} \right]_{(\hat{\alpha}, \hat{\beta})}. \quad (40)$$

The asymptotic $100(1 - \alpha)\%$ confidence interval for α , β , $S_{ML}(x)$, $h_{ML}(x)$ and $h_{1ML}(x)$ are given, respectively, by

$$L_{\omega} = \hat{\omega} - Z_{\frac{\tau}{2}} \sigma_{\hat{\omega}}, \quad \text{and} \quad U_{\omega} = \hat{\omega} + Z_{\frac{\tau}{2}} \sigma_{\hat{\omega}}, \quad (41)$$

where L_{ω} and U_{ω} are the lower and upper bound respectively, $\hat{\omega}$ is $\hat{\alpha}, \hat{\beta}, \hat{S}(x), \hat{h}(x)$ or $\hat{h}_1(x)$, where Z is the $100\% \left(1 - \frac{\tau}{2}\right)$ th standard normal percentile, $(1 - \tau)$ is confidence coefficient and $\sigma_{\hat{\omega}}$ is the standard deviation. [see Lehmann and Casella (1998)].

4.2 Bayesian Estimation

In this subsection, the Bayesian approach is considered, under two types of loss functions, SE and *linear exponential* (LINEX) loss functions to estimate the parameters, sf, hrf and ahrf of the DIKum (α, β) distribution based on Type II censored samples, using informative prior. Also credible intervals for the parameters, sf, hrf and ahrf are obtained.

$L(\underline{\varphi}; \underline{x})$ in (31) can be written as given below:

$$L(\alpha, \beta | \underline{x}) \propto \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r}, \quad (42)$$

where

$$w_{i1} = (1 - (1 + x_i)^{-\alpha})^{\beta}, w_{i2} = (1 - (2 + x_i)^{-\alpha})^{\beta} \\ \text{and} \quad w_r = 1 - (1 - (1 + x_r)^{-\alpha})^{\beta}. \quad (43)$$

Assuming that the parameters α and β of DIKum distribution are random variables with a joint bivariate prior density function that was used by AL-Hussaini and Jaheen (1992) as

$$\pi(\alpha, \beta) = g_1(\alpha | \beta) g_2(\beta), \quad \alpha, \beta > 0, \quad (44)$$

where

$$g_1(\alpha | \beta) = \frac{\beta^a}{\Gamma(a)} \alpha^{a-1} e^{-\beta \alpha}, \quad a, \alpha, \beta > 0, \quad (45)$$

and

$$g_2(\beta) = \frac{b^c}{\Gamma(c)} \beta^{c-1} e^{-b\beta}, \quad \beta, b, c > 0. \quad (46)$$

The joint prior density of α and β will be obtained by substituting (45) and (46) in (44) and it's given by

$$\pi(\alpha, \beta) \propto \alpha^{a-1} \beta^{a+c-1} e^{-\beta(\alpha+b)}, \quad \alpha, \beta, a, b, c > 0. \quad (47)$$

The joint posterior distribution for α and β can be derived using (31) and (47),

hence

$$\pi(\alpha, \beta | \underline{x}) \propto L(\alpha, \beta | \underline{x}) \pi(\alpha, \beta) \quad (48)$$

$$\pi(\alpha, \beta | \underline{x}) = k_1 \alpha^{a-1} \beta^{a+c-1} e^{-\beta(\alpha+b)} \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r}, \quad (49)$$

where

$$k_1^{-1} = \int_0^\infty \beta^{a+c-1} \int_0^\infty \alpha^{a-1} \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha d\beta, \quad (50)$$

which is a normalizing constant.

The marginal posterior distributions $\pi(\alpha | \underline{x})$ and $\pi(\beta | \underline{x})$ are given, respectively, by

$$\pi(\alpha | \underline{x}) = k_1 \alpha^{a-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} \beta^{a+c-1} e^{-\beta(\alpha+b)} d\beta, \quad (51)$$

and

$$\pi(\beta | \underline{x}) = k_1 \beta^{a+c-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha, \quad (52)$$

where k_1^{-1} is a normalizing constant given in (50) and w_{i1} , w_{i2} and w_r are defined in (43).

a. Point estimation

The Bayes point estimators of the parameters, sf, hrf and ahrf are considered based on informative prior and two different loss functions; SE and LINEX loss functions.

I. Bayesian estimation under squared error loss function

Under SE loss function the Bayes estimators of the parameters α and β are given by their marginal posterior expectations using (51) and (52), respectively, as shown below

$$\begin{aligned} \alpha_{(SE)}^* &= E(\alpha | \underline{x}) \\ &= k_1 \int_0^\infty \alpha^a \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} \beta^{a+c-1} e^{-\beta(\alpha+b)} d\beta d\alpha, \end{aligned} \quad (53)$$

and

$$\beta_{(SE)}^* = E(\beta|\underline{x}) = k_1 \int_0^1 \beta^{a+c} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha d\beta. \quad (54)$$

Also, the Bayes estimators of the sf, hrf and ahrf under SE loss function can be obtained using (7), (8), (9) and (49) as follows:

$$\begin{aligned} S_{(SE)}^*(x) &= E(S(x)|\underline{x}) \\ &= 1 - k_1^{-1} \int_0^\infty \beta^{a+c-1} \int_0^\infty (1 - (1+x)^{-\alpha})^\beta \alpha^{a-1} e^{-\beta(\alpha+b)} \\ &\quad \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} d\alpha d\beta, \end{aligned} \quad (55)$$

$$\begin{aligned} h_{(SE)}^*(x) &= E(h(x)|\underline{x}) \\ &= k_1^{-1} \int_0^\infty \beta^{a+c-1} \int_0^\infty \frac{(1-(2+x)^{-\alpha})^\beta - (1-(1+x)^{-\alpha})^\beta}{1-(1-(1+x)^{-\alpha})^\beta} \alpha^{a-1} e^{-\beta(\alpha+b)} \\ &\quad \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} d\alpha d\beta, \end{aligned} \quad (56)$$

and

$$\begin{aligned} h_{1(SE)}^*(x) &= E(h_1(x)|\underline{x}) \\ &= k_1^{-1} \int_0^\infty \beta^{a+c-1} \int_0^\infty \ln \left[\frac{1-(1-(1+x)^{-\alpha})^\beta}{1-(1-(2+x)^{-\alpha})^\beta} \right] \alpha^{a-1} e^{-\beta(\alpha+b)} \\ &\quad \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} d\alpha d\beta. \end{aligned} \quad (57)$$

II. Bayesian estimation under linear exponential loss function

Under the LINEX loss function, the Bayes estimators for the parameters α and β are given, respectively, by

$$\begin{aligned} \alpha_{(LINX)}^* &= \frac{-1}{\vartheta} \ln E(e^{-\vartheta\alpha}|\underline{x}) \\ &= k_1 \frac{-1}{\vartheta} \ln \left[\int_0^\infty \alpha^{a-1} e^{-\vartheta\alpha} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} \beta^{a+c-1} e^{-\beta(\alpha+b)} d\beta d\alpha \right], \end{aligned} \quad (58)$$

$$\begin{aligned} \beta_{(LINX)}^* &= \frac{-1}{\vartheta} \ln E(e^{-\vartheta\beta}|\underline{x}) \\ &= k_1 \frac{-1}{\vartheta} \ln \left[\int_0^\infty \beta^{a+c-1} e^{-\vartheta\beta} \int_0^\infty \alpha^{a-1} \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha d\beta \right], \end{aligned} \quad (59)$$

where $\vartheta \neq 0$.

Similarly, the Bayes estimators of the sf, hrf and ahrf under LINEX loss function can be obtained from (7), (8), (9) and (49) as follows:

$$\begin{aligned} S_{(LINX)}^*(x) &= \frac{-1}{\vartheta} \ln E(e^{-\vartheta S(x)}|\underline{x}) \\ &= 1 + \frac{k_1^{-1}}{\vartheta} \ln \left[\int_0^\infty \beta^{a+c-1} \int_0^\infty e^{-\vartheta(1-(1+x)^{-\alpha})^\beta} \alpha^{a-1} e^{-\beta(\alpha+b)} \right. \\ &\quad \left. \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} d\alpha d\beta \right], \end{aligned} \quad (60)$$

$$\begin{aligned} h_{(LINX)}^*(x) &= \frac{-1}{\vartheta} \ln E(e^{-\vartheta h(x)}|\underline{x}) \\ &= \frac{-1}{\vartheta} k_1^{-1} \ln \left[\int_0^\infty \beta^{a+c-1} \int_0^\infty e^{-\vartheta(1-(1+x)^{-\alpha})^\beta} \alpha^{a-1} e^{-\beta(\alpha+b)} \right. \\ &\quad \left. \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} d\alpha d\beta \right] \end{aligned}$$

$$\times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r}] d\alpha d\beta, \quad (61)$$

and

$$\begin{aligned} h_{1(LNX)}^*(x) &= \frac{-1}{\vartheta} \ln E(e^{-\vartheta h_1(x)} | \underline{x}) \\ &= \frac{-1}{\vartheta} k_1^{-1} \ln \left[\int_0^\infty \beta^{a+c-1} \int_0^\infty \left[\frac{1-(1-(1+x)^{-\alpha})^\beta}{1-(1-(2+x)^{-\alpha})^\beta} \right]^{-\vartheta} \alpha^{a-1} e^{-\beta(\alpha+b)} \right. \\ &\quad \times \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r}] d\alpha d\beta. \end{aligned} \quad (62)$$

To obtain the Bayes estimates of the parameters, sf, hrf and ahrf (53) - (62) should be evaluated numerically.

b. Credible interval for the parameters

In general, a two-sided $100(1-\omega)\%$ credible interval of φ is given by

$$P[L(\underline{x}) < \varphi < U(\underline{x}) | \underline{x}] = \int_{L(\underline{x})}^{U(\underline{x})} \pi(\varphi | \underline{x}) d\varphi = 1 - \omega,$$

where $L(\underline{x})$ and $U(\underline{x})$, are the *lower limit* (LL) and *upper limit* (UL).

Since, the marginal posterior distributions are given by (51) and (52), then a $100(1-\omega)\%$ credible interval for α ; $(L(\underline{x}), U(\underline{x}))$, are given by

$$\begin{aligned} P[\alpha > L(\underline{x}) | \underline{x}] \\ = k_1 \int_{L(\underline{x})}^\infty \alpha^{a-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} \beta^{a+c-1} e^{-\beta(\alpha+b)} d\beta d\alpha = 1 - \frac{\omega}{2}, \end{aligned} \quad (63)$$

and

$$\begin{aligned} P[\alpha > U(\underline{x}) | \underline{x}] \\ = k_1 \int_{U(\underline{x})}^\infty \alpha^{a-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} \beta^{a+c-1} e^{-\beta(\alpha+b)} d\beta d\alpha = \frac{\omega}{2}, \end{aligned} \quad (64)$$

Also, a $100(1-\omega)\%$ credible interval for β is $(L(\underline{x}), U(\underline{x}))$ and can be obtained as follows:

$$\begin{aligned} P[\beta > L(\underline{x}) | \underline{x}] \\ = k_1 \int_{L(\underline{x})}^\infty \beta^{a+c-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha d\beta = 1 - \frac{\omega}{2}, \end{aligned} \quad (65)$$

and

$$\begin{aligned} P[\beta > U(\underline{x}) | \underline{x}] \\ = k_1 \int_{U(\underline{x})}^\infty \beta^{a+c-1} \int_0^\infty \{\prod_{i=1}^r (w_{i2} - w_{i1})\} w_r^{n-r} e^{-\beta(\alpha+b)} d\alpha d\beta = \frac{\omega}{2}. \end{aligned} \quad (66)$$

Furthermore, a $100(1-\omega)\%$ credible interval for $S(x)$ is $(L(\underline{x}), U(\underline{x}))$, where

$$P[L(\underline{x}) < S(x) < U(\underline{x}) | \underline{x}] = \int_{L(\underline{x})}^{U(\underline{x})} \pi(S | \underline{x}) dS = 1 - \omega, \quad (67)$$

where $\pi(S | \underline{x})$ is the posterior distribution of sf and $\pi(S | \underline{x}) = \int_0^\infty \pi(S, Z | \underline{x}) dZ$.

Let $S = S(x) = 1 - (1 - (1+x)^{-\alpha})^\beta$ and $\alpha = z$,

$$\beta = \left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right), \text{ so that } \begin{vmatrix} \frac{\partial \alpha}{\partial z} & \frac{\partial \beta}{\partial z} \\ \frac{\partial \alpha}{\partial s} & \frac{\partial \beta}{\partial s} \end{vmatrix} = \frac{1}{(1-S)\ln(1-(1+x)^{-z})}.$$

The joint posterior distribution of s and z is

$$\pi(s, z | \underline{x}) = \frac{k_1 z^{a-1}}{(1-S)\ln(1-(1+x)^{-z})} \left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)^{a+c-1} e^{-\left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)(z+b)} \times \{\prod_{i=1}^r (w_{i2*} - w_{i1*})\} w_{r*}^{n-r}, \quad 0 < s < 1, z > 0, \quad (68)$$

where

$$w_{i2*} = (1 - (2 + x_i)^{-z})^{\left(\frac{\ln(1-S)}{\ln(1-(1+x_i)^{-z})} \right)}, w_{i1*} = (1 - (1 + x_i)^{-z})^{\left(\frac{\ln(1-S)}{\ln(1-(1+x_i)^{-z})} \right)},$$

and

$$w_{r*} = 1 - (1 - (1 + x_r)^{-z})^{\left(\frac{\ln(1-S)}{\ln(1-(1+x_r)^{-z})} \right)}. \quad (69)$$

Hence, the posterior density function for sf is given by

$$\pi(s | \underline{x}) = k_1 \int_0^\infty \frac{k_1 z^{a-1}}{(1-S)\ln(1-(1+x)^{-z})} \left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)^{a+c-1} e^{-\left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)(z+b)} \times \{\prod_{i=1}^r (w_{i2*} - w_{i1*})\} w_{r*}^{n-r} dz, \quad 0 < s < 1. \quad (70)$$

Then a 100 (1- ω) % credible interval for S is $(L(\underline{x}), U(\underline{x}))$,

$$P[S > L(\underline{x}) | \underline{x}] = k_1 \int_{L(\underline{x})}^1 \int_0^\infty \frac{k_1 z^{a-1}}{(1-S)\ln(1-(1+x)^{-z})} \left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)^{a+c-1} e^{-\left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)(z+b)} \times \{\prod_{i=1}^r (w_{i2*} - w_{i1*})\} w_{r*}^{n-r} dz ds = 1 - \frac{\omega}{2}, \quad (71)$$

and

$$P[S > U(\underline{x}) | \underline{x}] = k_1 \int_{U(\underline{x})}^1 \int_0^\infty \frac{k_1 z^{a-1}}{(1-S)\ln(1-(1+x)^{-z})} \left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)^{a+c-1} e^{-\left(\frac{\ln(1-S)}{\ln(1-(1+x)^{-z})} \right)(z+b)} \times \{\prod_{i=1}^r (w_{i2*} - w_{i1*})\} w_{r*}^{n-r} dz ds = \frac{\omega}{2}. \quad (72)$$

5. Numerical Results

This section aims to investigate the precision of the theoretical results based on simulated and real data, by evaluating *relative absolute biases* (RABs) and *relative errors* (REs).

5.1 Simulation study

In this subsection, a simulation study is conducted to illustrate the performance of the presented ML and Bayes estimates on the basis of generated data from the DIKum distribution. ML and Bayes averages of the parameters, sf, hrf and ahrf based on Type II censoring are computed. Moreover, confidence and credible intervals of the parameters, sf, hrf and ahrf are calculated. All simulation studies are performed using Mathematica 9 and R programming language. The numerical procedures are performed according to the following algorithm.

Step 1: a random sample $X_1, X_2, \dots, X_n$ of sizes ($n=30, 60, 120$) these random samples are generated from DIKum distribution using the following transformation:

$$x_i = \left\lceil \left[\left(1 - (u_i)^{\frac{1}{\beta}} \right)^{-\frac{1}{\alpha}} - 1 \right] - 2 \right\rceil, \quad i = 1, 2, \dots, n \text{ and } u_i \text{ are random sample from uniform } (0,1) \text{ and then taking the ceiling.}$$

Step 2: two different set values of the parameters are selected as,

Set 1 ($\alpha = 3, \beta = 5$) and Set 2 ($\alpha = 5, \beta = 50$).

Step 3: For each model parameters and for each sample size, the ML and Bayes estimates are computed.

Step 4: Steps from 1 to 3 are repeated 5000 times for each sample size and for selected sets of the parameters. Then the averages, RABs, REs and variances of the estimates of the unknown parameters are computed. The RABs, REs, variances of ML and Bayes estimates of the parameters, sf, hrf and ahrf are computed as follows:

$$1) \text{ Averages} = \frac{\sum_{i=1}^{NR} \text{estimator}}{NR},$$

$$2) \text{ RABs (estimate)} = \frac{|\text{bias (estimate)}|}{\text{true value}},$$

$$3) \text{ Relative errors (REs)} = \frac{\text{ER(estimate)}}{\text{true value}},$$

$$4) \text{ Variances (estimate)} = \text{ER(estimate)} - \text{bias}^2(\text{estimate}),$$

$$5) \text{ Estimated risk (estimate)} = \frac{\sum_{i=1}^{NR} (\text{estimate} - \text{true value})^2}{NR}.$$

- Table 2 shows the averages, RABs, REs, variances, sf, hrf and ahrf estimates, also 95% confidence intervals where the population parameter values are

$\alpha = 3, \beta = 5$ under three levels of $\frac{r}{n} \times 100$ percentage of uncensored observations Type II censoring 60%, 80% and 100%.

- Table 3 displays the same computational results, but for different population parameter values $\alpha = 5, \beta = 50$, from the DIKum distribution for different samples of size ($n=30, 60$ and 120) and also level of Type II censoring 60%, 80% and 100% and number of replications, $N = 5000$.
- Tables 4 and 5 present the Bayes averages for the parameters and their RABs, REs and credible intervals under three levels of $\frac{r}{n} \times 100$ percentage of uncensored observations Type II censoring 60%, 80% and 100% for different population parameter values for $\alpha = (3, 5), \beta = (5, 50)$ and replications $NR = 10000$.
- Table 6 displays the Bayes averages and 95% confidence intervals of the sf, hrf and ahrf at $t_0 = 1$, from DIKum distribution based on Type II censoring for different samples sizes and replications $NR = 10000$.
- Tables 7 to 9 represent the ML and Bayes estimates for real data.

5.2 Applications

The main aim of this subsection is to demonstrate how the proposed methods can be used in practice. A real lifetime data sets is analyzed to illustrate the theoretical results in ML and Bayesian inference.

The application is a real data set obtained from Hinkley (1977). It consists of thirty successive values of March precipitation (in inches) in Minneapolis/St Paul. The data is 0.77, 1.74, 0.81, 1.20, 1.95, 1.20, 0.47, 1.43, 3.37, 2.20, 3.00, 3.09, 1.51, 2.10, 0.52, 1.62, 1.31, 0.32, 0.59, 0.81, 2.81, 1.87, 1.18, 1.35, 4.75, 2.48, 0.96, 1.89, 0.90, 2.05.

To check the validity of the fitted model, Kolmogorov-Smirnov goodness of fit test is performed for data set. The p value is 0.799. The p value showed that the model fits the data very well.

Concluding remarks

- I. From Tables 2, 3 and 7 the RABs, variances and REs of the ML estimates of the parameters α and β decrease when the sample size n increases. Also, it is observed that as the level of censoring decreases the RABs, variances and REs of the ML estimates of the parameters, sf, hrf and the ahrf decrease. The lengths of the confidence intervals become narrower as the sample size increases.
- II. It is noticed, from Tables 4 - 6, 8 and 9 that the RABs, REs for the estimates of the parameters, sf, hrf, ahrf and the credible interval lengths of the parameters, sf, hrf and ahrf under LINEX loss function have less values than the corresponding values under the SE loss function.
- III. It is observed that less RABs and REs are, obtained for complete sample sizes, also the results are less than the corresponding for censored samples. Also, results perform better when n and r get larger.
- IV. The Bayes intervals include the estimates (between the LL and UL).

Acknowledgments:

The authors are thankful to the Editor and Referees for their constructive comments that led to improvements on the original manuscript.

References

- Abd AL-Fattah, A. M., EL-Helbawy, A. A. and AL-Dayian, G.R. (2017).** Inverted Kumaraswamy Properts and Estimation. *Pakistan Journal of Statistics*, 33(1): 37-61.
- Alamatsaz, M. H., Dey, S. Dey, T. and Shams Harandi, S. (2016).** Discrete Generalized Rayleigh Distribution. *Pakistan Journal of Statistics*, 32 (1): 1-20.
- AL-Huniti, A. A. and AL-Dayian, G. R. (2012).** Discrete Burr Type III Distribution. *American Journal of Mathematics and Statistics*. 2(5): 145-152.
- AL-Hussaini E K, Jaheen Z F (1992).** Bayesian prediction bounds for the Burr Type XII failure model. *Communication in Statistics-Theory and methods*, 24(7): 1829– 1842.
- Arnold, B., Balakrishnan, N. and Najaraja, H. N. (2008).** *A First Course in Order Statistics*. John-Wiley and Sons, New York.
- Borah, M. and Hazarika, J. (2017).** Discrete Shanker Distribution and Its Derived Distributions. *Biometrics and Biostatistics International Journal*, 5(4).
- Bracquemond, C. and Gaudoin, O. (2003).** A Survey on Discrete Lifetime Distributions. *International Journal of Reliability, Quality and Safety Engineering*, 10(1): 69-98.
- Chakraborty, S. (2015).** Generating Discrete Analogues of Continuous Probability Distributions – A Survey of Methods and Constructions. *Journal of Statistical Distributions and Applications*, 2(6):1-30.
- Chakraborty, S. and Chakravorty, D. (2016).** A new Discrete Probability Distribution with Integer Support on $(-\infty, \infty)$. *Communication in Statistics – Theory and Methods*, 45(2): 492-505.
- Gomez-Deniz, E. and Calderin-Ojeda, E. (2011).** The Discrete Lindley Distribution: Properties and Application. *Journal of Statistics Computation and Simulation*, 81(11): 1405-1416.
- Hegazy, M. A., EL-Helbawy, A. A., Abd EL-Kader, R. E. and AL-Dayian, G.R. (2018).** Discrete Gompertz Distribution: Properties and Estimation. *Journal of Biostatistics and Biometric Applications* 3(3): 307.
- Hinkley, D. (1977).** On quick choice of power transformations. *Journal of the Royal Statistical Society Series C (Applied Statistics)*, 26 (1): 67-69.
- Hussain T. and Ahmad M. (2014).** Discrete Inverse Rayleigh Distribution. *Pakistan Journal of Statistics*, 30(2): 203-222.
- Hussain, T. Aslam, M. and Ahmad, M. (2016).** A Two Parameter Discrete Lindley Distribution. *Revista Colombiana de Estadística*, 39 (1):45-61.
- Inusah, S. and Kozubowski, J. T. (2006).** A Discrete Analogue of the Laplace Distribution. *Journal of Statistics Planning Inference* 136: 1090-1102.

- Jazi, M. A., Lai, C. D. and Alamatsaz, M. H. (2010).** A Discrete Inverse Weibull Distribution and Estimation of its Parameters. *Elsevier Statistical Methodology*, 7 (2):121- 132.
- Kamari H, Bevrani H, Kamary K (2015).** Bayesian Estimation of Discrete Burr Distribution with Two Parameters. *Journal of Statistics and Mathematical*, 1(2): 62-68.
- Khan, M. S. A., Khalique, A. and Abouammoh, A. M. (1989).** On Estimating Parameters in a Discrete Weibull Distribution. *IEEE Transactions on Reliability*, 38(3): 348-350.
- Krishna, H. and Pundir, P. S. (2009).** Discrete Burr and Discrete Pareto Distributions. *Statistical Methodology*, 6: 177-188.
- Lai, C. D. (2013).** Issues Concerning Constructions of Discrete Lifetime Models. *Quality Technology and Quantitative Management*, 10(2): 251-262.
- Lai, C. D. (2014).** *Generalized Weibull Distribution*. Springer Heidelberg, New York, Dordrecht, London.
- Lehmann, E. L. and Casella, G. (1998).** *Theory of Point Estimation*, John-Wiley and Sons, New York.
- Lekshmi, S. and Sebastian, S. (2014).** A Skewed Generalized Discrete Laplace Distribution. *International Journal of Mathematics and Statistics Invention*, 2: 95-102.
- Migdadi, H. S. (2014).** Bayesian Inference for the Scale parameter of the Discrete Rayleigh Distribution. *MAGNT Research Report*, 3 (2), 1073-1080.
- Nakagawa, T. and Osaki, S. (1975).** The Discrete Weibull Distribution. *IEEE Transactions on Reliability*, 24 (5): 300-301.
- Nekoukhrou, V., Alamatsaz, M. H., and Bidram, H. (2012).** A Discrete Analog of Generalized Exponential Distribution. *Communication in Statistics-Theory and Methods*, 41 (11): 2000-2013.
- Para, B.A. and Jan, T. R. (2014).** Discrete Generalized Burr-Type XII Distribution. *Journal of Modern Applied Statistical Method*, 13 (2):244-258.
- Para, B. A. and Jan, T. R. (2016).** On Discrete Tree Parameter Burr Type XII and Discrete Lomax Distributions and Their Applications to Model count Data from Medical Science. *Biometrics and Biostatistics International Journal*, 4(2): 1-15.
- Rohatgi, V. K. and Saleh, E. A. K. (2001).** *An Introduction to Probability and Statistics*. 2nd Edition, John- Wiley and Sons, New York.
- Roy, D. (2003).** Discrete Normal Distribution. *Communication in Statistics-Theory and Methods*, 32 (10): 1871-1883.
- Roy, D. (2004).** Discrete Rayleigh Distribution. *IEEE Transactions on Reliability*, 53(2): 255-260.
- Roy, D. and Gupta, R. P. (1992).** Classifications of Discrete Lives. *Microelectronics and Reliability*, R-34(3): 253-255.

Sarhan, A. M. (2017). A two-Parameter Discrete Distribution with a Bathtub Hazard Shape. *Communications for Statistical Applications and Methods*, 24 (1): 15–27.

Xie, M., Gaudoin, O. and Bracquemond, C. (2002). Redefining Failure Rate Function for Discrete Distribution. *International Journal of Reliability, Quality and Safety Engineering*, 9(3): 275-286.

Table 2

Averages relative absolute biases, relative errors, variances of ML estimates, 95%confidence intervals of the parameters, survivals, hazard rate and alternative hazard rate functions from IKum distribution for different sample sizes n, censoring size r and the replications NR= 5000 ($\alpha = 3, \beta = 5$)

n	R	Parameter s	Average s	RABs	REs	Variance	UL	LL	Length
30	18	α	2.4058	0.1980	0.2090	0.0401	2.7984	2.0132	0.7852
		β	5.3932	0.0849	0.2755	0.7954	7.1413	3.6451	3.4962
		$R(t_0)$	0.6675	0.3703	0.4109	0.0075	0.8374	0.4976	0.3398
		$h(t_0)$	0.5129	0.2072	0.2315	0.0045	0.6437	0.3821	0.2617
		$h_1(t_0)$	0.7275	0.3012	0.3241	0.0155	0.9714	0.4837	0.4877
	24	α	2.4147	0.1951	0.2049	0.0351	2.7819	2.0475	0.7343
		β	5.1902	0.0786	0.1949	0.7182	6.6346	3.7459	2.8886
		$R(t_0)$	0.6575	0.3500	0.3859	0.0062	0.8128	0.5023	0.3105
		$h(t_0)$	0.5158	0.2026	0.2276	0.0045	0.6473	0.3924	0.2531
		$h_1(t_0)$	0.7334	0.2955	0.3180	0.0149	0.9731	0.4937	0.4793
	30	α	2.5302	0.1589	0.1593	0.0020	2.6198	2.4406	0.1791
		β	5.4246	0.0380	0.1522	0.5430	6.9938	4.8553	2.1384
		$R(t_0)$	0.6337	0.3011	0.3300	0.0043	0.7626	0.5049	0.2576
		$h(t_0)$	0.5446	0.1580	0.1607	0.0003	0.5815	0.5078	0.0736
		$h_1(t_0)$	0.7875	0.2435	0.2466	0.0016	0.8666	0.7084	0.1581
60	36	α	2.4268	0.1911	0.1983	0.0252	2.7380	2.1156	0.6223
		β	5.2371	0.0474	0.1952	0.483	6.2950	4.1792	2.1158
		$R(t_0)$	0.6561	0.347	0.3771	0.0052	0.7970	0.5152	0.2818
		$h(t_0)$	0.5205	0.1955	0.2150	0.0019	0.6339	0.4070	0.2269
		$h_1(t_0)$	0.7411	0.2882	0.3057	0.0113	0.9491	0.5330	0.4162
	48	α	2.4386	0.1871	0.1920	0.0166	2.6915	2.1857	0.5057
		β	5.0661	0.0247	0.1179	0.2913	5.7275	4.4048	1.3227
		$R(t_0)$	0.6425	0.3191	0.3381	0.0029	0.7492	0.5358	0.2134
		$h(t_0)$	0.5286	0.1829	0.1955	0.0009	0.6162	0.44105	0.1751
		$h_1(t_0)$	0.7558	0.2741	0.2852	0.0067	0.9164	0.5951	0.3212
	60	α	2.5232	0.1602	0.1604	0.0010	2.5879	2.4583	0.1296
		β	5.0658	0.0131	0.0688	0.1138	5.9745	5.1571	0.8173
		$R(t_0)$	0.6147	0.2620	0.2809	0.0024	0.7114	0.5180	0.1934
		$h(t_0)$	0.5506	0.1489	0.1505	0.0002	0.5788	0.5223	0.0564
		$h_1(t_0)$	0.8004	0.2312	0.2331	0.0009	0.8610	0.7396	0.1214
120	72	α	2.4542	0.1819	0.1852	0.0107	2.6566	2.2517	0.4049
		β	5.1233	0.0322	0.1202	0.0915	5.7162	4.5304	1.1858
		$R(t_0)$	0.6426	0.3193	0.3326	0.0021	0.7316	0.5537	0.1779
		$h(t_0)$	0.5315	0.1784	0.1872	0.0013	0.6033	0.4597	0.1436
		$h_1(t_0)$	0.7608	0.2693	0.2770	0.0046	0.8933	0.6283	0.2651
	96	α	2.4543	0.1819	0.1832	0.0041	2.5799	2.3287	0.2512
		β	5.0272	0.0132	0.0653	0.0335	5.3436	4.7107	0.6330
		$R(t_0)$	0.6365	0.3068	0.3122	0.0008	0.6919	0.5751	0.1168
		$h(t_0)$	0.5346	0.1736	0.1766	0.0004	0.5755	0.4937	0.0818
		$h_1(t_0)$	0.7658	0.2645	0.2674	0.0016	0.8453	0.6862	0.1590
	120	α	2.5193	0.1565	0.1573	0.0005	2.5625	2.4761	0.0865
		β	5.8389	0.0054	0.0327	0.0261	5.9742	5.7036	0.2706
		$R(t_0)$	0.6025	0.2369	0.2447	0.0008	0.6609	0.5502	0.1107
		$h(t_0)$	0.5545	0.1429	0.1436	8.32×10^{-5}	0.5724	0.5366	0.0358
		$h_1(t_0)$	0.8088	0.2232	0.2240	0.0004	0.8474	0.7701	0.0773

Table 3: Averages relative absolute biases, relative errors, variances of ML estimates, 95%confidence intervals of the parameters, survival, hazard rate and the alternative hazard rate functions from DIKum distribution for different sample sizes n, censoring size r and the replications NR= 5000 ($\alpha = 5, \beta = 50$)

n	R	Parameters	Averages	RABs	Res	Variance	UL	LL	Length
30	18	α	4.3601	0.1320	0.1329	0.0100	4.5564	4.1638	0.3925
		β	51.5789	0.0316	0.0210	0.9232	52.4621	49.6957	2.7664
		$R(t_0)$	0.9228	0.1613	0.1617	0.0001	0.9461	0.8994	0.0467
		$h(t_0)$	0.6202	0.1934	0.1966	0.0013	0.6908	0.5497	0.1412
		$h_1(t_0)$	0.9717	0.3363	0.3386	0.0060	1.1238	0.8195	0.3044
	24	α	4.3934	0.1280	0.1295	0.0071	4.5586	4.2281	0.3305
		β	51.4992	0.0300	0.0370	0.7023	52.1417	49.8566	2.2852
		$R(t_0)$	0.9181	0.1599	0.1606	0.0001	0.9387	0.8974	0.0414
		$h(t_0)$	0.6300	0.1901	0.1958	0.0009	0.6881	0.5719	0.1163
		$h_1(t_0)$	0.9967	0.3306	0.3349	0.0043	1.1246	0.8689	0.2557
	30	α	4.3401	0.1213	0.1225	0.0062	4.4944	4.1858	0.3086
		β	51.0067	0.0201	0.0343	0.0930	51.6046	50.4089	1.1957
		$R(t_0)$	0.9239	0.1540	0.1546	7.21×10^{-5}	0.9405	0.9073	0.0333
		$h(t_0)$	0.6177	0.1773	0.1815	0.0007	0.6705	0.5648	0.1057
		$h_1(t_0)$	0.9634	0.3134	0.3166	0.0033	1.0755	0.8512	0.2242
60	36	α	4.3838	0.1278	0.1280	0.0022	4.4751	4.2925	0.1827
		β	51.4433	0.0289	0.0191	0.2451	52.4136	50.4730	1.9406
		$R(t_0)$	0.9198	0.1579	0.1580	4.52×10^{-5}	0.9329	0.9066	0.0263
		$h(t_0)$	0.6282	0.1849	0.1854	0.0002	0.6552	0.6012	0.0541
		$h_1(t_0)$	0.9900	0.3255	0.3260	0.0012	1.0587	0.9214	0.1372
	48	α	4.4162	0.1232	0.1236	0.0016	4.4948	4.3375	0.1573
		β	51.3599	0.0272	0.0305	0.0909	51.9508	50.7690	1.1818
		$R(t_0)$	0.9147	0.1561	0.1564	3.82×10^{-5}	0.9269	0.9026	0.0242
		$h(t_0)$	0.6373	0.1797	0.1806	0.0001	0.6593	0.6154	0.0439
		$h_1(t_0)$	1.0147	0.3180	0.3189	0.0009	1.0741	0.9553	0.1188
	60	α	4.3611	0.1168	0.1170	0.0012	4.4278	4.2944	0.13336
		β	50.9473	0.0189	0.0279	0.0129	51.1703	50.7243	0.4460
		$R(t_0)$	0.9212	0.1498	0.1500	1.92×10^{-5}	0.9297	0.9126	0.0172
		$h(t_0)$	0.6242	0.1678	0.1684	0.0001	0.6446	0.6038	0.0408
		$h_1(t_0)$	0.9791	0.3010	0.3017	0.0006	1.0277	0.9304	0.0973
120	72	α	4.3959	0.1259	0.1259	0.0009	4.4555	4.3364	0.1191
		β	51.4420	0.0288	0.0301	0.1853	52.2857	50.5984	1.6873
		$R(t_0)$	0.9181	0.1561	0.1562	2.07×10^{-5}	0.9270	0.9092	0.0178
		$h(t_0)$	0.6316	0.1812	0.1813	6.94×10^{-5}	0.6480	0.6153	0.0327
		$h_1(t_0)$	0.9989	0.3205	0.3206	0.0005	1.0428	0.9550	0.0877
	96	α	4.4303	0.1208	0.1210	0.0008	4.4854	4.3751	0.1103
		β	51.3119	0.0262	0.0265	0.0396	51.7021	50.9218	0.7803
		$R(t_0)$	0.9125	0.1541	0.1542	1.95×10^{-5}	0.9212	0.9039	0.0173
		$h(t_0)$	0.6413	0.1752	0.1755	5.92×10^{-5}	0.6564	0.6262	0.0302
		$h_1(t_0)$	1.0255	0.3118	0.3122	0.0004	1.0670	0.9840	0.0830
	120	α	4.371	0.1139	0.1141	0.0003	4.4060	4.3353	0.0707
		β	50.9235	0.0185	0.0185	0.0032	51.0341	50.8128	0.2214
		$R(t_0)$	0.9198	0.1470	0.1471	6.4220×10^{-5}	0.9247	0.9148	0.0099
		$h(t_0)$	0.6270	0.1626	0.1629	2.4319×10^{-5}	0.6367	0.6174	0.0193
		$h_1(t_0)$	0.9863	0.2935	0.2939	0.0002	1.0120	0.9607	0.0513

Table 4: Average, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II Censoring of DIKum ($NR = 10000, t_0 = 10, \alpha = 3, \beta = 5$)

n	r	Par	SE						LINEX($v = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	α	2.7699	0.0767	0.0114	2.9931	2.5890	0.4041	2.9710	0.0097	2.64×10^{-4}	3.0140	2.8723	0.1417
		β	4.7736	0.0452	0.0064	4.9676	4.5504	0.4171	5.0615	0.0156	0.0017	5.2274	4.8453	0.3821
	80% 24	α	3.1271	0.0424	0.0036	3.2451	2.9532	0.2919	2.9842	0.0053	2.54×10^{-4}	3.0483	2.9167	0.1316
		β	5.0582	0.0116	0.0015	5.2373	4.8931	0.3442	4.9219	0.0123	2.35×10^{-4}	5.0022	4.7960	0.2062
	100% 30	α	2.9883	0.0039	2.96×10^{-5}	2.9988	2.9750	0.0238	3.0041	0.0014	4.80×10^{-6}	3.0096	2.9968	0.0127
		β	5.0119	0.0024	2.50×10^{-5}	5.0278	4.9935	0.0343	5.0064	0.0013	4.89×10^{-6}	5.0103	4.9996	0.0107
60	60% 36	α	2.9934	0.0022	1.36×10^{-5}	3.0030	2.9790	0.0240	3.0019	6.65×10^{-4}	3.88×10^{-6}	3.0086	2.9944	0.0141
		β	4.9994	0.0027	3.34×10^{-5}	5.0179	4.9793	0.0386	4.9926	0.0014	7.38×10^{-6}	5.0003	4.9850	0.0153
	80% 48	α	3.0063	0.0021	9.97×10^{-6}	3.0134	2.9955	0.0179	3.0012	3.96×10^{-4}	1.48×10^{-6}	3.0077	2.9964	0.0113
		β	4.9863	3.17×10^{-4}	1.72×10^{-5}	5.0024	4.9678	0.0345	4.9992	1.59×10^{-4}	1.22×10^{-6}	5.0052	4.9926	0.0126
	100% 60	α	3.0015	5.12×10^{-4}	7.63×10^{-7}	3.0037	2.9990	0.0047	3.0008	2.80×10^{-4}	1.56×10^{-7}	3.0015	3.0000	0.0015
		β	4.9986	2.84×10^{-4}	2.37×10^{-7}	4.9997	4.9974	0.0023	4.9997	6.66×10^{-5}	2.16×10^{-8}	5.0002	4.9991	0.0011
120	60% 72	α	2.9989	3.76×10^{-4}	3.46×10^{-7}	2.9998	2.9968	0.0030	2.9994	1.91×10^{-4}	8.40×10^{-8}	3.0001	2.9986	0.0014
		β	5.0011	2.27×10^{-4}	3.46×10^{-7}	5.0031	4.9990	0.0041	4.9989	2.23×10^{-4}	1.57×10^{-7}	4.9998	4.9974	0.0024
	80% 96	α	3.0002	6.12×10^{-5}	4.66×10^{-8}	3.0012	2.9992	0.0020	2.9999	4.90×10^{-5}	1.64×10^{-8}	3.0003	2.9993	0.0010
		β	5.0009	1.84×10^{-4}	1.32×10^{-7}	5.0018	4.9991	0.0027	4.9995	1.04×10^{-4}	4.98×10^{-8}	5.0003	4.9986	0.0017
	100% 120	α	3.0001	4.39×10^{-5}	3.80×10^{-9}	3.0003	3.0000	0.0003	3.0000	2.60×10^{-5}	1.15×10^{-9}	3.0000	2.9999	0.0001
		β	4.9999	2.05×10^{-5}	1.38×10^{-9}	5.0000	4.9997	0.0003	5.0001	1.68×10^{-5}	1.16×10^{-9}	5.0002	5.0000	0.0002

Table 5: Average, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II Censoring of DIKum ($NR = 10000, t_0 = 10, \alpha = 5, \beta = 50$)

n	r	Par	SE						LINEX($\nu = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	α	3.7607	0.2478	0.1945	4.7620	2.5532	2.2088	5.5311	0.1062	0.0282	5.9715	4.8043	1.1671
		β	45.7949	0.0841	0.2240	49.9432	42.4794	7.4637	48.8031	0.0239	0.01432	49.9723	47.6042	2.3681
	80% 24	α	4.9690	6.20×10^{-3}	5.48×10^{-4}	5.0567	4.7860	0.2707	5.0077	3.53×10^{-3}	5.93×10^{-6}	5.0742	4.9484	0.1258
		β	50.1109	2.21×10^{-3}	1.76×10^{-4}	50.2608	49.9402	0.3206	49.9865	2.70×10^{-4}	1.83×10^{-6}	50.0296	49.9394	0.0902
	100% 30	α	5.0154	3.08×10^{-3}	2.99×10^{-5}	5.0265	4.9987	0.0278	4.9895	2.09×10^{-3}	1.09×10^{-5}	4.9984	4.9789	0.0195
		β	49.9949	1.02×10^{-4}	4.14×10^{-7}	50.0038	49.9856	0.0182	50.0019	3.75×10^{-5}	5.64×10^{-8}	50.0059	49.9972	0.0087
60	60% 36	α	4.9698	6.03×10^{-3}	3.38×10^{-4}	5.0386	4.8380	0.2006	4.9866	2.66×10^{-3}	1.77×10^{-7}	5.0370	4.9306	0.1064
		β	49.9037	1.92×10^{-3}	1.20×10^{-4}	49.9911	49.7770	0.2141	49.9851	2.98×10^{-4}	2.22×10^{-6}	50.0729	49.8602	0.2127
	80% 48	α	4.9893	2.13×10^{-3}	1.61×10^{-5}	4.9988	4.9751	0.0237	5.0055	1.10×10^{-3}	3.03×10^{-6}	5.0123	4.9988	0.0134
		β	49.9838	3.23×10^{-4}	4.97×10^{-6}	50.0080	49.9589	0.0491	49.9974	5.15×10^{-5}	6.66×10^{-8}	50.0021	49.9922	0.0098
	100% 60	α	4.9995	1.01×10^{-4}	3.84×10^{-8}	5.0001	4.9985	0.0016	5.0003	5.23×10^{-5}	1.50×10^{-8}	5.0007	4.9995	0.0012
		β	49.9974	5.20×10^{-5}	7.39×10^{-8}	49.9994	49.9959	0.0035	49.9998	3.44×10^{-6}	3.55×10^{-8}	50.0001	49.9993	0.0007
120	60% 72	α	4.9978	4.32×10^{-4}	4.86×10^{-6}	5.0081	4.9812	0.0269	4.9993	1.47×10^{-4}	6.28×10^{-8}	5.0021	4.9946	0.0074
		β	50.0153	3.06×10^{-4}	3.03×10^{-6}	50.0291	49.9934	0.0357	50.0057	1.13×10^{-4}	3.22×10^{-7}	50.0167	49.9980	0.0187
	80% 96	α	5.0016	3.27×10^{-4}	4.47×10^{-7}	5.0040	4.9994	0.0046	5.0001	1.57×10^{-5}	8.78×10^{-9}	5.0012	4.9992	0.0019
		β	50.0021	4.33×10^{-5}	5.98×10^{-8}	50.0035	49.9997	0.0038	50.0014	2.91×10^{-5}	2.14×10^{-8}	50.0023	49.9999	0.0024
	100% 120	α	4.9999	1.51×10^{-5}	1.66×10^{-9}	5.0001	4.9997	0.0003	5.0000	5.79×10^{-6}	1.27×10^{-10}	5.0000	4.9999	0.0001
		β	49.9999	2.85×10^{-6}	2.48×10^{-8}	50.0000	49.9997	0.0003	49.9999	1.18×10^{-6}	4.53×10^{-9}	50.0000	49.9999	0.0001

Table 6: Average, relative absolute biases, relative errors of the Bayes estimates and 95% credible intervals of the $s(x)$, $hr(x)$ and $ah(x)$ based on Type II censoring of DIKum ($NR = 10000, t_0 = 1, \alpha = 5, \beta = 50$)

n	r	Par	SE						LINEX($v = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	$s(x)$	0.6191	0.2216	0.0268	0.7561	0.3716	0.3845	0.7347	0.0764	3.01×10^{-3}	0.7961	0.6694	0.1267
		$h(x)$	0.5653	0.7727	0.2752	0.7459	0.3704	0.3755	0.6503	0.1507	0.0109	0.7538	0.5545	0.1992
		$ah(x)$	2.3258	1.0747	0.9443	3.0717	1.8260	1.2457	1.8209	0.7269	0.4164	2.3576	1.2653	1.0923
	80% 24	$s(x)$	0.7839	0.0145	3.62×10^{-3}	0.8911	0.6289	0.2621	0.7894	7.72×10^{-3}	2.64×10^{-4}	0.8277	0.7424	0.0853
		$h(x)$	0.6995	0.0864	4.42×10^{-3}	0.7869	0.6105	0.1763	0.7415	0.0316	6.32×10^{-4}	0.7717	0.6967	0.0749
		$ah(x)$	1.5985	0.5736	0.2444	1.7861	1.3879	0.3982	1.5890	0.5671	0.2355	1.6912	1.4473	0.2439
	100% 30	$s(x)$	0.7968	3.97×10^{-3}	1.05×10^{-5}	0.8041	0.7884	0.0157	0.7987	1.59×10^{-3}	9.87×10^{-6}	0.8022	0.7931	0.0091
		$h(x)$	0.7539	1.54×10^{-2}	1.30×10^{-4}	0.7682	0.7408	0.0274	0.7599	7.63×10^{-3}	2.58×10^{-5}	0.7631	0.7511	0.0120
		$ah(x)$	1.4573	0.4764	0.1647	1.4699	1.4463	0.0236	1.4550	0.4749	0.1636	1.4617	1.4441	0.0176
60	60% 36	$s(x)$	0.6970	0.1237	8.19×10^{-3}	0.7878	0.5400	0.2477	0.8543	0.0739	2.64×10^{-3}	0.8963	0.7926	0.1037
		$h(x)$	0.6758	0.2617	0.0359	0.7459	0.6807	0.0652	0.6352	0.1704	0.0164	0.6678	0.6243	0.0435
		$ah(x)$	1.5267	0.5242	0.2005	1.6119	1.3908	0.2210	1.5111	0.5134	0.1925	1.6224	1.4116	0.2108
	80% 48	$s(x)$	0.8019	8.89×10^{-3}	1.66×10^{-4}	0.8250	0.7809	0.0441	0.7969	1.77×10^{-3}	6.48×10^{-6}	0.8005	0.7880	0.0125
		$h(x)$	0.7628	4.06×10^{-3}	1.93×10^{-5}	0.7719	0.7543	0.0176	0.7614	3.81×10^{-3}	1.71×10^{-5}	0.7667	0.7555	0.0112
		$ah(x)$	1.4811	0.4928	0.1763	1.5004	1.4460	0.0544	1.4571	0.4762	0.1646	1.4605	1.4503	0.0102
	100% 60	$s(x)$	0.7963	1.03×10^{-3}	5.68×10^{-7}	0.7971	0.7952	0.0019	0.7962	9.37×10^{-4}	4.27×10^{-7}	0.7956	0.7944	0.0012
		$h(x)$	0.7653	9.60×10^{-4}	6.71×10^{-7}	0.7664	0.7640	0.0024	0.7668	4.81×10^{-4}	3.95×10^{-7}	0.7669	0.7658	0.0011
		$ah(x)$	1.4528	0.4733	0.1626	1.4537	1.4514	0.0023	1.4514	0.4727	0.1622	1.4523	1.4511	0.0012
120	60% 72	$s(x)$	0.8048	0.0116	1.42×10^{-4}	0.8234	0.7831	0.0403	0.7930	3.13×10^{-3}	6.29×10^{-6}	0.7959	0.7885	0.0073
		$h(x)$	0.7756	0.0128	8.77×10^{-5}	0.7869	0.7638	0.0230	0.7669	1.52×10^{-3}	8.04×10^{-6}	0.7729	0.7581	0.0147
		$ah(x)$	1.4668	0.4829	0.1693	1.4850	1.4436	0.0414	1.4569	0.4761	0.1645	1.4633	1.4484	0.0148
	80% 96	$s(x)$	0.7951	9.01×10^{-4}	4.72×10^{-7}	0.7958	0.7937	0.0021	0.7948	4.67×10^{-4}	2.29×10^{-7}	0.7955	0.7939	0.0016
		$h(x)$	0.7671	1.69×10^{-3}	1.75×10^{-6}	0.7689	0.7654	0.0035	0.7662	5.61×10^{-4}	1.55×10^{-7}	0.7666	0.7657	0.0009
		$ah(x)$	1.4531	0.4735	0.1627	1.4538	1.4514	0.0024	1.4521	0.4728	0.1622	1.4526	1.4514	0.0012
	100% 120	$s(x)$	0.7954	1.87×10^{-4}	1.67×10^{-8}	0.7955	0.7953	0.0002	0.7955	8.16×10^{-6}	2.64×10^{-10}	0.7956	0.7955	0.0000
		$h(x)$	0.7657	3.96×10^{-5}	4.83×10^{-9}	0.7658	0.7655	0.0003	0.7658	1.90×10^{-5}	1.21×10^{-9}	0.7659	0.7658	0.0001
		$ah(x)$	1.4515	0.4724	0.1620	1.4516	1.4514	0.0002	1.4516	0.4724	0.1619	1.4516	1.4515	0.0001

Table 7: ML estimates of the parameters, sf, hrf and ahrf and their relative absolute biases and estimated risk for the real data based on Type II censoring

n	r	Parameters	Averages	RAB	ER
30	21	α	4.3155	0.1369	0.4686
		β	51.4347	0.0287	2.0582
		$s(x)$	0.9294	0.16822	0.0179
		$h(x)$	0.6094	0.2042	0.0245
		$h_1(x)$	0.940	0.3523	0.2616
	27	α	4.4383	0.1123	0.3155
		β	50.4820	0.0096	0.2323
		$s(x)$	0.9078	0.1411	0.0126
		$h(x)$	0.6469	0.1553	0.0141
		$h_1(x)$	1.0409	0.2829	0.1687
	30	α	4.4503	0.1099	0.3021
		β	50.3975	0.0079	0.1580
		$s(x)$	0.9055	0.1383	0.0121
		$h(x)$	0.6503	0.1508	0.0133
		$h_1(x)$	1.0506	0.2762	0.1608

Table 8: Estimates, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II censoring for real data ($NR = 10000, \alpha = 5, \beta=50$)

Real Data	n	r	Par	SEL						LINEX($v = 0.1$)					
				Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
Application	29	60% 17	α	5.0107	2.14×10^{-3}	1.45×10^{-5}	5.0170	4.9954	0.0216	4.9990	1.97×10^{-4}	1.09×10^{-6}	5.0039	4.9910	0.01291
			β	50.0294	5.89×10^{-4}	1.17×10^{-5}	50.0516	49.9981	0.0535	50.0041	8.26×10^{-5}	4.36×10^{-7}	50.0109	49.9938	0.0171
		80% 23	α	5.0019	3.87×10^{-4}	4.96×10^{-7}	5.0036	4.9995	0.0040	5.0001	2.66×10^{-5}	1.32×10^{-8}	5.0007	4.9994	0.0012
			β	49.9994	1.04×10^{-5}	5.64×10^{-9}	50.0004	49.9984	0.0020	50.0005	1.02×10^{-5}	4.00×10^{-9}	50.0011	49.9995	0.0015
		100% 29	α	4.9999	1.97×10^{-5}	1.57×10^{-9}	5.0000	4.9997	0.0003	5.0000	3.43×10^{-6}	9.04×10^{-11}	5.0001	5.0000	0.0001
			β	50.0000	5.70×10^{-7}	1.20×10^{-10}	50.0001	49.9997	0.0003	50.0000	1.48×10^{-8}	1.06×10^{-11}	50.0000	49.9999	0.0001

Table 9: Estimates, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the $s(x)$, $h(x)$ and $ah(x)$ of DIKum based on Type II censoring for real data ($NR = 10000, \alpha = 5, \beta = 50$)

Real Data	n	r	Par	SEL						LINEX($v = 0.1$)					
				Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
Application	29	60% 17	$s(x)$	0.4507	0.0175	2.37×10^{-4}	0.4676	0.4282	0.0394	0.4592	1.04×10^{-3}	6.13×10^{-6}	0.4637	0.4552	0.0085
			$h(x)$	0.2605	0.0976	1.30×10^{-3}	0.2718	0.2350	0.0368	0.2347	0.0110	3.10×10^{-5}	0.2389	0.2275	0.0113
			$ah(x)$	0.2776	0.1485	3.06×10^{-3}	0.2905	0.2650	0.0255	0.2739	0.1348	2.48×10^{-3}	0.2797	0.2690	0.0107
		80% 23	$s(x)$	0.4585	4.81×10^{-4}	4.91×10^{-7}	0.4593	0.4568	0.0025	0.4589	2.65×10^{-4}	7.11×10^{-8}	0.4593	0.4584	0.0009
			$h(x)$	0.2352	8.97×10^{-3}	1.16×10^{-5}	0.2371	0.2335	0.0036	0.2365	3.77×10^{-3}	3.09×10^{-6}	0.2377	0.2353	0.0024
			$ah(x)$	0.2726	0.1301	2.29×10^{-3}	0.2733	0.2708	0.0025	0.2713	0.1252	2.12×10^{-3}	0.2716	0.2708	0.0008
		100% 29	$s(x)$	0.4587	4.40×10^{-5}	1.33×10^{-8}	0.4589	0.4585	0.0004	0.4588	1.73×10^{-5}	3.51×10^{-10}	0.4588	0.4587	0.0001
			$h(x)$	0.2374	1.99×10^{-4}	2.97×10^{-8}	0.2376	0.2372	0.0004	0.2373	1.45×10^{-4}	3.98×10^{-9}	0.2374	0.2373	0.0001
			$ah(x)$	0.2710	0.1240	2.17×10^{-3}	0.2711	0.2708	0.0003	0.2710	0.1239	2.08×10^{-3}	0.2711	0.2709	0.0002

Table 4: Average, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II

Censoring of DIKum ($NR = 10000, t_0 = 10, \alpha = 3, \beta = 5$)

n	r	Par	SE						LINEX($v = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	α	2.7699	0.0767	0.0114	2.9931	2.5890	0.4041	2.9710	0.0097	2.64×10^{-4}	3.0140	2.8723	0.1417
		β	4.7736	0.0452	0.0064	4.9676	4.5504	0.4171	5.0615	0.0156	0.0017	5.2274	4.8453	0.3821
	80% 24	α	3.1271	0.0424	0.0036	3.2451	2.9532	0.2919	2.9842	0.0053	2.54×10^{-4}	3.0483	2.9167	0.1316
		β	5.0582	0.0116	0.0015	5.2373	4.8931	0.3442	4.9219	0.0123	2.35×10^{-4}	5.0022	4.7960	0.2062
	100% 30	α	2.9883	0.0039	2.96×10^{-5}	2.9988	2.9750	0.0238	3.0041	0.0014	4.80×10^{-6}	3.0096	2.9968	0.0127
		β	5.0119	0.0024	2.50×10^{-5}	5.0278	4.9935	0.0343	5.0064	0.0013	4.89×10^{-6}	5.0103	4.9996	0.0107
60	60% 36	α	2.9934	0.0022	1.36×10^{-5}	3.0030	2.9790	0.0240	3.0019	6.65×10^{-4}	3.88×10^{-6}	3.0086	2.9944	0.0141
		β	4.9994	0.0027	3.34×10^{-5}	5.0179	4.9793	0.0386	4.9926	0.0014	7.38×10^{-6}	5.0003	4.9850	0.0153
	80% 48	α	3.0063	0.0021	9.97×10^{-6}	3.0134	2.9955	0.0179	3.0012	3.96×10^{-4}	1.48×10^{-6}	3.0077	2.9964	0.0113
		β	4.9863	3.17×10^{-4}	1.72×10^{-5}	5.0024	4.9678	0.0345	4.9992	1.59×10^{-4}	1.22×10^{-6}	5.0052	4.9926	0.0126
	100% 60	α	3.0015	5.12×10^{-4}	7.63×10^{-7}	3.0037	2.9990	0.0047	3.0008	2.80×10^{-4}	1.56×10^{-7}	3.0015	3.0000	0.0015
		β	4.9986	2.84×10^{-4}	2.37×10^{-7}	4.9997	4.9974	0.0023	4.9997	6.66×10^{-5}	2.16×10^{-8}	5.0002	4.9991	0.0011
120	60% 72	α	2.9989	3.76×10^{-4}	3.46×10^{-7}	2.9998	2.9968	0.0030	2.9994	1.91×10^{-4}	8.40×10^{-8}	3.0001	2.9986	0.0014
		β	5.0011	2.27×10^{-4}	3.46×10^{-7}	5.0031	4.9990	0.0041	4.9989	2.23×10^{-4}	1.57×10^{-7}	4.9998	4.9974	0.0024
	80% 96	α	3.0002	6.12×10^{-5}	4.66×10^{-8}	3.0012	2.9992	0.0020	2.9999	4.90×10^{-5}	1.64×10^{-8}	3.0003	2.9993	0.0010
		β	5.0009	1.84×10^{-4}	1.32×10^{-7}	5.0018	4.9991	0.0027	4.9995	1.04×10^{-4}	4.98×10^{-8}	5.0003	4.9986	0.0017
	100% 120	α	3.0001	4.39×10^{-5}	3.80×10^{-9}	3.0003	3.0000	0.0003	3.0000	2.60×10^{-5}	1.15×10^{-9}	3.0000	2.9999	0.0001
		β	4.9999	2.05×10^{-5}	1.38×10^{-9}	5.0000	4.9997	0.0003	5.0001	1.68×10^{-5}	1.16×10^{-9}	5.0002	5.0000	0.0002

Table 5: Average, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II

n	r	Par	SE						LINEX($\nu = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	α	3.7607	0.2478	0.1945	4.7620	2.5532	2.2088	5.5311	0.1062	0.0282	5.9715	4.8043	1.1671
		β	45.7949	0.0841	0.2240	49.9432	42.4794	7.4637	48.8031	0.0239	0.01432	49.9723	47.6042	2.3681
	80% 24	α	4.9690	6.20×10^{-3}	5.48×10^{-4}	5.0567	4.7860	0.2707	5.0077	3.53×10^{-3}	5.93×10^{-6}	5.0742	4.9484	0.1258
		β	50.1109	2.21×10^{-3}	1.76×10^{-4}	50.2608	49.9402	0.3206	49.9865	2.70×10^{-4}	1.83×10^{-6}	50.0296	49.9394	0.0902
	100% 30	α	5.0154	3.08×10^{-3}	2.99×10^{-5}	5.0265	4.9987	0.0278	4.9895	2.09×10^{-3}	1.09×10^{-5}	4.9984	4.9789	0.0195
		β	49.9949	1.02×10^{-4}	4.14×10^{-7}	50.0038	49.9856	0.0182	50.0019	3.75×10^{-5}	5.64×10^{-8}	50.0059	49.9972	0.0087
60	60% 36	α	4.9698	6.03×10^{-3}	3.38×10^{-4}	5.0386	4.8380	0.2006	4.9866	2.66×10^{-3}	1.77×10^{-7}	5.0370	4.9306	0.1064
		β	49.9037	1.92×10^{-3}	1.20×10^{-4}	49.9911	49.7770	0.2141	49.9851	2.98×10^{-4}	2.22×10^{-6}	50.0729	49.8602	0.2127
	80% 48	α	4.9893	2.13×10^{-3}	1.61×10^{-5}	4.9988	4.9751	0.0237	5.0055	1.10×10^{-3}	3.03×10^{-6}	5.0123	4.9988	0.0134
		β	49.9838	3.23×10^{-4}	4.97×10^{-6}	50.0080	49.9589	0.0491	49.9974	5.15×10^{-5}	6.66×10^{-8}	50.0021	49.9922	0.0098
	100% 60	α	4.9995	1.01×10^{-4}	3.84×10^{-8}	5.0001	4.9985	0.0016	5.0003	5.23×10^{-5}	1.50×10^{-8}	5.0007	4.9995	0.0012
		β	49.9974	5.20×10^{-5}	7.39×10^{-8}	49.9994	49.9959	0.0035	49.9998	3.44×10^{-6}	3.55×10^{-8}	50.0001	49.9993	0.0007
120	60% 72	α	4.9978	4.32×10^{-4}	4.86×10^{-6}	5.0081	4.9812	0.0269	4.9993	1.47×10^{-4}	6.28×10^{-8}	5.0021	4.9946	0.0074
		β	50.0153	3.06×10^{-4}	3.03×10^{-6}	50.0291	49.9934	0.0357	50.0057	1.13×10^{-4}	3.22×10^{-7}	50.0167	49.9980	0.0187
	80% 96	α	5.0016	3.27×10^{-4}	4.47×10^{-7}	5.0040	4.9994	0.0046	5.0001	1.57×10^{-5}	8.78×10^{-9}	5.0012	4.9992	0.0019
		β	50.0021	4.33×10^{-5}	5.98×10^{-8}	50.0035	49.9997	0.0038	50.0014	2.91×10^{-5}	2.14×10^{-8}	50.0023	49.9999	0.0024
	100% 120	α	4.9999	1.51×10^{-5}	1.66×10^{-9}	5.0001	4.9997	0.0003	5.0000	5.79×10^{-6}	1.27×10^{-10}	5.0000	4.9999	0.0001
		β	49.9999	2.85×10^{-6}	2.48×10^{-8}	50.0000	49.9997	0.0003	49.9999	1.18×10^{-6}	4.53×10^{-9}	50.0000	49.9999	0.0001

Censoring of DIKum ($NR = 10000, t_0 = 10, \alpha = 5, \beta = 50$)

Table 6: Average, relative absolute biases, relative errors of the Bayes estimates and 95% credible intervals of the $s(x)$, $hr(x)$ and $ah(x)$
based on Type II censoring of DIKum ($NR = 10000, t_0 = 1, \alpha = 5, \beta = 50$)

n	r	Par	SE						LINEX($v = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
30	60% 18	$s(x)$	0.6191	0.2216	0.0268	0.7561	0.3716	0.3845	0.7347	0.0764	3.01×10^{-3}	0.7961	0.6694	0.1267
		$h(x)$	0.5653	0.7727	0.2752	0.7459	0.3704	0.3755	0.6503	0.1507	0.0109	0.7538	0.5545	0.1992
		$ah(x)$	2.3258	1.0747	0.9443	3.0717	1.8260	1.2457	1.8209	0.7269	0.4164	2.3576	1.2653	1.0923
	80% 24	$s(x)$	0.7839	0.0145	3.62×10^{-3}	0.8911	0.6289	0.2621	0.7894	7.72×10^{-3}	2.64×10^{-4}	0.8277	0.7424	0.0853
		$h(x)$	0.6995	0.0864	4.42×10^{-3}	0.7869	0.6105	0.1763	0.7415	0.0316	6.32×10^{-4}	0.7717	0.6967	0.0749
		$ah(x)$	1.5985	0.5736	0.2444	1.7861	1.3879	0.3982	1.5890	0.5671	0.2355	1.6912	1.4473	0.2439
	100% 30	$s(x)$	0.7968	3.97×10^{-3}	1.05×10^{-5}	0.8041	0.7884	0.0157	0.7987	1.59×10^{-3}	9.87×10^{-6}	0.8022	0.7931	0.0091
		$h(x)$	0.7539	1.54×10^{-2}	1.30×10^{-4}	0.7682	0.7408	0.0274	0.7599	7.63×10^{-3}	2.58×10^{-5}	0.7631	0.7511	0.0120
		$ah(x)$	1.4573	0.4764	0.1647	1.4699	1.4463	0.0236	1.4550	0.4749	0.1636	1.4617	1.4441	0.0176
60	60% 36	$s(x)$	0.6970	0.1237	8.19×10^{-3}	0.7878	0.5400	0.2477	0.8543	0.0739	2.64×10^{-3}	0.8963	0.7926	0.1037
		$h(x)$	0.6758	0.2617	0.0359	0.7459	0.6807	0.0652	0.6352	0.1704	0.0164	0.6678	0.6243	0.0435
		$ah(x)$	1.5267	0.5242	0.2005	1.6119	1.3908	0.2210	1.5111	0.5134	0.1925	1.6224	1.4116	0.2108
	80% 48	$s(x)$	0.8019	8.89×10^{-3}	1.66×10^{-4}	0.8250	0.7809	0.0441	0.7969	1.77×10^{-3}	6.48×10^{-6}	0.8005	0.7880	0.0125
		$h(x)$	0.7628	4.06×10^{-3}	1.93×10^{-5}	0.7719	0.7543	0.0176	0.7614	3.81×10^{-3}	1.71×10^{-5}	0.7667	0.7555	0.0112
		$ah(x)$	1.4811	0.4928	0.1763	1.5004	1.4460	0.0544	1.4571	0.4762	0.1646	1.4605	1.4503	0.0102
	100% 60	$s(x)$	0.7963	1.03×10^{-3}	5.68×10^{-7}	0.7971	0.7952	0.0019	0.7962	9.37×10^{-4}	4.27×10^{-7}	0.7956	0.7944	0.0012
		$h(x)$	0.7653	9.60×10^{-4}	6.71×10^{-7}	0.7664	0.7640	0.0024	0.7668	4.81×10^{-4}	3.95×10^{-7}	0.7669	0.7658	0.0011
		$ah(x)$	1.4528	0.4733	0.1626	1.4537	1.4514	0.0023	1.4514	0.4727	0.1622	1.4523	1.4511	0.0012

Table 6 continued

n	r	Par	SE						LINEX($v = 0.1$)					
			Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
120	60% 72	$s(x)$	0.8048	0.0116	1.42×10^{-4}	0.8234	0.7831	0.0403	0.7930	3.13×10^{-3}	6.29×10^{-6}	0.7959	0.7885	0.0073
		$h(x)$	0.7756	0.0128	8.77×10^{-5}	0.7869	0.7638	0.0230	0.7669	1.52×10^{-3}	8.04×10^{-6}	0.7729	0.7581	0.0147
		$ah(x)$	1.4668	0.4829	0.1693	1.4850	1.4436	0.0414	1.4569	0.4761	0.1645	1.4633	1.4484	0.0148
	80% 96	$s(x)$	0.7951	9.01×10^{-4}	4.72×10^{-7}	0.7958	0.7937	0.0021	0.7948	4.67×10^{-4}	2.29×10^{-7}	0.7955	0.7939	0.0016
		$h(x)$	0.7671	1.69×10^{-3}	1.75×10^{-6}	0.7689	0.7654	0.0035	0.7662	5.61×10^{-4}	1.55×10^{-7}	0.7666	0.7657	0.0009
		$ah(x)$	1.4531	0.4735	0.1627	1.4538	1.4514	0.0024	1.4521	0.4728	0.1622	1.4526	1.4514	0.0012
	100% 120	$s(x)$	0.7954	1.87×10^{-4}	1.67×10^{-8}	0.7955	0.7953	0.0002	0.7955	8.16×10^{-6}	2.64×10^{-10}	0.7956	0.7955	0.0000
		$h(x)$	0.7657	3.96×10^{-5}	4.83×10^{-9}	0.7658	0.7655	0.0003	0.7658	1.90×10^{-5}	1.21×10^{-9}	0.7659	0.7658	0.0001
		$ah(x)$	1.4515	0.4724	0.1620	1.4516	1.4514	0.0002	1.4516	0.4724	0.1619	1.4516	1.4515	0.0001

Table 7

ML estimates of the parameters, sf, hrf and ahf and their relative absolute biases and estimated risk for the real data based on Type II censoring

n	r	Parameters	Averages	RAB	ER
30	21	α	4.3155	0.1369	0.4686
		β	51.4347	0.0287	2.0582
		$s(x)$	0.9294	0.16822	0.0179
		$h(x)$	0.6094	0.2042	0.0245
		$h_1(x)$	0.940	0.3523	0.2616
	27	α	4.4383	0.1123	0.3155
		β	50.4820	0.0096	0.2323
		$s(x)$	0.9078	0.1411	0.0126
		$h(x)$	0.6469	0.1553	0.0141
		$h_1(x)$	1.0409	0.2829	0.1687
	30	α	4.4503	0.1099	0.3021
		β	50.3975	0.0079	0.1580
		$s(x)$	0.9055	0.1383	0.0121
		$h(x)$	0.6503	0.1508	0.0133
		$h_1(x)$	1.0506	0.2762	0.1608

Table 8: Estimates, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the parameters λ , and θ based on Type II censoring for real data ($NR = 10000, \alpha = 5, \beta=50$)

Real Data	n	r	Par	SEL						LINEX($v = 0.1$)					
				Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
Application	29	60% 17	α		2.14×10^{-3}						1.97×10^{-4}				
					5.89×10^{-4}						8.26×10^{-5}				
			β	5.0107		1.45×10^{-5}	5.0170	4.9954	0.0216	4.9990		1.09×10^{-6}			
				50.0294		1.17×10^{-5}	50.0516	49.9981	0.0535	50.0041		4.36×10^{-7}	5.0039	4.9910	0.01291
		80% 23	α		3.87×10^{-4}						2.66×10^{-5}				
					1.04×10^{-5}						1.02×10^{-5}				
			β	5.0019		4.96×10^{-7}	5.0036	4.9995	0.0040	5.0001		1.32×10^{-8}			
				49.9994		5.64×10^{-9}	50.0004	49.9984	0.0020	50.0005		4.00×10^{-9}	5.0007	4.9994	0.0012
		100% 29	α												
			β	4.9999		1.57×10^{-9}	5.0000	4.9997	0.0003	5.0000		9.04×10^{-11}	5.0001	5.0000	0.0001
				50.0000		1.20×10^{-10}	50.0001	49.9997	0.0003	50.0000		1.06×10^{-11}	50.0000	49.9999	0.0001

Table 9: Estimates, relative absolute biases, relative error of the Bayes estimates and 95% credible intervals of the $s(x)$, $h(x)$ and $ah(x)$ of DIKum based on Type II censoring for real data ($NR = 10000, \alpha = 5, \beta = 50$)

Real Data	n	r	Par	SEL						LINEX($v = 0.1$)					
				Average	RAB	RE	UL	LL	Length	Average	RAB	RE	UL	LL	Length
Application	29	60% 17	$s(x)$	0.4507	0.0175	2.37×10^{-4}	0.4676	0.4282	0.0394	0.4592	1.04×10^{-3}	6.13×10^{-6}	0.4637	0.4552	0.0085
			$h(x)$	0.2605	0.0976	1.30×10^{-3}	0.2718	0.2350	0.0368	0.2347	0.0110	3.10×10^{-5}	0.2389	0.2275	0.0113
			$ah(x)$	0.2776	0.1485	3.06×10^{-3}	0.2905	0.2650	0.0255	0.2739	0.1348	2.48×10^{-3}	0.2797	0.2690	0.0107
		80% 23	$s(x)$	0.4585	4.81×10^{-4}	4.91×10^{-7}	0.4593	0.4568	0.0025	0.4589	2.65×10^{-4}	7.11×10^{-8}	0.4593	0.4584	0.0009
			$h(x)$	0.2352	8.97×10^{-3}	1.16×10^{-5}	0.2371	0.2335	0.0036	0.2365	3.77×10^{-3}	3.09×10^{-6}	0.2377	0.2353	0.0024
			$ah(x)$	0.2726	0.1301	2.29×10^{-3}	0.2733	0.2708	0.0025	0.2713	0.1252	2.12×10^{-3}	0.2716	0.2708	0.0008
		100% 29	$s(x)$	0.4587	4.40×10^{-5}	1.33×10^{-8}	0.4589	0.4585	0.0004	0.4588	1.73×10^{-5}	3.51×10^{-10}	0.4588	0.4587	0.0001
			$h(x)$	0.2374	1.99×10^{-4}	2.97×10^{-8}	0.2376	0.2372	0.0004	0.2373	1.45×10^{-4}	3.98×10^{-9}	0.2374	0.2373	0.0001
			$ah(x)$	0.2710	0.1240	2.17×10^{-3}	0.2711	0.2708	0.0003	0.2710	0.1239	2.08×10^{-3}	0.2711	0.2709	0.0002

Parameter Estimation of Xgamma Distribution based on Grouped Data

Marwa Abd-Alla Abd-Elghafar¹, Eman Al- Saeed Shehata Sharaf¹

¹Department of Mathematical Statistics, Faculty of Graduate Studies for Statistical Research, Cairo University, Egypt.

Abstract

In many situations, it is often impossible to obtain the measurements of a statistical experiment exactly, but it is possible to classify them into intervals, or disjoint subsets. The resulting data are known as grouped data. This paper presents study the parameter estimation of the xgamma distribution via grouped data. Maximum likelihood, minimum chi-square, modified minimum chi-square, least squares and least lines estimators are derived in equi and unequi-spaced grouping. Numerical study is carried out to compare the performance of different estimates in each case.

Keywords: xgamma distribution; grouped data; maximum likelihood; least squares; minimum chi square.

1. Introduction

Grouped data arise in the case where the observations grouped into intervals, where individual observations are not recorded but only the class intervals in which they fall, and the number of observations falling in each interval are recorded.

Grouped data occur frequently in industry, when variables measurements are difficult or expensive. It may be more economical to gauge observations into groups than to measure their quantities exactly. Stevens (1948) was the first to use of a two-step gauge, which divides observations into three groups. He showed that the grouped data are an excellent alternative to exact measurements, since the small loss in statistical efficiency may be more than offset by savings in the cost of measurements.

Grouped data used for economize on data transmissions and storage, reduce computation, and protect the privacy of individual records, or account for limitations of measurement instrument (see Zhang (1997)).

Several authors discussed the estimation problem and the optimum group limits of the unknown parameters of various distributions from grouped data. The ML estimators for grouped data were first found for equi-spaced grouping by Lindley (1950) utilizing Taylor's expansions. The expansions have been determined and programmed for the location and scale parameters of a number of different distributions. Kulldorff (1961) obtained the maximum likelihood (ML) estimate of the unknown parameters from exponential distribution based on grouped data. Also, he extensively studied the problem of asymptotically optimum grouping for the exponential and normal distributions.

More recently, Hassan et al. (2014) presented estimation of population parameters for the exponentiated inverted Weibull distribution based on grouped data with equi and

unequi-spaced grouping. Several alternative estimation schemes, such as, the method of ML, minimum chi-square (MCS), modified minimum chi-square (MMCS), least squares (LS) and least lines (LL) were considered. The root MSE of resulting estimators was used as comparison criterion to measure both the accuracy and the precision for each parameter. Kim (2016) considered statistical inferences on the estimation of the parameters of a Weibull distribution when data are randomly grouped and censored. ML estimators and approximate ML estimators were derived to estimate the parameters. Kim (2018) studied statistical inferences on the ML estimation of a normal distribution when data are randomly grouped and censored. Nowak (2019) defined the moment-type estimators for distribution parameters under grouped data. These estimators were compared in terms of asymptotic variance to ML estimator's. These estimators were the same as the ML estimator's in the case of exponential families. Moreover, the obtained results were illustrated by examples jointly with simulation study.

The xgamma distribution has been proposed by Sen *et al.* (2016) with the following probability density function (pdf)

$$f(x; \theta) = \frac{\theta^2}{(1+\theta)} \left(1 + \frac{\theta}{2} x^2\right) e^{-\theta x}, \quad x > 0, \theta > 0. \quad (1)$$

Therefore, the cumulative distribution function (cdf) is

$$F(x; \theta) = 1 - \frac{\left(1 + \theta x + \frac{\theta^2}{2} x^2\right)}{(1+\theta)} e^{-\theta x}, \quad x > 0, \theta > 0. \quad (2)$$

This paper is organized as follows: ML estimator of xgamma distribution in case of equi and unequi spaced grouping is provided in Section 2. MCS and MMCS estimators in case of equi and unequi grouped samples are obtained in Section 3. On the other hand, LS and LL estimators in case of equi and unequi-grouped samples are provided in Section 4. Simulation study as well as simulation results in case of equi- and unequi- grouped samples are given in Section 5. Finally, this paper ends with some concluding remarks.

2. Maximum Likelihood Estimator

In this section, the ML estimator of xgamma distribution will be obtained based on unequi- and equi-spaced grouped data.

Suppose that $x_1, x_2, \dots, x_n$ be n observations on a random sample $X_1, X_2, \dots, X_n$ of size n independent and identical drawn from the xgamma distribution, pdf (1) and cdf (2) as defined above.

Let x_0 be the lower bound of the distribution, such that $F(x_0; \theta) = 0$, and let x_k analogously, be the upper bound of the distribution, such that $F(x_k; \theta) = 1$. It will be assumed that the range of variation (x_0, x_k) is divided (partitioned) into k intervals, and let $x_1, x_2, \dots, x_{k-1}$ be the dividing points, such that $x_0 < x_1 < \dots < x_{k-1} < x_k$.

The interval (x_{i-1}, x_i) will be referred to as the i^{th} interval or the i^{th} group, and the points x_i ($i = 0, 1, \dots, k$) will be called group limits. At the first investigation $x_0 > 0$, the

number of failed items, say, $n_1 (\leq n)$ is recorded. If n_1 is equal to n , the experiment ends. However, usually n is much larger compared with n_1 and the experiment moves on.

At the second investigation $x_1 > x_0$, the number of newly failed items n_2 is recorded, and so on. At the last investigation x_{k-1} , the number of failed items n_k is recorded. In addition, let $n_k = n - n_1 - \dots - n_{k-1}$ then a grouped data set is obtained, denoted by $\{(x_i, n_i), i = 1, 2, \dots, k\}$. Here, $(n_1, n_2, \dots, n_k)$ are the numbers of failed items at sampling $(x_0 < x_1 < \dots < x_k)$ with $n_1 + n_2 + \dots + n_k = n$, i.e. the basic feature of the grouped sample is the fact that the only information available is the number of observations falling into one of the k groups, and it will be denoted by n_i such that, $\sum_{i=1}^k n_i = n$.

Then according to Kulldorff (1961), the probability of an observation falling in the i^{th} group, denoted by $p_i \equiv p_i(\theta)$, and the likelihood function in case of grouped samples, denoted by $L(x; \theta)$ respectively is given by

$$p_i \equiv p_i(\theta) = \int_{x_{i-1}}^{x_i} f(x; \theta) dx = F(x_i; \theta) - F(x_{i-1}; \theta), \quad i = 1, 2, \dots, k \quad (3)$$

$$\begin{aligned} L(x; \theta) &= C \prod_{i=1}^k p_i^{n_i} \\ &= C \prod_{i=1}^k [F(x_i; \theta) - F(x_{i-1}; \theta)]^{n_i}, \end{aligned} \quad (4)$$

where,

$$C = \frac{n!}{n_1! n_2! \dots n_k!} \text{ is a constant of proportionality, } 0 \leq n_i \leq n,$$

Depending on equations (2), (3); the p_i of xgamma is given by

$$\begin{aligned} p_i \equiv p_i(\theta) &= \frac{\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2\right)}{(1 + \theta)} e^{-\theta x_{i-1}} - \frac{\left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2\right)}{(1 + \theta)} e^{-\theta x_i} \\ &= A_{i-1} - A_i, \end{aligned} \quad (5)$$

where;

$$\begin{aligned} A_{i-1} &= \left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2\right) \frac{e^{-\theta x_{i-1}}}{(1 + \theta)}, \\ A_i &= \left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2\right) \frac{e^{-\theta x_i}}{(1 + \theta)}, \end{aligned}$$

Moreover, the likelihood function of the xgamma distribution based on grouped data is given by

$$L(x; \theta) = C \prod_{i=1}^k \left[\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2\right) \frac{e^{-\theta x_{i-1}}}{(1 + \theta)} - \left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2\right) \frac{e^{-\theta x_i}}{(1 + \theta)} \right]^{n_i}. \quad (6)$$

Taking the logarithmic of likelihood function (6), then the log-likelihood function, denoted by $\ln L(x; \theta)$ is written as

$$\ln L(x; \theta) = \ln C + \sum_{i=1}^k n_i \ln [A_{i-1} - A_i]. \quad (7)$$

where, A_{i-1} and A_i as defined before.

The first derivative of the log-likelihood function with respect to θ is given by

$$\frac{d \ln L(x; \theta)}{d \theta} = \sum_{i=1}^k \frac{n_i}{A_{i-1} - A_i} \left(\frac{d A_{i-1}}{d \theta} - \frac{d A_i}{d \theta} \right), \quad (8)$$

where,

$$\frac{d A_i}{d \theta} = \frac{-2(2\theta + \theta^2)x_i - \theta^2 x_i^2 - (\theta^2 + \theta^3)x_i^3}{2(1 + \theta)^2} e^{-\theta x_i},$$

and,

$$\frac{d A_{i-1}}{d \theta} = \frac{-2(2\theta + \theta^2)x_{i-1} - \theta^2 x_{i-1}^2 - (\theta^2 + \theta^3)x_{i-1}^3}{2(1 + \theta)^2} e^{-\theta x_{i-1}}.$$

The ML estimator in case of unequi-spaced grouping, denoted by $\hat{\theta}_{UG}$, is determined by setting (8) to be zero. The solution to non-linear equation (8) can be obtained numerically.

As a special case when the distances between the groups limits $x_0, x_1, x_2, \dots, x_{k-1}, x_k$ are equi i.e. $x_i = i\delta$, $\delta > 0$. This status called equi-spaced grouped samples. The log-likelihood function in equi-spaced grouped sample, denoted by $\ln L^*(x; \theta)$, is obtained by setting $x_i = i\delta$ in (7) as follows

$$\ln L^*(x; \theta) = \ln C + \sum_{i=1}^k n_i \ln [A_{i-1}^* - A_i^*], \quad (9)$$

where,

$$A_{i-1}^* = \left(1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2}((i-1)\delta)^2 \right) \frac{e^{-\theta(i-1)\delta}}{1 + \theta},$$

and,

$$A_i^* = \left(1 + \theta + \theta(i\delta) + \frac{\theta^2}{2}(i\delta)^2 \right) \frac{e^{-\theta(i\delta)}}{1 + \theta}.$$

Then the first derivative with respect to θ is obtained as follows

$$\frac{d \ln L^*(x; \theta)}{d \theta} = \sum_{i=1}^k \frac{n_i}{(A_{i-1}^* - A_i^*)} \left(\frac{d A_{i-1}^*}{d \theta} - \frac{d A_i^*}{d \theta} \right), \quad (10)$$

where,

$$\frac{dA_i^*}{d\theta} = \frac{-2(2\theta + \theta^2)(i\delta) - \theta^2(i\delta)^2 - (\theta^2 + \theta^3)(i\delta)^3}{2(1 + \theta)^2} e^{-\theta(i\delta)},$$

and,

$$\frac{dA_{i-1}^*}{d\theta} = \frac{-2(2\theta + \theta^2)(i-1)\delta - \theta^2((i-1)\delta)^2 - (\theta^2 + \theta^3)((i-1)\delta)^3}{2(1 + \theta)^2} e^{-\theta(i-1)\delta}.$$

The ML estimator of θ in the case of equi-spaced grouping, denoted by $\hat{\theta}_{EG}$, is determined by setting (10) to zero. As mentioned above, the solution to non-linear equation (10) is obtained by using numerical method.

3. Minimum Chi-Square and Modified Minimum Chi-Square Estimators

According to MacDonald & Ransom (1979); the MCS estimator of θ in the case of unequi-spaced grouped sampled, denoted by MCS_{UG} , is obtained by minimizing the following quadratic form

$$\chi^2_{UG}(\theta) = \sum_{i=1}^k \frac{D_i}{E_i}, \quad (11)$$

where,

$$D_i = \left(n_i - n \left[\frac{\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2 \right)}{(1 + \theta)} e^{-\theta x_{i-1}} - \frac{\left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2 \right)}{(1 + \theta)} e^{-\theta x_i} \right] \right)^2,$$

and,

$$E_i = n \left[\frac{\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2 \right)}{(1 + \theta)} e^{-\theta x_{i-1}} - \frac{\left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2 \right)}{(1 + \theta)} e^{-\theta x_i} \right],$$

with respect to θ . Furthermore, MCS estimator of θ in the case of equi-spaced grouped sampled, denoted by MCS_{EG} , is obtained by minimizing the following quadratic form

$$\chi^2_{EG}(\theta) = \sum_{i=1}^k \frac{L_i}{M_i}, \quad (12)$$

where,

$$L_i = \left(n_i - n \left[\frac{1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2}((i-1)\delta)^2}{1 + \theta} e^{-\theta(i-1)\delta} - \frac{1 + \theta + \theta(i\delta) + \frac{\theta^2}{2}(i\delta)^2}{1 + \theta} e^{-\theta(i\delta)} \right] \right)^2,$$

and,

$$M_i = n \left[\frac{1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2}((i-1)\delta)^2}{1 + \theta} e^{-\theta(i-1)\delta} - \frac{1 + \theta + \theta(i\delta) + \frac{\theta^2}{2}(i\delta)^2}{1 + \theta} e^{-\theta(i\delta)} \right],$$

with respect to θ . The MCS estimator of θ in case of unequid- and equi-spaced grouped sampled is obtained by solving numerically Equations (11) and (12).

On other hand, the MMCS method is used to obtain MMCS estimator of xgamma distribution in the case of unequid- and equi-spaced grouped sampled.

According to MacDonald & Ransom (1979); the MMCS estimator of θ in the case of unequid-spaced grouped sampled, denoted by $MMCS_{UG}$, is obtained by minimizing the following quadratic form

$$\chi^2_1(\theta_{UG}) = \sum_{i=1}^k \frac{1}{n_i} \left[n_i - n \left\{ \left[\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2 \right) \frac{e^{-\theta x_{i-1}}}{(1 + \theta)} \right] \right. \right. \\ \left. \left. - \left[\left(1 + \theta + \theta x_i + \frac{\theta^2}{2} x_i^2 \right) \frac{e^{-\theta x_i}}{(1 + \theta)} \right] \right\} \right]^2, \quad (13)$$

with respect to θ . Similarly, in case of equi-spaced grouped samples, the MMCS estimator of θ in the case of equi-spaced grouped sampled, denoted by $MMCS_{EG}$, is obtained by minimizing the following quadratic form

$$\chi^2_1(\theta_{EG}) = \sum_{i=1}^k \frac{1}{n_i} \left[n_i - n \left\{ \left[\left(1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2}((i-1)\delta)^2 \right) \frac{e^{-\theta(i-1)\delta}}{1 + \theta} \right] \right. \right. \\ \left. \left. - \left[\left(1 + \theta + \theta(i\delta) + \frac{\theta^2}{2}(i\delta)^2 \right) \frac{e^{-\theta(i\delta)}}{1 + \theta} \right] \right\} \right]^2, \quad (14)$$

with respect to θ . The MMCS estimator of θ in case of unequid- and equi-spaced grouped sampled is obtained by solving numerically Equations (13) and (14).

4. Least Squares and Least Lines Estimators

This section provides the LS and LL estimators of θ for xgamma distribution based on unequid- and equi-spaced grouped data.

Firstly; the LS estimator of θ in the case of unequid-spaced grouped sampled, denoted by LS_{UG} , can be obtained by minimizing the following quadratic form as follows

$$SE_{UG}(\theta, 2) = \sum_{i=1}^k \left\{ \left[\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2 \right) \frac{e^{-\theta x_{i-1}}}{(1+\theta)} \right] - \frac{n_i}{n} \right\}^2, \quad (15)$$

with respect to θ . In the case of equi-spaced grouping, the LS estimator of θ , denoted by LS_{EG} , can be obtained by minimizing the following quadratic form as follows

$$SE_{EG}(\theta, 2) = \sum_{i=1}^k \left\{ \left[\left(1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2} ((i-1)\delta)^2 \right) \frac{e^{-\theta(i-1)\delta}}{1+\theta} \right] - \frac{n_i}{n} \right\}^2, \quad (16)$$

with respect to θ .

Secondly; the LL estimator of θ for xgamma distribution under two different cases, the unequi-spaced grouping and the equi-spaced grouping is obtained.

The LL estimator of θ in the case of unequi-spaced grouped sampled, denoted by LL_{UG} , will be estimated by minimizing the following equation

$$SE_{UG}(\theta, 1) = \sum_{i=1}^k \left\{ \left[\left(1 + \theta + \theta x_{i-1} + \frac{\theta^2}{2} x_{i-1}^2 \right) \frac{e^{-\theta x_{i-1}}}{(1+\theta)} \right] - \frac{n_i}{n} \right\}, \quad (17)$$

with respect to θ . Likewise the LL estimator of θ in the case of equi-spaced grouped sampled, denoted by LL_{EG} , will be estimated by minimizing the following equation

$$SE_{EG}(\theta, 1) = \sum_{i=1}^k \left\{ \left[\left(1 + \theta + \theta(i-1)\delta + \frac{\theta^2}{2} ((i-1)\delta)^2 \right) \frac{e^{-\theta(i-1)\delta}}{1+\theta} \right] - \frac{n_i}{n} \right\}, \quad (18)$$

with respect to θ . The LL and LS estimators of θ in case of unequi- and equi-spaced grouped sampled are obtained numerically by solving Equations (15), (16), (17) and (18).

5. Simulation Study

It is very difficult to compare the theoretical results for different estimators proposed in Sections 2-4. Therefore, a simulation study will be carried out to compare the performances of different estimators, mainly, with respect to their biases and MSEs, for different sample sizes. All simulation studies are done via MathCAD 14. The simulation procedures are done through the two following algorithms.

5.1 Simulation study in case of unequi-spaced grouped samples

In this subsection, a numerical study is carried out to compare the performance of different estimates in case of unequi-spaced grouped samples. The simulation study is designed through the following steps

Step 1: A random data $X_1, X_2, \dots, X_n$ of size $n = 50(50)250$ is generated from the xgamma distribution that by using the following algorithm

- i. Generate $U_i \sim \text{uniform}(0, 1), i = 1, 2, \dots, n$.
- ii. Generate $V_i \sim \text{exponential}(\theta), i = 1, 2, \dots, n$.
- iii. Generate $W_i \sim \text{gamma}(3, \theta), i = 1, 2, \dots, n$.
- iv. If $U_i \leq \frac{\theta}{(1+\theta)}$, then set $X_i = V_i$.
- v. Otherwise, set $X_i = W_i$. (see Sen *et al.* (2016))

Step 2: Different values of the parameter θ are taken as $\theta = 2, 4$ and 6. The number of groups selected as $k = 5(1)10$.

Step 3: Construct a frequency distribution (grouping) from the generated random numbers by entering a pre-specified group limits (not equal), and determine the number of observations falling in each group sample n_i .

Step 4: The frequency distribution obtained in Step 3 is used to determine the probability of an observation falling in the i^{th} group p_i .

Step 5: Depending on the values of n_i, p_i and Equation (8) the ML estimator $\hat{\theta}_{UG}$ in the case of unequi-grouping will be calculated.

Step 6: The MCS and MMCS estimators of θ are calculated by using (11) and (13) for different value of n_i and p_i .

Step 7: The LS and LL estimators are obtained by using Equations (15) and (17) for different value of n_i and p_i .

Step 8: The simulation study is carried out considering $N = 1000$ times for selected values of n and θ .

Step 9: The following two measures are computed for the five methods of estimation and will be compared to decide which method is preferable for estimating the unknown parameter.

Step 10: All the above steps are done via MathCAD 14.

Step 11: Simulation results will be reported in Table 1 and represented in Figures 1-3.

5.2 Numerical results for first algorithm

In this subsection, some observations about the behavior of different estimates are given based on Table 1 and Figures 1-3.

1. The MSEs for all estimates decrease as k increases in almost cases.
2. The MSEs of all estimates are decreasing as n increases in almost cases.
3. The MCS and ML estimation methods yield the smallest sample MSEs in almost cases (see Table 1).
4. The biases decrease as the number of group k increases for all estimates in almost cases.
5. The biases of all estimates are decreasing as n increases in almost all of the cases.
6. Generally, the biases of MCS and ML estimates take the smallest values compared to other estimates while the biases of ML estimate take the largest values compared to other estimates.
7. The biases and MSEs of MMCS estimates are larger than the biases and MSEs associated with other estimates in almost of situations. (see Table 1).

5.3 Simulation study in case of equi-spaced grouped samples

Here, numerical study is performed to compare the behavior of different estimates in case of equi-spaced grouped samples. The numerical study is done through the following steps.

Step 1 and Step 2: will be performed as in the previous algorithm.

Step 3: Construct a frequency distribution (grouping) from the generated random numbers by evaluating equi-spaced group limits, and determine the number of observations falling in each group sample.

Step 4: The frequency distribution obtained in Step 3 is used to determine the probability of an observation falling in the i^{th} group (p_i).

Step 5: Depending on the values of n_i , p_i and equation (9) the ML, $\hat{\theta}_{EG}$, in the case of equi-grouping will be calculated.

Step 6: The MCS and MMCS estimators of θ are evaluated by using Equations (12) and (14) respectively depend on the values of n_i and p_i .

Step 7: Based on Equations (16) and (18), the LS and LL estimators are obtained for different values of n_i and p_i .

Step 8: The simulation study is carried out considering $N = 1000$ times for selected values of n and θ .

Step 9: The biases and MSEs are computed for the five methods of estimation and will be compared to decide which method is preferable for estimating the unknown parameter.

Step 10: All the above steps are executed in Program 2 via MathCAD 14.

Step 11: Simulation results are reported in Table 2 and represented in Figures 4 and 5.

1.4 Numerical results for second algorithm

Here, observations about the behavior of different estimates in case of equi-spaced grouping sampled are provided based on Table 2 and Figures 4 and 5.

1. In the case of equi-spaced grouping; the MSEs of ML estimate take the smallest values compared to LL, LS, MCS and MMCS estimates (see Table 2).
2. The MSEs of MMCS estimate take the largest values compared to other estimates.
3. The MSEs of LL, LS, MMC, MMCS and ML estimate are decreasing as k increases.
4. The MSEs of all estimates are decreasing as n increases in almost cases.
5. The biases of ML estimate have the smallest values compared to biases of other estimates in case of equi-spaces grouped sampling in almost cases.
6. The biases of MMCS estimate take the largest values compared to the biases of other estimates in almost of situations (see Table 2).
7. As k increases the biases of all estimates are decreasing.
8. For all estimates, the biases are decreasing as n increases in almost cases.

Table 1. *The MSEs and Biases of Estimates in Unequi-spaced Grouping at n=50*

K	Method	$\theta=2$		$\theta=4$		$\theta=6$	
		Biases(θ)	MSE(θ)	Biases(θ)	MSE(θ)	Biases(θ)	MSE(θ)
5	LL	-1.6041	2.76064	-2.08742	4.39929	-2.14764	4.64696
	LS	-0.8344	0.70282	-1.25966	1.5884	-1.33664	1.79003
	MCS	-0.84899	0.7408	-1.36132	1.86594	-1.43499	2.07625
	MMCS	-1.88389	3.68653	-2.47417	6.17494	-2.58581	6.72102
	ML	1.14228	1.32203	0.70788	0.60097	0.49096	0.45064
6	LL	-1.29703	1.85107	-2.0061	4.09134	-2.05528	4.27228
	LS	-0.75366	0.57378	-1.12904	1.27584	-1.09395	1.19765
	MCS	-0.66634	0.45313	-1.15107	1.34096	-1.12208	1.26895
	MMCS	-1.56246	2.57472	-2.19784	4.87284	-2.25229	5.09417
	ML	1.12798	1.27172	0.79064	0.65329	0.54478	0.3681
7	LL	-1.02498	1.17401	-1.83475	3.4358	-1.9803	3.98601
	LS	-0.65748	0.44295	-1.02859	1.05909	-0.98372	0.96834
	MCS	-0.48475	0.24435	-0.98431	0.97967	-0.95296	0.93086
	MMCS	-1.29125	1.7757	-1.98837	3.98878	-2.04231	4.18893
	ML	0.95086	0.90899	0.68305	0.47989	0.49133	0.26628
8	LL	-0.83322	0.8302	-1.68809	2.91983	-1.85354	3.49521
	LS	-0.53808	0.31418	-0.91262	0.8342	-0.86477	0.7485
	MCS	-0.32116	0.11961	-0.84336	0.7159	-0.80791	0.65574
	MMCS	-1.1013	1.32964	-1.8014	3.27932	-1.8651	3.49403
	ML	0.84066	0.70562	0.63795	0.4136	0.45703	0.22353
9	LL	-0.71323	0.68146	-1.6113	2.66468	-1.79078	3.25829
	LS	-0.49255	0.28721	-0.89268	0.79847	-0.841	0.70802
	MCS	-0.2461	0.07872	-0.78369	0.63085	-0.75495	0.58008
	MMCS	-1.00164	1.13004	-1.71598	2.97832	-1.7824	3.19186
	ML	0.55169	0.30514	0.38944	0.15557	0.23561	0.06368
10	LL	-0.60837	0.55152	-1.54134	2.44418	-1.73634	3.07185
	LS	-0.43355	0.24365	-0.86783	0.75506	-0.81333	0.66233
	MCS	-0.17633	0.04579	-0.72625	0.53034	-0.70061	0.494
	MMCS	-0.9109	0.94379	-1.63465	2.70386	-1.69863	2.89938
	ML	0.34332	0.11822	0.20984	0.04628	0.06543	0.00949

Table 2. *The MSEs and Biases of Estimates in Equi-spaced Grouping at n=50*

K	Method	$\theta=2$		$\theta=4$		$\theta=6$	
		Biases(θ)	MSE(θ)	Biases(θ)	MSE(θ)	Biases(θ)	MSE(θ)
5	LL	-1.6846	3.03361	-1.6148	2.74724	-1.8963	3.69723
	LS	-1.5412	2.43273	-1.2595	1.6096	-1.4066	1.99997
	MCS	-1.2805	1.65063	-0.9487	0.90115	-1.0517	1.10634
	MMCS	-1.7866	3.33326	-1.6538	2.82784	-1.9536	3.90001
	ML	-0.9365	0.89067	-0.8975	0.8217	-0.9870	0.99987
6	LL	-1.6496	2.96005	-1.4815	2.2898	-1.7347	3.10394
	LS	-1.4507	2.15991	-1.1668	1.38451	-1.3171	1.75594
	MCS	-1.1861	1.41831	-0.8620	0.74426	-0.9681	0.93753
	MMCS	-1.7344	3.15947	-1.5919	2.62535	-1.8786	3.5983
	ML	-0.8493	0.73453	-0.8061	0.66544	-0.8943	0.82461
7	LL	-1.5515	2.5991	1.35524	1.9412	-1.6525	2.83409
	LS	-1.3540	1.88966	-1.0715	1.17105	-1.2223	1.51501
	MCS	-1.0874	1.19429	-0.7697	0.59365	-0.8789	0.77277
	MMCS	-1.7022	3.05252	-1.5244	2.40514	-1.7847	3.24114
	ML	-0.7570	0.58592	-0.7112	0.52118	-0.7986	0.66219
8	LL	-1.4315	2.23652	-1.2654	1.70896	-1.5381	2.45972
	LS	-1.2572	1.6353	-0.9743	0.97217	-1.1258	1.2885
	MCS	-0.9873	0.98691	-0.6738	0.45525	-0.7863	0.61859
	MMCS	-1.6458	2.86064	-1.4423	2.14966	-1.6758	2.85842
	ML	-0.6619	0.45086	-0.6145	0.39279	-0.7014	0.51608
9	LL	-1.3209	1.94526	-1.1597	1.44463	-1.4311	2.15248
	LS	-1.1569	1.39434	-0.8763	0.7909	-1.0283	1.07837
	MCS	-0.8847	0.79529	-0.5759	0.33305	-0.6918	0.47886
	MMCS	-1.5873	2.66425	-1.3447	1.86851	-1.5565	2.46709
	ML	-0.5653	0.33215	-0.5168	0.28207	-0.6034	0.38789
10	LL	-1.2135	1.65892	-1.0483	1.1973	-1.3347	1.87948
	LS	-1.0575	1.17429	-0.7763	0.62569	-0.9303	0.88641
	MCS	-0.7824	0.62483	-0.4770	0.22899	-0.5959	0.35533
	MMCS	-1.5150	2.42818	-1.2385	1.5864	-1.4316	2.09286
	ML	-0.4677	0.23124	-0.4184	0.18995	-0.5047	0.27845

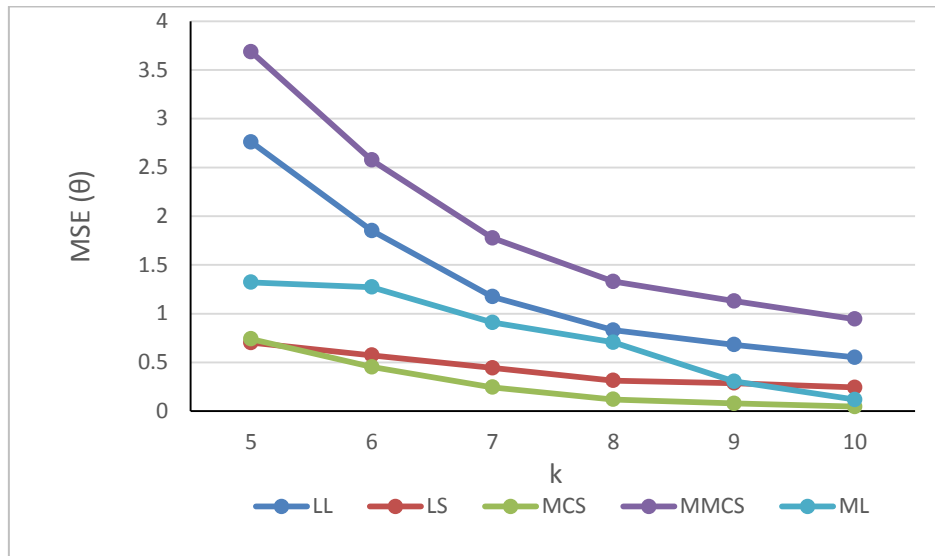


Figure 1. MSEs of all estimates in unequi-spaced grouping for $\theta = 2$ and $n = 50$

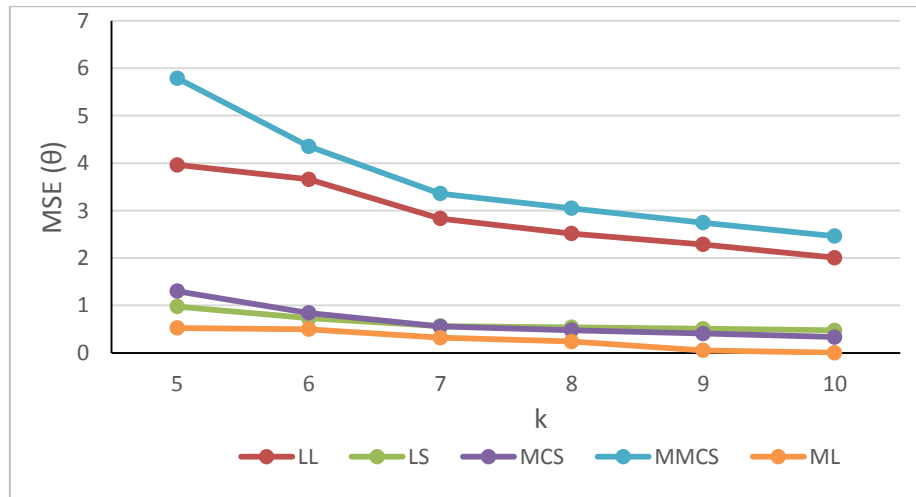


Figure 2. MSEs of all estimates in unequi-spaced grouping for $\theta = 4$ and $n = 150$

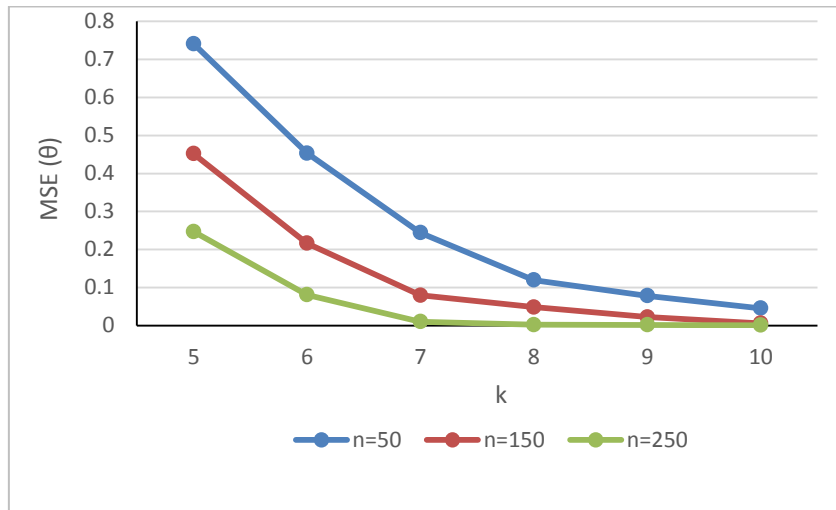


Figure 3. The MSEs of MCS estimate in unequi-spaced grouping at $\theta = 2$ for different values of n

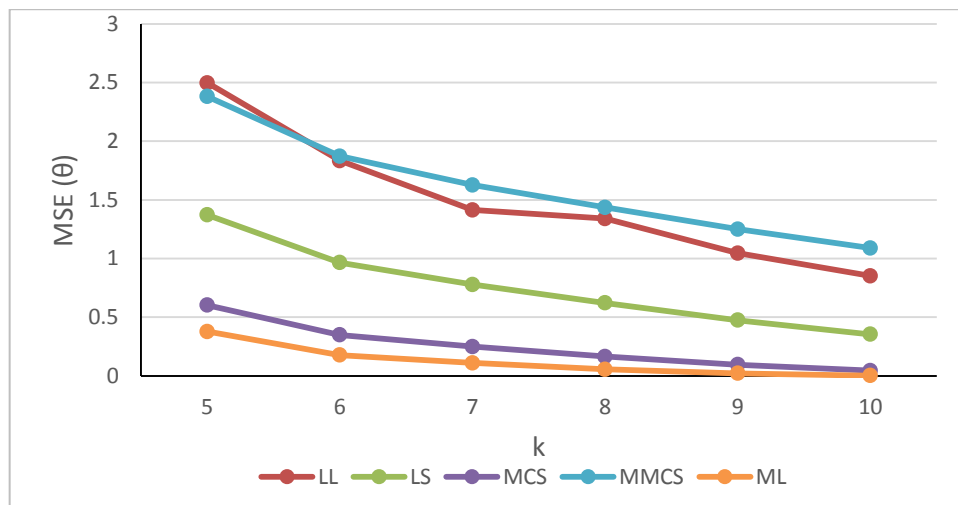


Figure 4. The MSEs of all estimates in equi-spaced grouping for $\theta = 4$ and $n = 150$

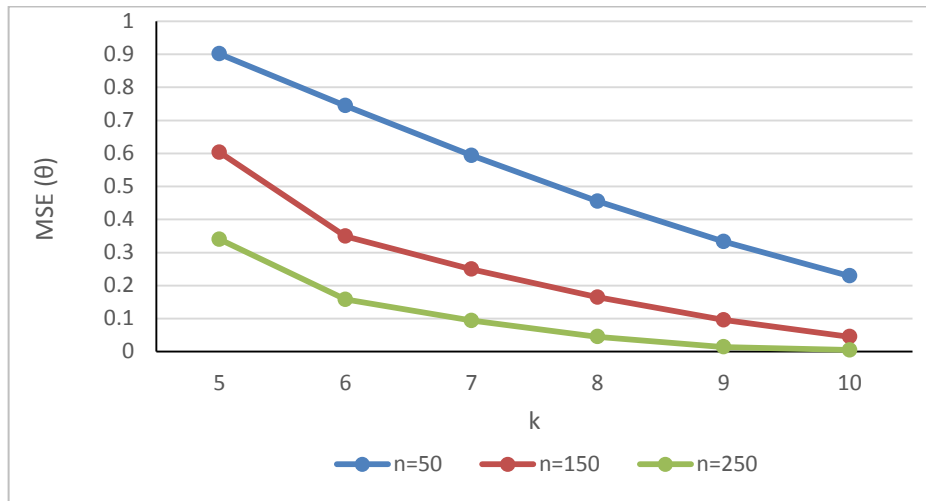


Figure 5. The MSEs of MCS estimate of equi-spaced grouping at $\theta = 4$ for different values n

6. Conclusion

This chapter is devoted to estimate the unknown parameter of the xgamma distribution based on grouped data with equi- and unequi-spaced grouping depending on five different estimation methods. A simulation study is performed to compare the performance of different estimates in both cases.

Based on simulation study, it can be observed that; the estimates with unequi-spaced sample have smaller MSE's compared with the equi-spaced grouped sample in almost situations. The MCS and ML estimates are better than the other estimates relative to their biases and MSEs in the case of unequi-spaced grouping. While in case of equi-spaced grouping the ML estimates are better than the other estimate relative to their biases and MSEs. This study revealed that the MCS and ML estimates perform the best among rest estimates, in approximately, most of situations.

Acknowledgements:

First of all, all gratitude is to "Allah" who guided and aided me to facilitate this thesis and his guidance and uncountable gifts to me through all my life.

Foremost, I would like to express my deepest and gratitude to my supervisor Prof. Dr. Amal Soliman Hassan for her continued guidance, offered facilities, suggestions, support, advice, encouragement patience and help during the preparation of this paper.

I wish to express my heartfelt thanks to my parents for their patience, especially my father for his encouragement and moral support. Also I would like to thank my sisters, my brothers for their support and encouragement although part of them had no idea what I was doing.

Finally I dedicate this work to my husband and my little Children Jody and Farida God bless their life.

References

- Hassan, A. S., Marwa, A. A., Zaher, H., & Elsherpieny, E. A. (2014). Comparison of estimators for exponentiated inverted Weibull distribution based on grouped data, *Journal of Engineering Research and Applications*, 4(1), 477-90.
- Kim, N. (2016). On the maximum likelihood estimators for parameters of a Weibull distribution under random censoring. *Communications for Statistical Applications and Methods*, 23(3), 241-250.
- Kim, N. (2018). On the maximum likelihood estimation for a normal distribution under random censoring. *Communications for Statistical Applications and Methods*, 25(6), 647-658.
- Kulldorff, G. (1961). *Contributions to the Theory of Estimation from Grouped and Partially Grouped Samples*. John Wiley Sons, New York.
- Lindley, D. V. (1950). Grouping corrections and maximum likelihood equations. In *Mathematical Proceedings of the Cambridge Philosophical Society*, Cambridge University Press, 46 (1), 106-110.
- McDonald, J. B., & Ransom, M. R. (1979). Functional forms, estimation techniques and the distribution of income. *Econometrica: Journal of the Econometric Society*, 47(6), 1513-1525.
- Nowak, P. B. (2019). Moment-type estimation from grouped samples. *Statistics and Probability Letters*, 149, 80-85.
- Sen, S., Maiti, S. S., & Chandra, N. (2016). The xgamma distribution: statistical properties and application. *Journal of Modern Applied Statistical Methods*, 15(1), 774-788.
- Stevens, W. L. (1948). Control by gauging. *Journal of the Royal Statistical Society. Series B (Methodological)*, 10(1), 54-108.
- Zhang, M. J. (1997). Grouped failure times tied failure times two contributions to the encyclopedia of biostatistics. Ph.D. thesis. Division of Biostatistics. Medical College of Wisconsin.

New Transformation for Generating Families of Continuous Distributions with Application to Weibull Distribution

Samah Ali El-Taweel

Department of Mathematical Statistics,
Cairo University/ Faculty of Graduate Studies for Statistical Research,
Giza, Egypt
E_mail: h_s_elsafty@hotmail.com

Abstract

In reliability study, social sciences and others, the classical statistical distributions are not suitable for describing and predicting real world phenomena. Generated families of probability distributions are recent development which extend and get better fits than the usual distributions. In this paper, we introduce for the first time, a new transformation which is more flexible to generate new families of probability distributions. We call the new family as a new Kumaraswamy-Generated(NK_w-G), depending on the Kumaraswamy distribution as a generator. We hope that this new family yields a better fit in certain practical situations. Some special sub-models are illustrated. We introduce and study some of its statistical properties, including; useful expansions, probability weighted moments, moments, and Rényi entropy. Maximum likelihood estimation is proposed for estimating the model parameters. Furthermore, two real data sets are employed to show the usefulness of the proposed family.

Keywords: *Kumaraswamy Distribution; Moments; Quantile function; Rényi Entropy; Maximum Likelihood Estimation.*

1. Introduction

Kumaraswamy (1980) presented a two-parameter distribution on the interval $[0,1]$, the so called Kumaraswamy (K_w) distribution, as continuous probability for double bounded random processes of hydrological applications. The probability density function (pdf) and cumulative distribution function (cdf) of K_w distribution are given, respectively, as follows

$$f_{K_w}(t) = abt^{a-1} (1-t^a)^{b-1}; a, b > 0, 0 < t < 1, \quad (1)$$

and

$$F_{K_w}(t) = 1 - (1-t^a)^b.$$

Recently, several ways of generating new distributions from classic ones have been attracted theoretical and applied statisticians due to their flexible properties. Some of the well-known generators are, the beta-G (Eugene et al. (2002)), Kumaraswamy-G (Cordeiro and de Castro (2011)), gamma-G (Zografos and Balakrishnan (2009)), transformed-transformer (T-X) (Alzaatreh et al. (2013)), exponentiated T-X (Alzaghal et al. (2013)), Weibull-G (Bourguignon et al. (2014)), new K_w - K_w (Mahmoud et al. (2015)), K_w -G Poisson (Ramos et al. (2015)), K_w odd log-logistic-G (Alizadeh et al. (2015)), K_w transmuted-G (Afify et al. (2016)), logistic-X (Tahir et al. (2016)), K_w Weibull-G (Hassan and Elgarhy (2016a)), exponentiated Weibull-G (Hassan and Elgarhy (2016b)), additive Weibull-G (Hassan and Hemeda (2016)), Odd Lindley-G (Silva et al. (2017)), Type-II half logistic (Hassan et al. (2017a)), generalized additive Weibull-G (Hassan et al. (2017b)), Marsall-Olkin K_w -G (Handique et al. (2017a)), the generalized Marsall-Olkin K_w -G (Chakraborty and Handique (2017)), beta generated K_w -G (Handique et al. (2017b)), and power Lindley-G (Hassan and Nasser (2018)) among others.

Alzaatreh et al. (2013) proposed a general method that allows the use of any non-negative continuous random variable T as a generator, the *cdf* of the T - X family is as follows

$$F(x) = \int_0^{\omega(G(x;\xi))} r(t) dt, \quad (2)$$

where $r(t)$ be the *pdf* of a random variable $T \in [c, d]$ for $-\infty \leq c < d \leq \infty$, $G(x; \xi)$ is the baseline *cdf* of any random variable X , ξ is the parameter vector of the baseline distribution and $\omega(G(x; \xi))$ satisfies the following conditions:

- 1- $\omega(G(x, \xi)) \in [c, d]$.
- 2- $\omega(G(x, \xi))$ is differentiable and monotonically non decreasing function.
- 3- $\omega(G(x, \xi)) \rightarrow c, \text{ as } x \rightarrow -\infty, \omega(G(x, \xi)) \rightarrow d, \text{ as } x \rightarrow \infty$.

Several generated families with different transformations for the upper limit $\omega(G(x, \xi))$ and for different generators $r(t)$ have been discussed by several authors.

Our motivation of this paper is introducing for the first time, a new transformation which is defined as follows;

$$\omega(G(x; \xi)) = 1 - e^{-\left(\frac{G(x; \xi)}{1 - G(x; \xi)}\right)}.$$

This transformation is more flexible and at the same time satisfies the above three conditions; $\omega(G(x; \xi)) \in [0, 1]$. Based on this new transformation and based on the *Kw* distribution as a generator, we propose a new family of probability distribution, named as new Kumaraswamy-generated (*NKw-G*) family. The *NKw-G* is provides more flexibility to model specific real data.

This paper is unfolded as follows. In Section 2, the construction of the *NKw-G* family of distributions is introduced. In Section 3, some special sub-models of the *NKw-G* family are defined. Section 4 provides some general statistical properties of the family. Maximum likelihood estimation of the model parameters is addressed in Section 5. Section 6 includes numerical study for the New Kumaraswamy-Weibull (*NKw-W*) model. An illustrative application based on two real data is examined in Section 7. Concluding remarks are presented in Section 8.

2. The New Kumaraswamy-Generated Family

This section provides the construction of the new family of probability distributions, called New Kumaraswamy-Generated (*NKw-G*) family of distributions. Expansions of its *pdf* and *cdf* are also provided.

2.1 The Distribution and Density Function

Based on (2), the *cdf* of the *NKw-G* family can be given by substituting:

- (i) The generator $r(t)$ by the *pdf* of the *Kw* distribution given in (1), and
- (ii) $\omega(G(x; \xi))$ with $1 - e^{-\left(\frac{G(x; \xi)}{1 - G(x; \xi)}\right)}$,

Then, the *cdf* of the *NKw-G* family is obtained as follows

$$F_{NKw-G}(x) = \int_0^{1-e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)}} abt^{a-1}(1-t)^{b-1} dt = 1 - \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^b; a, b > 0. \quad (3)$$

The corresponding *pdf* is as follows

$$f_{NKw-G}(x) = ab \frac{g(x;\xi)}{(1-G(x;\xi))^2} e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^{a-1} \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^{b-1}, \quad (4)$$

where $a, b > 0$ are the two shape parameters. Hereafter, a random variable X with *pdf* (4) is denoted by $X \sim NKw-G(a, b, \xi)$.

2.2 Reliability Functions

The survival, hazard and reversed hazard rate functions for $NKw-G$ family are given, respectively by

$$S_{NKw-G}(x) = \bar{F}_{NKw-G}(x) = \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^b,$$

$$H_{NKw-G}(x) = \frac{f_{NKw-G}(x)}{S_{NKw-G}(x)} = \frac{abg(x;\xi)}{(1-G(x;\xi))^2} e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^{a-1} \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^{-1},$$

and

$$\tau_{NKw-G}(x) = \frac{f_{NKw-G}(x)}{F_{NKw-G}(x)} = \frac{abg(x;\xi) e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^{a-1} \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^{b-1}}{(1-G(x;\xi))^2 \left[1 - \left(1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right)^a \right]^b}.$$

2.3 Useful Expansions

For simplicity, some representations of the *pdf* and *cdf* for the $NKw-G$ family are deduced. If $b > 0$ is a real non-integer and $|z| < 1$, the generalized binomial theorem is written as

$$[1-z]^{b-1} = \sum_{j=0}^{\infty} (-1)^j \binom{b-1}{j} (z)^j. \quad (5)$$

By applying binomial expansion (5) in (4), the *pdf* of the $NKw-G$ family, where b is a real non integer, can be written as

$$f_{NKw-G}(x) = ab \frac{g(x;\xi)}{(1-G(x;\xi))^2} e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \sum_{j=0}^{\infty} (-1)^j \binom{b-1}{j} \left[1 - e^{-\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)} \right]^{a(j+1)-1}.$$

If b is an integer, the index j in the previous sum stop at $b-1$. By repeating the binomial expansion, where a is a real non-integer, leads to:

$$f_{NKw-G}(x) = ab \frac{g(x;\xi)}{(1-G(x;\xi))^2} \sum_{j,k=0}^{\infty} (-1)^{j+k} \binom{b-1}{j} \binom{a(j+1)-1}{k} e^{-(k+1)\left(\frac{G(x;\xi)}{1-G(x;\xi)}\right)}. \quad (6)$$

By using the power series and the generalized binomial theorem, then (6) can be written as follows

$$f_{NKw-G}(x) = abg(x; \xi) \sum_{j,k,l,m=0}^{\infty} \frac{(-1)^{j+k+l} (k+1)^l}{l!} \binom{b-1}{j} \binom{a(j+1)-1}{k} \binom{l+1+m}{m} (G(x; \xi))^{l+m}. \quad (7)$$

The *pdf* of the *NKw-G* distribution in (7) can be expressed as an infinite linear combination of exponentiated-G density function, that is,

$$f_{NKw-G}(x) = \sum_{j,k,l,m=0}^{\infty} W_{j,k,l,m} g(x; \xi) (G(x; \xi))^{l+m}, \quad (8)$$

Then

$$f_{NKw-G}(x) = \sum_{j,k,l,m=0}^{\infty} \eta_{j,k,l,m} h_{(l+m+1)}(x; \xi),$$

where

$$\eta_{j,k,l,m} = \frac{W_{j,k,l,m}}{l+m+1} = \frac{ab(-1)^{j+k+l} (k+1)^l}{(l+m+1)l!} \binom{b-1}{j} \binom{a(j+1)-1}{k} \binom{l+m+1}{m}; \quad \sum_{j,k,l,m=0}^{\infty} \eta_{j,k,l,m} = 1,$$

and

$$h_{(l+m+1)}(x; \xi) = (l+m+1)g(x; \xi) (G(x; \xi))^{l+m}.$$

If a, b and l, m are integers, the index j stop at $b-1$, the index k stops at $a(j+1)-1$, and m stops at $l+m+1$.

Additionally, an expansion for the *cdf* in (3) can be obtained as follows; by using binomial expansion for $[F_{NKw-G}(x)]^h$, where h is an integer, leads to

$$[F_{NKw-G}(x)]^h = \sum_{g=0}^h (-1)^g \binom{h}{g} \left(1 - \left(1 - e^{-\left(\frac{G(x; \xi)}{1-G(x; \xi)} \right)^a} \right)^{bg} \right).$$

Using binomial expansion again, where b is a real non integer, also, using binomial expansion another time, where a is a real non integer, yields

$$[F_{NKw-G}(x)]^h = \sum_{g=0}^h \sum_{p,q=0}^{\infty} (-1)^{g+p+q} \binom{h}{g} \binom{bg}{p} \binom{ap}{q} e^{-q \left(\frac{G(x; \xi)}{1-G(x; \xi)} \right)^a}.$$

By using power series for the exponential function and the generalized binomial theorem, then, we obtain

$$[F_{NKw-G}(x)]^h = \sum_{g=0}^h \sum_{p,q,s=0}^{\infty} \sum_{t=0}^s \frac{(-1)^{g+p+q+s} q^s}{s!} \binom{h}{g} \binom{bg}{p} \binom{ap}{q} \binom{s}{t} (G(x; \xi))^{s+t},$$

Replacing $\sum_{s=0}^{\infty} \sum_{t=0}^s$ with $\sum_{s=t}^{\infty} \sum_{t=0}^{\infty}$ yields

$$[F_{NKw-G}(x)]^h = \sum_{g=0}^h \sum_{p,q=0}^{\infty} \sum_{s=t}^{\infty} \sum_{t=0}^{\infty} \frac{(-1)^{g+p+q+s} q^s}{s!} \binom{h}{g} \binom{bg}{p} \binom{ap}{q} \binom{s}{t} (G(x; \xi))^{s+t}.$$

Finally,

$$[F_{NKw-G}(x)]^h = \sum_{s=t}^{\infty} \sum_{t=0}^{\infty} S_{s,t} (G(x; \xi))^{s+t}, \quad S_{s,t} = \sum_{g=0}^h \sum_{p,q=0}^{\infty} \frac{(-1)^{g+p+q+s} q^s}{s!} \binom{h}{g} \binom{bg}{p} \binom{ap}{q} \binom{s}{t}. \quad (9)$$

3. Special Sub-Models

In present section, we give some interesting sub models from *NKw-G* family.

3.1 The New Kumaraswamy-Weibull Distribution

Consider the random variable X follows the Weibull (W) distribution with *pdf* and *cdf* given by $g(x; \alpha, \beta) = \alpha \beta x^{\beta-1} e^{-\alpha x^\beta}$; $x, \alpha, \beta > 0$ and $G(x; \alpha, \beta) = 1 - e^{-\alpha x^\beta}$. The New Kumaraswamy-Weibull ($NKw-W$) distribution has following *cdf* and *pdf* respectively

$$F_{NKw-W}(x) = 1 - \left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right); x > 0, a, b, \alpha, \beta > 0, \quad (10)$$

and

$$f_{NKw-W}(x) = ab \alpha \beta x^{\beta-1} e^{\alpha x^\beta} e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^{a-1} \left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right)^{b-1}; x > 0, a, b, \alpha, \beta > 0. \quad (11)$$

The survival, hazard and reversed hazard rate functions for the $NKw-W$ distribution are obtaining as follows:

$$S_{NKw-W}(x) = \left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right),$$

$$H_{NKw-W}(x) = ab \alpha \beta x^{\beta-1} e^{\alpha x^\beta} e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^{a-1} \left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right)^{-1},$$

and

$$\tau_{NKw-W}(x) = \frac{ab \alpha \beta x^{\beta-1} e^{\alpha x^\beta} e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^{a-1} \left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right)^{b-1}}{\left(1 - \left(1 - e^{-\left(e^{\alpha x^\beta} - 1 \right)^a} \right)^b \right)}.$$

Note that the $NKw-W$ distribution is very flexible model that approaches to different distributions when its parameters are changed. The $NKw-W$ distribution contains as special-models in the following table:

Table 1. Special-models for the $NKw-W$

Distribution	a	b	α	β
New Exponentiated-Weibull ($NE-W$)	-	1	-	-
New Kumaraswamy-Exponential ($NKw-E$)	-	-	-	1
New Kumaraswamy-Rayleigh ($NKw-R$)	-	-	-	2

3.2 The New Kumaraswamy-Kumaraswamy Distribution

Suppose that the parent distribution is Kw distribution in the interval $0 < x < 1$ with the following *pdf* and *cdf* respectively $g(x; a_1, b_1) = a_1 b_1 x^{a_1-1} (1-x^{a_1})^{b_1-1}$; $a_1, b_1 > 0$ and $G(x; a_1, b_1) = 1 - (1-x^{a_1})^{b_1}$. The New Kumaraswamy-Kumaraswamy ($NKw-Kw$) distribution has the following *cdf* and *pdf* respectively as follows;

$$F_{NKw-Kw}(x) = 1 - \left[1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1} - 1 \right)^a} \right)^b \right]; 0 < x < 1, a, b, a_1, b_1 > 0,$$

and

$$f_{NKw-Kw}(x) = ab a_1 b_1 \frac{x^{a_1-1}}{(1-x^{a_1})^{b_1+1}} e^{-\left((1-x^{a_1})^{-b_1}-1\right)} \left[1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right]^{a-1} \left(1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right)^a\right)^{b-1}; 0 < x < 1, a, b, a_1, b_1 > 0.$$

Therefore, the survival, hazard and reversed hazard rate functions for the $NKw-Kw$ distribution are obtained, respectively by

$$S_{NKw-Kw}(x) = \left[1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right)^a\right]^b,$$

$$H_{NKw-Kw}(x) = \frac{ab a_1 b_1 x^{a_1-1}}{(1-x^{a_1})^{b_1+1}} e^{-\left((1-x^{a_1})^{-b_1}-1\right)} \left[1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right]^{a-1} \left(1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right)^a\right)^{-1},$$

and

$$\tau_{NKw-Kw}(x) = \frac{ab a_1 b_1 x^{a_1-1} e^{-\left((1-x^{a_1})^{-b_1}-1\right)} \left[1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right]^{a-1} \left(1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right)^a\right)^{b-1}}{(1-x^{a_1})^{b_1+1} \left[1 - \left(1 - \left(1 - e^{-\left((1-x^{a_1})^{-b_1}-1\right)}\right)^a\right)^b\right]}.$$

3.3. The New Kumaraswamy-Uniform Distribution

Suppose that the parent distribution is a uniform in the interval $(0, \theta)$ with the $cdf G(x; \theta) = x / \theta$. The New Kumaraswamy-Uniform ($NKw-U$) distribution has the following cdf and pdf :

$$F_{NKw-U}(x) = 1 - \left(1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right)^b; 0 < x < \theta, a, b, \theta > 0,$$

and

$$f_{NKw-U}(x) = ab \frac{\theta}{(\theta-x)^2} e^{-\left(\frac{x}{\theta-x}\right)} \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^{a-1} \left(1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right)^{b-1}; 0 < x < \theta, a, b, \theta > 0.$$

Therefore, the survival, hazard and reversed hazard rate functions for $NKw-U$ distribution are as follows:

$$S_{NKw-U}(x) = \left[1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right]^b,$$

$$H_{NKw-U}(x) = \frac{ab \theta}{(\theta-x)^2} e^{-\left(\frac{x}{\theta-x}\right)} \left[1 - e^{-\left(\frac{x}{\theta-x}\right)}\right]^{a-1} \left(1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right)^{-1},$$

and

$$\tau_{NKw-U}(x) = \frac{ab\theta e^{-\left(\frac{x}{\theta-x}\right)} \left[1 - e^{-\left(\frac{x}{\theta-x}\right)}\right]^{a-1} \left(1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right)^{b-1}}{(\theta-x)^2 \left[1 - \left(1 - \left(1 - e^{-\left(\frac{x}{\theta-x}\right)}\right)^a\right)^b\right]}.$$

Figures 1 and 2 below show the *pdf* curves and the hazard rate curves for the *NKw –W* , *NKw –Kw* and *Kw –U* distributions for some selected values of parameters.

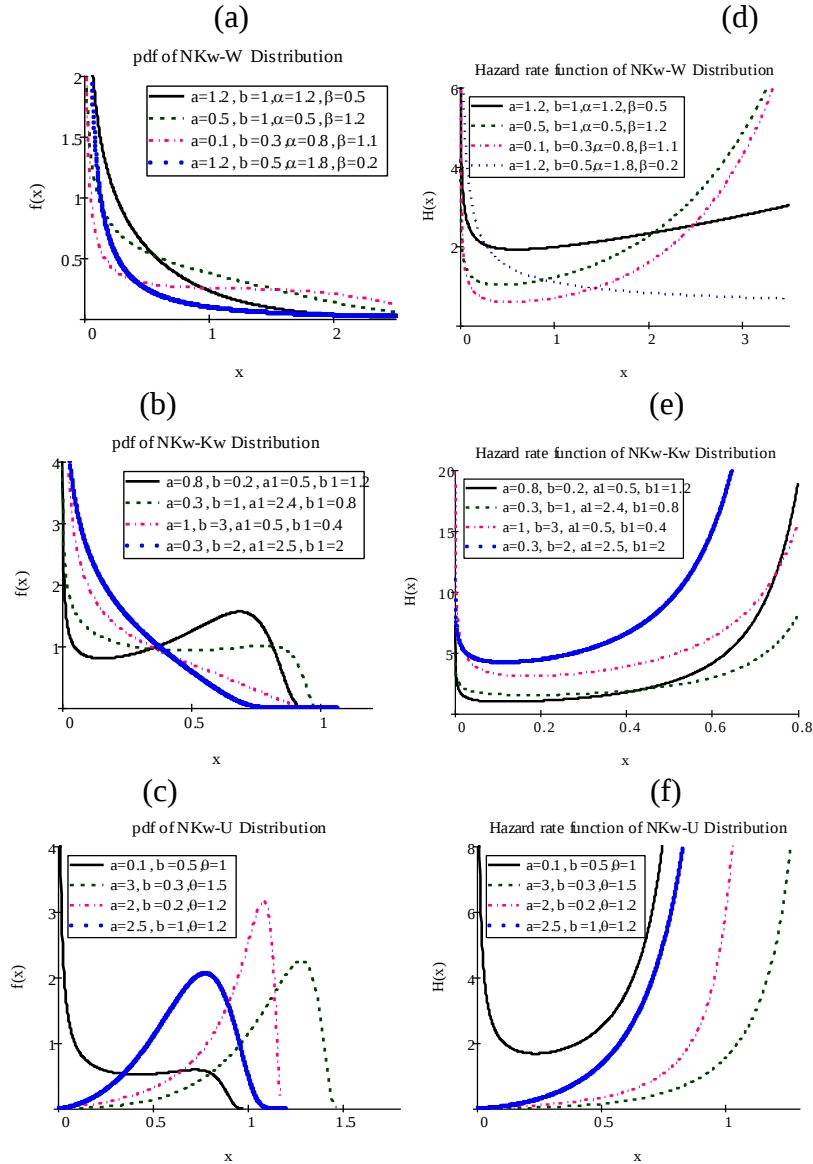


Figure 1. Density plots (a) *NKw –W* (a, b, α, β),
(b) *NKw –Kw* (a, b, a_1, b_1),
(c) *NKw –U* (a, b, θ).

Figure 2. Hazard rate plots (d) *NKw –W* (a, b, α, β),
(e) *NKw –Kw* (a, b, a_1, b_1),
(f) *NKw –U* (a, b, θ).

4. Statistical Properties

This section concerns with some properties of the *NKw-G* family of distributions.

4.1 Quantile function

Quantile functions are considered as an important point of probability theory, statistical application and simulation. Simulation methods are used to simulate random variable for classical and new continuous distributions. The quantile function of the $NKw -G$ family, say $Q(u) = F^{-1}(u)$ of X is given by

$$Q_{NKw-G}(u) = G^{-1} \left[\ln \left[1 - \left[1 - (1-u)^{1/b} \right]^{1/a} \right]^{-1} / 1 - \ln \left[1 - \left[1 - (1-u)^{1/b} \right]^{1/a} \right] \right]; 0 < u < 1, \quad (12)$$

where $G^{-1}(\cdot)$ is the inverse function of $G(\cdot)$, u is defined on the interval $(0,1)$ and it considered as a uniform random variable.

Example 4.1.1. The quantile function of the $NKw-W$ distribution discussed in subsection (3.1) is straightforward to be computed by inverting (10), then

$$Q_{NKw-W}(u) = \left[\ln \left(1 - \ln \left[1 - \left(1 - (1-u)^{1/b} \right)^{1/a} \right] \right)^{1/\alpha} \right]^{1/\beta}. \quad (13)$$

We can obtain the first quartile, the median, and the third quartile of $NKw-W$ distribution by setting $u = 0.25, 0.5$ and 0.75 , respectively, in (13).

4.2 The Probability Weighted Moments

The probability weighted moments (PWMs) proposed by Greenwood et al. (1979). A general theory for PWMs covers the summarization and description of theoretical probability distributions. The PWMs method can generally be used for estimating parameter of distribution whose inverse form cannot be expressed explicitly. For r and d are positive integers, the d th PWMs, denoted by $O_{r,d}$ can be generally defined as follows

$$O_{r,d} = E \left\{ X^r (F(x))^d \right\} = \int_{-\infty}^{\infty} x^r f(x) (F(x))^d dx. \quad (14)$$

Inserting (8) and (9) into (14), the PWMs of the $NKw -G$ family can be given by

$$O_{r,d}(NKw -G) = \sum_{j,k,l,m=0}^{\infty} \sum_{s=t}^{\infty} S_{s,t} W_{j,k,l,m} \int_{-\infty}^{\infty} x^r g(x; \xi) (G(x; \xi))^{l+m+s+t} dx.$$

Then

$$O_{r,d}(NKw -G) = \sum_{j,k,l,m=0}^{\infty} \sum_{s=t}^{\infty} S_{s,t} W_{j,k,l,m} O_{r,l+m+s+t},$$

where $O_{r,l+m+s+t} = \int_{-\infty}^{\infty} x^r g(x; \xi) (G(x; \xi))^{l+m+s+t} dx$ is the PWMs for $r, l+m+s+t$ of the baseline distribution.

Example 4.2.1. The PWMs of the $NKw -W$ distribution is obtained as follows

$$O_{r,d}(NKw -W) = \sum_{j,k,l,m=0}^{\infty} \sum_{s=t}^{\infty} \sum_{u=0}^{l+m+s+t} S_{s,t} W_{j,k,l,m,u} \left[\Gamma \left(\frac{r}{\beta} + 1 \right) / \alpha^{\frac{r}{\beta}} (u+1)^{\frac{r}{\beta}+1} \right],$$

where $W_{j,k,l,m,u} = (-1)^u \binom{l+m+s+t}{u} W_{j,k,l,m}$.

4.3 Moments

Moments are necessary and important in any statistical analysis, especially in applications. It can be used to study the most important features and characteristics of a distribution (e.g., dispersion, skewness and kurtosis). The r th moment, denoted by μ'_r can be generally defined as follows

$$\mu'_r = E(X^r) = \int_{-\infty}^{\infty} x^r f(x) dx. \quad (15)$$

Inserting (8) into (15), the r th moment of the $NKw-G$ family can be given by

$$\mu'_r(NKw-G) = \sum_{j,k,l,m=0}^{\infty} W_{j,k,l,m} O_{r,l+m},$$

where $O_{r,l+m}$ is PWMs which defined in (14), replacing d with $l+m$.

Example 4.3.1: The r th moment of the $NKw-W$ distribution can be obtained as follows

$$\mu'_r(NKw-W) = \sum_{j,k,l,m=0}^{\infty} \sum_{u=0}^{l+m} W'_{j,k,l,m,u} \left[\Gamma\left(\frac{r}{\beta} + 1\right) / \alpha^{\frac{r}{\beta}} (u+1)^{\frac{r}{\beta}+1} \right],$$

where $W'_{j,k,l,m,u} = (-1)^u \binom{l+m}{u} W_{j,k,l,m}$.

4.4 Rényi Entropy

The entropy of a random variable is a measure of uncertainty. Entropy has been used in various situations in science; engineering and numerous measures of entropy have been studied and compared in the literature. The Rényi entropy is defined as

$$R_e(\rho) = (1-\rho)^{-1} \log \left[\int_{-\infty}^{\infty} f(x)^{\rho} dx \right]; \rho > 0, \rho \neq 0. \quad (16)$$

Inserting (4) into (16), after some representations, the Rényi entropy of the $NKw-G$ family is as follows:

$$R_e(\rho)(NKw-G) = (1-\rho)^{-1} \log \left[\sum_{j,k,l,m=0}^{\infty} W_{j,k,l,m,\rho} \int_{-\infty}^{\infty} (g(x;\xi))^{\rho} (G(x;\xi))^{(l+m)} dx \right],$$

where

$$W_{j,k,l,m,\rho} = (ab)^{\rho} \frac{(-1)^{j+k+l} (\rho+k)^l}{l!} \binom{\rho(b-1)}{j} \binom{aj + \rho(a-1)}{k} \binom{l+2\rho+m-1}{m}.$$

Example 4.4.1: The Rényi entropy of the $NKw-W$ distribution can be obtained as follows

$$R_e(\rho)(NKw-W) = (1-\rho)^{-1} \log \left[\sum_{j,k,l,m=0}^{\infty} \sum_{u=0}^{l+m} W'_{j,k,l,m,\rho} \left[\Gamma\left(\rho\left(1-\frac{1}{\beta}\right) + \frac{1}{\beta}\right) / \beta (\alpha(\rho+u))^{\rho\left(1-\frac{1}{\beta}\right) + \frac{1}{\beta}} \right] \right],$$

where

$$W'_{j,k,l,m,\rho} = (\alpha\beta)^{\rho} (-1)^u \binom{l+m}{u} W_{j,k,l,m,\rho}.$$

5. Maximum Likelihood Estimation

Let $X_1, X_2, \dots, X_n$ be the independent identically (iid) random variables from the $NKw-G$ distribution (4) with parameter a, b and ξ , let Θ be the parameter vector, the log-likelihood function for the vector of parameters $\Theta = (a, b, \xi)^T$ can be expressed as

$$\ln L(\Theta) = n \ln(ab) + \sum_{i=1}^n \ln[g(x_i, \xi)] - 2 \sum_{i=1}^n \ln[1-G(x_i, \xi)] - \sum_{i=1}^n H(x_i, \xi) + (a-1) \sum_{i=1}^n \ln[1-e^{-(H(x_i, \xi))}] + (b-1) \sum_{i=1}^n \ln[1-(1-e^{-(H(x_i, \xi))})^a].$$

where $H(x_i, \xi) = \left(\frac{G(x_i, \xi)}{1-G(x_i, \xi)} \right)$. The elements of the score function $U(\Theta) = (U_a, U_b, U_\xi)$ are given by

$$U_a = \frac{n}{a} + \sum_{i=1}^n \ln[1-e^{-(H(x_i, \xi))}] - (b-1) \sum_{i=1}^n \left[\ln(1-e^{-(H(x_i, \xi))}) / (1-e^{-(H(x_i, \xi))})^{-a} - 1 \right],$$

$$U_b = \frac{n}{b} + \sum_{i=1}^n \ln[1-(1-e^{-(H(x_i, \xi))})^a],$$

and

$$U_\xi = \sum_{i=1}^n \frac{\partial g(x_i, \xi) / \partial \xi_k}{g(x_i, \xi)} + 2 \sum_{i=1}^n \frac{\partial G(x_i, \xi) / \partial \xi_k}{[1-G(x_i, \xi)]} - \sum_{i=1}^n \frac{\partial G(x_i, \xi) / \partial \xi_k}{[1-G(x_i, \xi)]^2} - (a-1) \sum_{i=1}^n \frac{\left(\frac{\partial G(x_i, \xi) / \partial \xi_k}{(1-G(x_i, \xi))^2} \right) e^{-(H(x_i, \xi))}}{1-e^{-(H(x_i, \xi))}} - a(b-1) \sum_{i=1}^n \left(\frac{\partial G(x_i, \xi) / \partial \xi_k}{(1-G(x_i, \xi))^2} \right) e^{-(H(x_i, \xi))} [1-e^{-(H(x_i, \xi))}]^{a-1} \left(1-(1-e^{-(H(x_i, \xi))})^a \right)^{-1}.$$

Setting U_a, U_b and U_ξ equal to zero and solving these equations simultaneously yield the maximum likelihood (ML) estimator $\hat{\Theta} = (\hat{a}, \hat{b}, \hat{\xi})$. These equations can be solved numerically.

Example 5.1: The maximum likelihood estimator (MLEs) of the $NKw-W$ distribution with pdf (11) can be obtained as follows:

The log likelihood function for the vector of parameters $\Theta = (a, b, \alpha, \beta)$ can be given by:

$$\ln L(\Theta) = n \ln(ab\alpha\beta) + (\beta-1) \sum_{i=1}^n \ln(x_i) + \alpha \sum_{i=1}^n x_i^\beta - \sum_{i=1}^n (1-e^{\alpha x_i^\beta}) + (a-1) \sum_{i=1}^n \ln(E_i) + (b-1) \sum_{i=1}^n \ln(1-E_i^a)$$

$$\text{where } E_i = \left(1 - e^{-\left(e^{\alpha x_i^\beta} - 1 \right)} \right).$$

The elements of the score function $U(\Theta) = (U_a, U_b, U_\alpha, U_\beta)$ are given by:

$$U_a = \frac{n}{a} + \sum_{i=1}^n \ln(E_i) - (b-1) \sum_{i=1}^n \frac{E_i^a \ln(E_i)}{(1-E_i^a)},$$

$$U_b = \frac{n}{b} + \sum_{i=1}^n \ln(1-E_i^a),$$

$$U_\alpha = \frac{n}{\alpha} + \sum_{i=1}^n x_i^\beta + \sum_{i=1}^n x_i^\beta e^{\alpha x_i^\beta} + (a-1) \sum_{i=1}^n \frac{x_i^\beta e^{\alpha x_i^\beta} (1-E_i)}{E_i} - a(b-1) \sum_{i=1}^n \frac{x_i^\beta e^{\alpha x_i^\beta} (1-E_i) E_i^{a-1}}{(1-E_i^a)},$$

and

$$U_\beta = \frac{n}{\beta} + \sum_{i=1}^n \ln(x_i) + \alpha \sum_{i=1}^n x_i^\beta \ln(x_i) + \alpha \sum_{i=1}^n x_i^\beta e^{\alpha x_i^\beta} \ln(x_i) + \alpha(a-1) \sum_{i=1}^n \frac{x_i^\beta \ln(x_i) e^{\alpha x_i^\beta} (1-E_i)}{E_i}$$

$$-a\alpha(b-1)\sum_{i=1}^n \frac{x_i^\beta \ln(x_i) e^{\alpha x_i^\beta} (1-E_i) E_i^{a-1}}{(1-E_i^a)}.$$

The MLEs of $\hat{\Theta}=(\hat{a},\hat{b},\hat{\alpha},\hat{\beta})$ can be obtained by solving the nonlinear likelihood equations $U_a=0, U_b=0, U_\alpha=0$ and $U_\beta=0$.

6. Simulation Study

In this section, simulation study will be carried for the *NKw-W* model, random samples of sizes $n=10,30,50,100$ are generated from the model and the parameter values have been estimated by ML method I : $a=2, b=0.2, \alpha=0.5, \beta=0.5$ and II : $a=1, b=0.5, \alpha=1.2, \beta=0.4$. 10000 replications are made to calculate the bias, variance and mean square error (MSE), respectively. All the computation are made using Mathcad 14 Software.

From the Table 2, it is observed that

1. The accuracy of the ML estimates of the parameters, where, as sample size n increases, bias decreases.
2. The consistency of the ML estimates of the parameters, where, as sample size n increases, MSE decreases.

Table 2: The estimates, bias, variance and MSE of parameters of $NKw-W$ model.

Groups	n	Parameters	Estimates	<i>bias</i>	Variance	MSE
I	10	$a = 2$	2.402	0.402	0.075	0.237
		$b = 0.2$	1.157	0.957	0.06	0.977
		$\alpha = 0.5$	0.999	0.499	0.023	0.273
		$\beta = 0.5$	0.164	-0.336	0.019	0.132
	30	$a = 2$	2.379	0.379	0.017	0.161
		$b = 0.2$	0.953	0.753	0.007	0.574
		$\alpha = 0.5$	0.932	0.432	0.006	0.193
		$\beta = 0.5$	0.068	-0.432	0.004	0.191
	50	$a = 2$	2.358	0.358	0.01	0.139
		$b = 0.2$	0.901	0.701	0.003	0.495
		$\alpha = 0.5$	0.93	0.43	0.004	0.188
		$\beta = 0.5$	0.048	-0.452	0.002	0.207
	100	$a = 2$	2.294	0.294	0.004	0.09
		$b = 0.2$	0.842	0.642	0.002	0.413
		$\alpha = 0.5$	0.916	0.416	0.002	0.175
		$\beta = 0.5$	0.018	-0.482	0.001	0.233
II	10	$a = 1$	1.039	0.039	0.002	0.0035
		$b = 0.5$	1.514	1.014	0.043	1.072
		$\alpha = 1.2$	1.452	0.252	0.0046	0.068
		$\beta = 0.4$	0.803	0.403	0.0047	0.167
	30	$a = 1$	1.001	0.0009	0.000008	0.000009
		$b = 0.5$	1.295	0.795	0.0028	0.635
		$\alpha = 1.2$	1.391	0.191	0.00021	0.037
		$\beta = 0.4$	0.736	0.336	0.0005	0.114
	50	$a = 1$	1	0.00002	0.00000005	0.00000005
		$b = 0.5$	1.244	0.744	0.001	0.555
		$\alpha = 1.2$	1.391	0.191	0.00008	0.036
		$\beta = 0.4$	0.718	0.318	0.0002	0.101
	100	$a = 1$	1	0	0	0
		$b = 0.5$	1.195	0.695	0.0003	0.483
		$\alpha = 1.2$	1.401	0.201	0.00003	0.04
		$\beta = 0.4$	0.705	0.305	0.00004	0.093

7. Applications to Real Data

In this section, two real data sets are employed to illustrate the importance of $NKw-W$ distribution as a particular case from the suggested family. The $NKw-W$ distribution is compared with Kumaraswamy-Weibull ($Kw-W$) presented by (Coreldiro, et al. 2010), and Weibull (W) distribution. The *pdfs* of $Kw-W$ and W distributions are defined by

$$f_{Kw-W}(x) = ab\alpha\beta x^{\beta-1}e^{-\alpha x^\beta} \left(1 - e^{-\alpha x^\beta}\right)^{a-1} \left(1 - \left(1 - e^{-\alpha x^\beta}\right)^a\right)^{b-1}; a, b, \alpha, \beta > 0, x > 0,$$

$$f_W(x) = \alpha\beta x^{\beta-1}e^{-\alpha x^\beta}; \alpha, \beta > 0, x > 0.$$

For both data sets, the following criteria are used to compare the three suggested models,

- Akaike information criterion (AIC): $AIC = 2k - 2\ln L$
- The correct AIC (CAIC): $CAIC = AIC + [2k(k+1)/n - k - 1]$
- Bayesian information criterion (BIC): $BIC = k\ln(n) - 2\ln L$
- Hannan-Quinn information criterion (HQIC): $HQIC = 2k\ln(\ln(n)) - 2\ln L$
- Kolmogorov-Smirnov (K-S) and p -value statistics: $K - S = \sup_y [F_n(y) - F(y)]$

where k is the number of parameters in each statistical model, n is the sample size and $\ln L$ is the maximized value of the log-likelihood function under the considered model, $F_n(y) = 1/n(\text{number of observation} \leq y)$, is the empirical distribution function and $F(y)$ denotes the *cdf* for each distribution. The "best" distribution corresponds to the smallest values of $-2\ln L$, AIC , $CAIC$, BIC , $HQIC$, $K-S$ and the biggest value of p -value.

7.1. Data set 1

The following data was taken from Murthy et al. (2004), which refer to the time between failures for a repairable item:

1.43, 0.11, 0.71, 0.77, 2.63, 1.49, 3.46, 2.46, 0.59, 0.74, 1.23, 0.94, 4.36, 0.40, 1.74, 4.73, 2.23, 0.45, 0.70, 1.06, 1.46, 0.30, 1.82, 2.37, 0.63, 1.23, 1.24, 1.97, 1.86, 1.17.

Figure 3 provides the plots of estimated densities and estimated cumulative of the fitted $NKw-W$, $Kw-W$ and W models for data set 1.

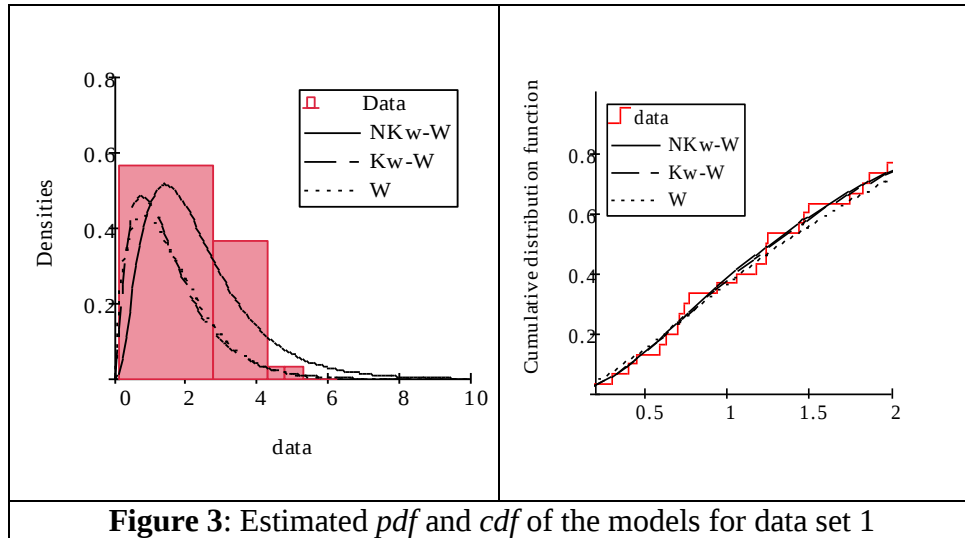


Figure 3: Estimated *pdf* and *cdf* of the models for data set 1

The following table shows the ML estimates and its standard error (S.E) of the model parameters for data set 1.

Table 3. The ML estimates and S.E of the model parameters for data set 1

Model	ML estimates and (S.E)			
$NKw - W(a, b, \alpha, \beta)$	67.626 (152.92)	35.8 (85.121)	1.331 (0.568)	0.07 (0.073)
$Kw - W(a, b, \alpha, \beta)$	1.864 (1.744)	0.517 (2.354)	1.433 (4.528)	1.143 (0.83)
$W(\alpha, \beta)$	- -	- -	0.456 (0.114)	1.463 (0.203)

The following table gives the values of mesurments for the data set 1.

Table 4. Model Selection Criteria for data set 1

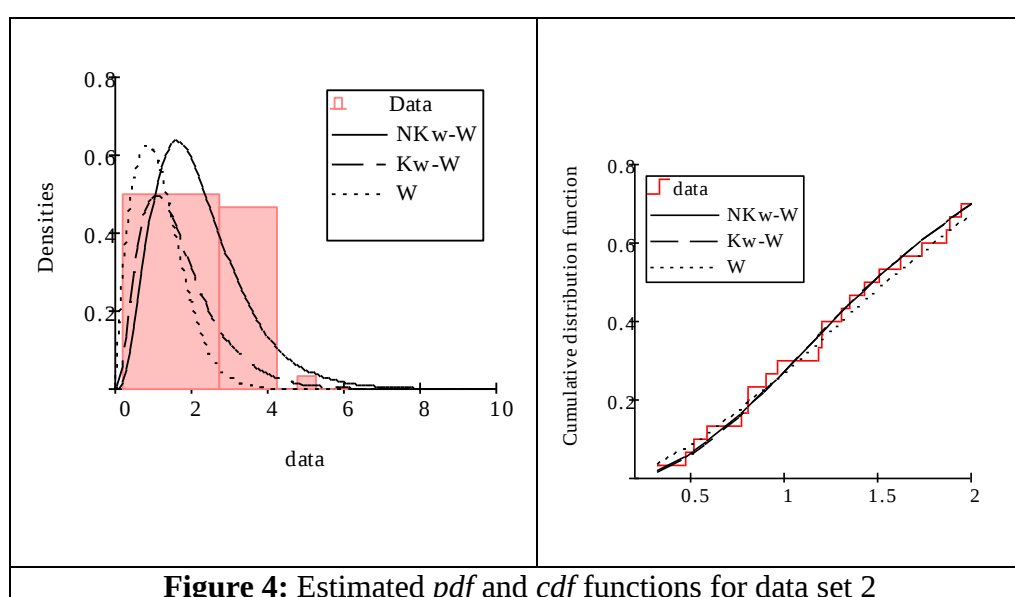
Model	$-2\ln L$	AIC	BIC	CAIC	HQIC	K-S	p-value
$NKw - W$	70.389	78.389	83.994	79.989	80.182	0.061	1
$Kw - W$	79.214	87.214	92.819	88.814	89.007	0.065	1
W	79.821	83.821	86.623	84.265	84.717	0.075	0.996

7.2. Data set 2

The following data is obtained from (Hinkley, (1977)), it consists of thirty successive value of March precipitation (in inches) in Minneapolis/St Paul:

0.77, 1.74, 0.81, 1.20, 1.95, 1.20, 0.47, 1.43, 3.37, 2.20, 3.00, 3.09, 1.51, 2.10, 0.52, 1.62, 1.31, 0.32, 0.59, 0.81, 2.81, 1.87, 1.18, 1.35, 4.75, 2.48, 0.96, 1.89, 0.90, 2.05

Figure 4 provides the plots of estimated densities and estimated cumulative of the fitted $NKw - W$, $Kw - W$ and W models for data set 2.

**Figure 4:** Estimated *pdf* and *cdf* functions for data set 2

The following table shows the MLEstimates and S.E of the model parameters for data set 2.

Table 5. The ML estimates and S.E of the model parameters for data set 2

Model	ML estimates and (S.E)			
$NKwW(a, b, \alpha, \beta)$	123.869 (189.603)	34.26 (77.735)	1.458 (0.366)	0.07 (0.046)
$KwW(a, b, \alpha, \beta)$	2.582 (2.612)	0.401 (3.836)	1.579 (9.589)	1.311 (1.955)
$W(\alpha, \beta)$	- -	- -	0.315 (0.091)	1.809 (0.249)

The following table gives the values of mesurments for the data set 2

Table 6: Model Section Criteria for data set 2

Model	$-2\ln L$	AIC	BIC	$CAIC$	$HQIC$	$K-S$	$p-value$
$NKw-W$	57.31	65.31	70.914	66.91	67.103	0.063	1
$Kw-W$	76.088	84.088	89.693	85.688	85.881	0.064	1
W	77.287	85.287	84.089	85.731	82.183	0.069	0.998

The values in Tables 4 and 6, indicate that the $NKw-W$ distribution is a strong competitor among other distributions used here for fitting data sets 1 and 2. Some plots of the estimated pdf and cdf compare the fitted densities of the models with the empirical histogram of the observed data (Figures 3 and 4). The fitted density for $NKw-W$ model is the closest to the empirical histogram than the other fitted models.

8. Concluding Remark

In the present paper, the New Kumaraswamy-Generated ($NKw-G$) family of distributions based on a new transformation is proposed. More specifically, the $NKw-G$ family covers several new distributions. We wish a broadly statistical application in some area for this new generation. Some characteristics of the $NKw-G$, such as, useful expansions for the density function and explicit expressions for the quantile function, probability weighted moments, moments, and Rényi entropy are discussed. The maximum likelihood method is employed for estimating the model parameters. Applications to two real data sets validate the priority of the new family.

References

- Afify, A. Z., Cordeiro, G. M., Yousef, H. M., Alzaatreh, A., and Nofal, Z. M. (2016). The Kumaraswamy transmuted-G family of distributions: properties and applications. *Journal of Data Science*, 14(2), 245-270.
- Alizadeh, M., Emadi, M., Doostparast, M., Cordeiro, G. M., Ortega, E. M., Pescim, R. R. (2015). "A new Family of Distributions: The Kumaraswamy odd log-logistic, Properties and Applications", *Journal of Mathematics and Statistics*, 44(6), 1491-1512.
- Alzaatreh, A., Lee, C. and Famoye, F. (2013). A new method for generating families of continuous distributions, *Metron*, 71, 63–79.
- Alzaghal, A., Famoye, F. and Lee, C. (2013). Exponentiated T-X family of distributions with some applications, *International Journal of Probability and Statistics*, 2, 31–49.
- Bourguignon, M., Silva, R. B. and Cordeiro, G. M. (2014). The Weibull-G family of probability distributions, *Journal of Data Science*, 12, 53–68.
- Chakraborty, S. and Handique, L. (2017). The Generalized Marshall-Olkin-Kumaraswamy-G family of distributions, *Journal of Data Science*, 15, 391-422.
- Cordeiro, G. M., and de Castro, M. (2011). A new family of generalized distributions, *Journal of Statistical Computation and Simulation*, 81, 883–898.
- Cordeiro, G. M., Ortega, E. M. M. and Nadarajah, S. (2010). The Kumaraswamy Weibull distribution with application to failure data. *Journal of the Franklin Institute-Engineering and Applied Mathematics*, 347, 1399-1429.
- Eugene, N., Lee, C., Famoye, F. (2002). Beta-normal distribution and its applications, *Communication in Statistics- Theory and Methods*, 31(4), 497–512.
- Greenwood, J. A., Landwehr, J. M., Matalas, N. C., and Wallis, J. R. (1979). Probability weighted moments: definition and relation to parameters of several distributions expressible in inverse form, *Water Resources Research*, 15(5), 1049-1054.
- Handique, L., Chakraborty, S. and Hamedani, G. G. (2017 a). The Marshall-Olkin-Kumaraswamy-G family of distributions. *Journal of Statistical Theory and Applications*, 16(4), 427-447.
- Handique, L., Chakraborty, S., and Ali, M. M. (2017 b). Beta Generated Kumaraswamy-G family of distributions, *Pakistan Journal of Statistics*, 33(6), 467-490.
- Hassan, A. S., and Elgarhy, M. (2016 a). Kumaraswamy Weibull- Generated Family Of Distributions with Applications, *Advances and Applications in Statistics*, 48, 205-239.
- Hassan, A. S., and Elgarhy, M. (2016 b). A New Family of Exponentiated Weibull-Generated Distributions, *International Journal of Mathematics and its applications*, 4, 135-148.

- Hassan, A. S., and Hemeda, S. E. (2016). A New Family of Additive Weibull-Generated Distributions, *International Journal of Mathematics and its Applications*, 4(2), 151-164.
- Hassan, A. S., and Nasser, S. G. (2018). Power Lindley-G Family of Distributions, *Annals of Data Science*, 1-22.
- Hassan, A. S., Elgarhy, M., and Shakil, M. (2017a). Type II Half Logistic Family of Distributions with Applications, *Pakistan Journal of Statistics and Operation Research*, 13(2), 245-264.
- Hassan, A. S., Hemeda, S. E., Maiti, S. S., and Pramanik, S. (2017b). The Generalized Additive Weibull-G Family of Distributions, *International Journal of Statistics and Probability*, 6(5), 65-83.
- Hinkley (1977). On quick choice of power transformations, *Journal of the Royal Statistical, Series (c), Applied Statistics*, 26, 67-69.
- Kumaraswamy, P. (1980). A generalized probability density functions for double-bounded random processes, *Journal of Hydrology*, 46, 79-88.
- Mahmoud, M. R., El-Sherbini, E. A., and Ahmed, M. A., and (2015). The new Kumaraswamy Kumaraswamy family of generalized distributions with application. *Pakistan Journal of Statistic and Operation Research*, 11(2), 159-180.
- Murthy, D. N. P., Xie, M. and Jiang. R. (2004). Weibull Models, John Wiley, New Jersey.
- Ramos, M. W., Marinho, P. R., Cordeiro, G. M., da Silva, R., and Hamedani, G. G. (2015). The Kumaraswamy-G Poisson Family of Distributions, *Journal of Statistical Theory and Applications*, 14(3), 222-239.
- Silva, F. G., Percontini, A., de Brito, E., W. Ramos, M., Venancio, R., and Cordeiro, G. M. (2017). The odd lindley-G family of distributions, *Austrain Journal of Statistics*, 46, 65-87.
- Tahir, M.H., Cordeiro, G.M., Alzaatreh, A., Mansoor, M., and Zubair, M. (2016). The Logistic-X family of distributions and its applications, *Communications in Statistics–Theory and Methods*, 45(24), 7326-7349.
- Zografos K. and Balakrishnan, N. (2009). On families of beta- and generalized gamma generated distributions and associated inference, *Statistical Methodology*, 6, 344–362.

Right Truncated Lomax Distribution: Properties and Estimation

¹Mohamed A. H. Sabry and ²A. Elsehetry

^{1,2}Faculty of Graduate Studies for Statistical Research,
Cairo University, Giza, Egypt

Abstract

In this paper, we propose the truncated Lomax distribution. Fundamental properties of the new distribution are studied such as; moments, moment generating and characteristic function, quantile function, incomplete moments, Lorenz and Bonferroni curves, order statistics and Rényi entropy. Maximum likelihood estimators are derived in case of complete sample, Type I and Type II censored samples. An approximate confidence interval of the parameters is obtained for large sample sizes. Simulation issue is executed in order to investigate the performance of estimators.

Keywords: *Lomax distribution; Maximum likelihood method; Type I censoring and Type II censoring.*

1. Introduction

A truncated distribution is defined as a conditional distribution that results from restricting the $f(t|a \leq T < b)$ domain of the statistical distribution. Hence, truncated distributions are used in cases where occurrences are limited to values which lie above or below a given threshold or within a specified range.

Let T be a random variable from a distribution with a probability density function (pdf), say $f(t)$, cumulative distribution function (cdf), say $F(t)$, and the range of the support $(-\infty, \infty)$. The density function of T defined in $a < T < b$ is given by

$$f(t|a \leq T < b) = \begin{cases} \frac{g(t)}{G(b) - G(a)} & a \leq t < b \\ 0 & \text{otherwise,} \end{cases} \quad (1)$$

(see, Lawless (2003)). Because is scaled up to account for the probability of being in the restricted support, this function is a density function. The restriction can occur either on a single side of the range which is called singly truncated or on both sides of the range which is called doubly truncated. If occurrences are limited to values which lie below a given threshold, the lower (left) truncated distribution is obtained. Similarly, if occurrences are limited to values which lie above a given threshold, the upper (right) truncated distribution arises.

Several truncated distributions have been provided by many authors. Some of the recent truncated distributions are; truncated Weibull distribution (Kantar and Usta, 2015); doubly truncated Fréchet distribution (Abid, 2016); generalized exponential truncated negative binomial distribution (Jayakumar and Sankaran, 2017); truncated inverted generalized exponential distribution (Genç, 2017); right truncated normal distribution (Thomopoulos, 2018), truncated Weibull Fréchet distribution (Hassan *et al.*, 2019 a), and truncated power Lomax distribution (Hassan *et al.*, 2019 b) among others.

Lomax (1954) provided a heavy-tail probability distribution, the so called Lomax (L) distribution, often used in business, economics, and actuarial modeling. It is widely applied in many areas, for instance, analysis of income and wealth data, modeling business failure data, biological sciences, model firm size and queuing problems (see Harris, 1968; Atkinson and Harrison, 1978; and Hassan and Al-Ghamdi, 2009). The (cdf) and (pdf) of the Lomax distribution are given, respectively, by

$$F_L(t; \alpha, \lambda) = 1 - \lambda^\alpha (\lambda + t)^{-\alpha}, \quad t > 0, \quad (2)$$

and

$$f_L(t; \alpha, \lambda) = \alpha \lambda^\alpha (\lambda + t)^{-(\alpha+1)}, \quad (3)$$

where, $\alpha, \lambda > 0$ are the shape and scale parameters(Ps) respectively.

In this article, we are motivated to define a new truncated distribution referred to right truncated L (RTL) distribution. We obtain some main properties of the new distribution and discuss maximum likelihood estimation of its Ps based on complete and censored samples. This article can be organized as follows. The pdf, cdf, and hazard rate function (hrf) of the RTL model are defined in Section 2. Section 3 provides some statistical properties of RTL distribution. The maximum likelihood estimators and simulation study are provided in Section 4. At the end, concluding remarks are outlined in Section 5.

2. Right Truncated Lomax Distribution

In this section, we introduce the information of the RTL distribution

Definition: A random variable X is said to have the RTL distribution (or $[0,1]$ RTL distribution) with shape Ps α and scale parameter $\lambda=1$, if its pdf is constructed by employing (1) for $a=0, b=1$ with cdf (2) and pdf (3) as follows

$$f_{RTL}(x; \alpha) = \frac{g(x; \alpha, 1)}{G(1; \alpha, 1) - G(0; \alpha, 1)} = \frac{\alpha(1+x)^{-(\alpha+1)}}{1 - 2^{-\alpha}}, \quad 0 < x < 1. \quad (4)$$

A random variable X with distribution (4) is denoted by $X \sim \text{RTL}(\alpha)$. The cdf related

to (4) is given by

$$F_{RTL}(x; \alpha) = \frac{G(x; \alpha, 1) - G(0; \alpha, 1)}{G(1; \alpha, 1) - G(0; \alpha, 1)} = \frac{1 - (1+x)^{-\alpha}}{1 - 2^{-\alpha}}. \quad (5)$$

Depending on pdf (4) and cdf (5), the survival function and hazard rate function (hrf) are given by

$$\bar{F}_{RTL}(x; \alpha) = \frac{(1+x)^{-\alpha} - 2^{-\alpha}}{1 - 2^{-\alpha}},$$

and,

$$h_{RTL}(x; \alpha) = \frac{\alpha(1+x)^{-(\alpha+1)}}{(1+x)^{-\alpha} - 2^{-\alpha}}.$$

A variety of possible shapes of the pdf and the hrf of the TL distribution for some choices of values of parameter are represented in Figure 1.

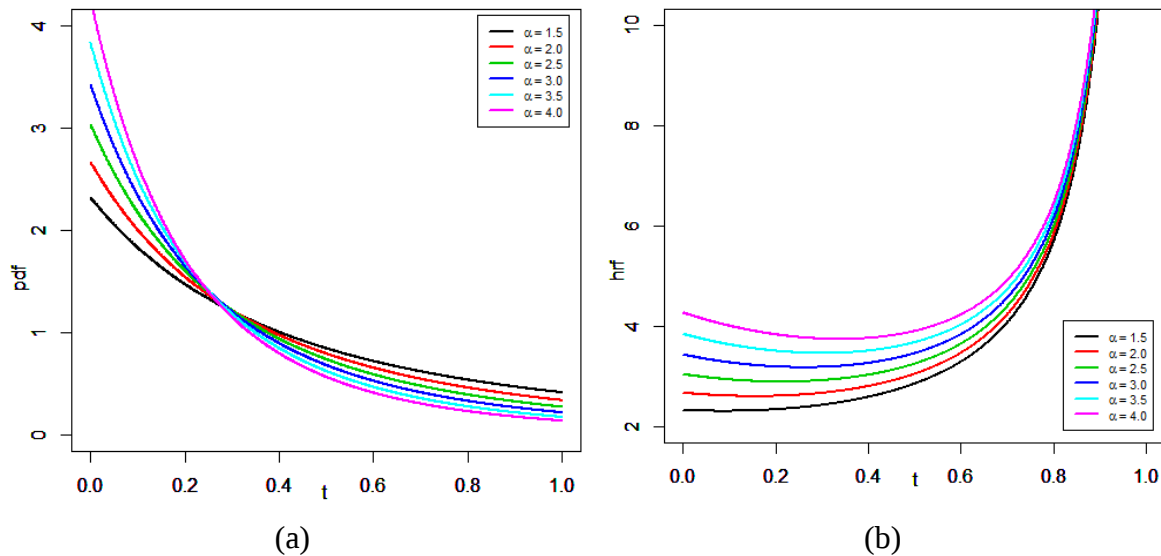


Figure 1. The pdf and hrf of RTL distribution for different values of parameters

It can be detected from Figure 1 that the pdf shape can be uni-model, reversed J-shaped and right skewed. Also, the shape of the hrf of the TL distribution could be increasing, and J -shaped.

In addition, the reversed hrf and cumulative hrf are obtained respectively as follows

$$\tau_{RTL}(x; \alpha) = \frac{\alpha(1+x)^{-(\alpha+1)}}{1-(1+x)^{-\alpha}},$$

and,

$$H_{RTL}(x; \alpha) = -\ln \left(\frac{(1+x)^{-\alpha} - 2^{-\alpha}}{1 - 2^{-\alpha}} \right).$$

The quantile function denoted by $Q(u)$ of random variable X has a cdf (5) is obtained as follows:

$$Q(u) = \left[1 - u(1 - 2^{-\alpha}) \right]^{\frac{-1}{\alpha}} - 1, \quad (6)$$

where, the random variable u belongs to the uniform distribution on $[0, 1]$. The 1st, 2nd, and 3rd quartiles are obtained from (6) by taking and $u = 0.25, 0.5$ and 0.75 respectively.

3. Basic Properties

In this section we give some statistical properties of RTL distribution.

3.1 Moments

Since the moments are necessary and important in any statistical analysis. So, we derive the r^{th} moment of the RTL distribution. If X has the pdf (4), then μ'_r is obtained as follows:

$$\mu'_r = \int_0^1 x^r f_{RTL}(x; \alpha) dx = \frac{\alpha}{1 - 2^{-\alpha}} \int_0^1 x^r (1+x)^{-(\alpha+1)} dx. \quad (7)$$

We employ the following binomial expansion:

$$(1+Z)^{-\theta} = \sum_{j=0}^{\infty} (-1)^j \binom{\theta+j-1}{j} Z^j, \quad (8)$$

for $(1+x)^{-(\alpha+1)}$ as follows:

$$(1+x)^{-(\alpha+1)} = \sum_{j=0}^{\infty} (-1)^j \binom{\alpha+j}{j} x^j. \quad (9)$$

Substituting (9) in (7) then,

$$\mu'_r = A^* \int_0^1 x^{r+j} dx = \frac{A^*}{r+j+1}, \quad (10)$$

$$\text{where, } A^* = \frac{\alpha}{1-2^{-\alpha}} \sum_{j=0}^{\infty} (-1)^j \binom{\alpha+j}{j}.$$

Furthermore, the moment generating function (mgf) and characteristic function (chf) of the RTL distribution respectively are derived as follows

$$M_x(t) = E(e^{tx}) = \sum_{r=0}^{\infty} \frac{t^r}{r!} \mu'_r = \sum_{r=0}^{\infty} \frac{t^r}{r!} \frac{A^*}{r+j+1},$$

and,

$$\phi_x(t) = E(e^{itx}) = \sum_{r=0}^{\infty} \frac{(it)^r}{r!} \frac{A^*}{r+j+1}.$$

3.2 Incomplete moments and Inequality measures

The incomplete moments are used in many statistical areas especially in the income distribution for measuring inequality like; income quintiles, the Lorenz curve, Pietra and Gini measures of inequality. The s^{th} incomplete moment of the RTPL distribution is obtained as follows

$$\Phi_s(t) = \int_0^t x^s f_{RTL}(x; \alpha) dx = A^* \int_0^t x^{s+j} dx,$$

then,

$$\Phi_s(t) = A^* \frac{t^{s+j+1}}{s+j+1}.$$

Therefore, inequality measures are often calculated for distributions other than expenditure. The Lorenz $[L_F(z)]$ and Bonferroni $[B_F(z)]$ curves with

are calculated as below

$$L_F(z) = \frac{\int_0^z x f_{RTL}(x; \alpha) dx}{E(Z)} = Z^{1+j},$$

and

$$B_F(Z) = \frac{L_F(Z)}{F_{RTL}(Z; \alpha)} = \frac{Z^{1+j} 1 - 2^{-\alpha}}{1 - (1+Z)^{-\alpha}}.$$

3.3 Rényi entropy

Entropy is a measure of the variation of the uncertainty associated with a distribution of a random variable X . The entropy of a random variable X is defined by

$$I_\delta(x) = \frac{1}{1-\delta} \log \left[\int_R f(x)^\delta dx \right], \quad (11)$$

where, $\delta > 0, \delta \neq 1$. Based on pdf (4), $f_{RTL}(x; \alpha)^\delta$ can be formed as follows:

$$f_{RTL}(x; \alpha)^\delta = \left(\frac{\alpha}{1-2^{-\alpha}} \right)^\delta (1+x)^{-\delta(\alpha+1)}. \quad (12)$$

Thus, the Rényi entropy of RTL distribution is obtained by substituting (12) in (11) as follows

$$I_\delta(x) = \frac{1}{1-\delta} \log \left[\frac{\alpha^\delta}{(1-2^{-\alpha})^\delta} \int_0^1 (1+x)^{-\delta(\alpha+1)} dx \right] \quad (13)$$

Then the employing binomial expansion (8) in (13), then we get

$$I_\delta(x) = \frac{1}{1-\delta} \log \left[\sum_{i=0}^{\infty} (-1)^i \binom{\delta(\alpha+1)+i-1}{i} \frac{\alpha^\delta}{(1-2^{-\alpha})^\delta} \int_0^1 x^i dx \right].$$

Hence, the Rényi entropy of RTL distribution is given by

$$I_\delta(x) = \frac{1}{1-\delta} \log \left[\sum_{i=0}^{\infty} \frac{w_i}{i+1} \right],$$

where,

$$w_i = (-1)^i \binom{\delta(\alpha+1)+i-1}{i} \frac{\alpha^\delta}{(1-2^{-\alpha})^\delta}.$$

3.4 Order statistics

Let $X_{(1)} < X_{(2)} < \dots < X_{(n)}$ be the order statistics of a random sample of size n then, the pdf of the k^{th} order statistic is given by:

$$f_{X_{(k)}}(x) = \frac{f(x)}{B(k, n-k+1)} \sum_{v=0}^{n-k} (-1)^v \binom{n-k}{v} F(x)^{v+k-1} \quad (14)$$

where, $B(.,.)$ is the beta function. The pdf of the k^{th} order statistic of RTL distribution is obtained by substituting (4) and (5) in (14) as follows

$$f_{RTL_{X(k)}}(x; \alpha) = \sum_{v=0}^{n-k} \frac{\alpha(-1)^v \binom{n-k}{v}}{(1-2^{-\alpha})^{v+k} B(k, n-k+1)} (1+x)^{-\alpha-1} (1-(1+x)^{-\alpha})^{v+k-1}.$$

Employing the following binomial expansion

$$(1-Z)^\theta = \sum_{u=0}^{\theta} (-1)^u \binom{\theta}{u} Z^u,$$

for $(1-(1+x)^{-\alpha})^{v+k-1}$ as follows

$$(1-(1+x)^{-\alpha})^{v+k-1} = \sum_{u=0}^{v+k-1} (-1)^u \binom{v+k-1}{u} (1+x)^{-\alpha u}$$

$$f_{RTL_{X(k)}}(x; \alpha) = \sum_{v=0}^{n-k} \sum_{u=0}^{v+k-1} \frac{\alpha(-1)^{v+u} \binom{n-k}{v} \binom{v+k-1}{u}}{(1-2^{-\alpha})^{v+k} B(k, n-k+1)} (1+x)^{-\alpha(u+1)-1}$$

By using binomial expansion in (8), we employ

$$f_{RTL_{X(k)}}(x; \alpha) = \sum_{v=0}^{n-k} \sum_{u=0}^{v+k-1} \sum_{m=0}^{\infty} \frac{\alpha(-1)^{v+u+m} \binom{n-k}{v} \binom{v+k-1}{u} \binom{\alpha(u+1)+m}{m}}{(1-2^{-\alpha})^{v+k} B(k, n-k+1)} (x)^m.$$

Hence, the pdf of the k^{th} order statistic of RTL distribution is given by

$$f_{RTL_{X(k)}}(x; \alpha) = \sum_{m=0}^{\infty} w_m x^m,$$

where,

$$w_m = \sum_{v=0}^{n-k} \sum_{u=0}^{v+k-1} \frac{\alpha(-1)^{v+u+m} \binom{n-k}{v} \binom{v+k-1}{u} \binom{\alpha(u+1)+m}{m}}{(1-2^{-\alpha})^{v+k} B(k, n-k+1)}.$$

Further, the r^{th} moment of k^{th} order statistics for RTL distribution is given by

$$\begin{aligned} E(X_{(k)}^r) &= \sum_{m=0}^{\infty} w_m \int_0^1 x^{r+m} dx \\ &= \sum_{m=0}^{\infty} \frac{w_m}{r+m+1}. \end{aligned}$$

4. Estimation and Simulation Study

In this section, maximum likelihood (ML) estimators of the model parameters are derived in case of complete, Type I and Type II censored samples. Approximate confidence intervals are also obtained. Further, numerical study is provided.

4.1 Parameter Estimation Based on Censoring Samples

In reliability or lifetime testing experiments, most of the encountered data are censored due to various reasons such as time limitation, cost or other resources. Here, we discuss estimation of population Ps of the RTL distribution based on two censoring schemes; namely, Type I and Type II. In Type-I censoring (TIC), we have a fixed time say; ω , but the number of items fail during the experiment is random. Whereas, in Type-II censoring (TIIC) scheme, the experiment is continued (i.e. time varies) until the specified number of failures c occurs.

4.1.1 ML estimators in case of TIC

Suppose that $X_{(1)} < X_{(2)} < \dots < X_{(n)}$ be a TIC sample of size n whose life time's follow RTL distribution (4) are placed on a life test and the test is terminated at determined time ω before all n items have failed. The number of failures c and all failure times are random variables. The likelihood function, based on TIC for RTL model is given by

$$L_1 = \frac{n!}{(n-c)!} \left[1 - \frac{1-S_{\omega}^{-\alpha}}{1-2^{-\alpha}} \right]^{n-c} \left\{ \prod_{i=1}^c \left[\frac{\alpha S_i^{-(\alpha+1)}}{1-2^{-\alpha}} \right] \right\},$$

where, $S_i = (1+x_i)$, and $S_{\omega} = (1+\omega)$, also for simplicity we write S_i instead of $S_{(i)}$.

Then, the log-likelihood function is obtained as follows

$$LnL_1 = \ln \left[\frac{n!}{(n-c)!} \right] + (n-c) \ln \left[1 - \frac{1-S_{\omega}^{-\alpha}}{1-2^{-\alpha}} \right] - c \ln(1-2^{-\alpha}) + c \ln \alpha - (\alpha+1) \sum_{i=1}^c \ln S_i.$$

Then, the first partial derivative of the log-likelihood with respect to the unknown Ps is given by

$$\frac{d \ln L_1}{d \alpha} = \frac{c}{\alpha} - \frac{2^{-\alpha} c \ln 2}{1-2^{-\alpha}} + \frac{(n-c) \left(-\frac{S_{\omega}^{-\alpha} \ln S_{\omega}}{1-2^{-\alpha}} + \frac{2^{-\alpha} \ln 2 (1-S_{\omega}^{-\alpha})}{(1-2^{-\alpha})^2} \right)}{1 - \frac{1-S_{\omega}^{-\alpha}}{1-2^{-\alpha}}} - \sum_{i=1}^c \ln S_i.$$

Equating this partial derivative with zero and solving simultaneously yield the ML estimator of α based on TIC samples.

4.1.2 ML estimators in case of TIIC

If we want to ensure that the resulting data set contains a fixed number c of observed lifetimes and furthermore we want to terminate the test as fast as possible. So, the design must allow for the test to terminate at the c^{th} failure such that $X_{(1)} < X_{(2)} < \dots < X_{(c)}$ be a TIIC sample of size n observed from lifetime testing experiment. The likelihood function of RTL model, based on TIIC, is given by

$$L_2 = \frac{n!}{(n-c)!} \left[1 - \frac{1-S_c^{-\alpha}}{1-2^{-\alpha}} \right]^{n-c} \left\{ \prod_{i=1}^c \left[\frac{\alpha S_i^{-(\alpha+1)}}{1-2^{-\alpha}} \right] \right\}.$$

where, $S_i = (1+x_i)$, and $S_c = (1+c)$ also for simplicity we write S_i instead of $S_{(i)}$ then, the log-likelihood function, based on TIIC, is given by

$$LnL_2 = \ln \left[\frac{n!}{(n-c)!} \right] + (n-c) \ln \left[1 - \frac{1-S_c^{-\alpha}}{1-2^{-\alpha}} \right] - c \ln(1-2^{-\alpha}) + c \ln \alpha - (\alpha+1) \sum_{i=1}^c \ln S_i.$$

Then, the first partial derivative of the log-likelihood are given by

$$\frac{d \ln L_2}{d \alpha} = \frac{c}{\alpha} - \frac{2^{-\alpha} c \ln 2}{1-2^{-\alpha}} + \frac{(n-c) \left(-\frac{S_c^{-\alpha} \ln S_c}{1-2^{-\alpha}} + \frac{2^{-\alpha} \ln 2 (1-S_c^{-\alpha})}{(1-2^{-\alpha})^2} \right)}{1 - \frac{1-S_c^{-\alpha}}{1-2^{-\alpha}}} - \sum_{i=1}^c \ln S_i.$$

Solving $\frac{d \ln L_2}{d \alpha} = 0$ numerically by using iteration technique, then the ML estimator of α is obtained via Mathematica 7.

Additionally, for $c = n$, we obtain the ML estimator under complete samples as seen in Tables 5-8.

For interval estimation of the Ps, it is known that under regularity condition, the asymptotic distribution of ML estimators of elements of unknown Ps for α is given by

$$(\hat{\alpha} - \alpha) \rightarrow N(0, I^{-1}(\alpha)),$$

where, $I^{-1}(\alpha)$ is the variance covariance matrix of unknown Ps α . The elements of Fisher information matrix are obtained for both censoring schemes. Therefore, the two-sided approximate γ 100 percent limits for the ML estimates of a population Ps for α can be obtained, as follows:

$$L_{\alpha} = \hat{\alpha} - z_{\frac{\gamma}{2}} \sqrt{\text{var}(\hat{\alpha})}, \quad U_{\alpha} = \hat{\alpha} + z_{\frac{\gamma}{2}} \sqrt{\text{var}(\hat{\alpha})},$$

where z is the $100(1-\gamma/2)\%$ th standard normal percentile and $\text{var}(\cdot)$'s denote the diagonal elements of variance covariance matrix corresponding to the model Ps.

4.2 Simulation Studies

Here, we provide numerical study to evaluate the behavior of the ML estimates of the RTL based on complete sample, TIC and TIIC schemes. The algorithm used here is outlined as follows:

- 1000 random sample of sizes $n = 30, 50$ and 100 are generated from the RTL distribution under TIC and TIIC.
- Exact values of Ps are chosen as; $\alpha = 0.7, 0.9, 1.5$ and 2.5 .
- Three termination times are selected as $\omega = 0.7, \omega = 0.9$ and $\omega = 1$ based on TIC and the number of failure items; c , based on TIIC are selected as 70%, 90% and 100% (complete sample).
- The following measures are calculated
 - i. Average ML of the simulated estimates $\hat{\alpha}_i, i = 1, 2, \dots, N$, where $N = 1000$

$$\frac{1}{N} \sum_{i=1}^N \hat{\alpha}_i$$

- ii. Average bias of the simulated estimates $\hat{\alpha}_i, i = 1, 2, \dots, N$

$$\frac{1}{N} \sum_{i=1}^N (\hat{\alpha}_i - \alpha).$$

- iii. Average mean square error (MSE) of the simulated estimates $\hat{\alpha}_i, i = 1, 2, \dots, N$

$$\frac{1}{N} \sum_{i=1}^N (\hat{\alpha}_i - \alpha)^2.$$

- iv. Average length of the N simulated confidence intervals and coverage probability with confidence level $\gamma = 90\%$ and 95% for $\hat{\alpha}_i, i = 1, 2, \dots, N$ are calculated.

- Numerical results are recorded in Tables 1, 2, 3, ... and 8.

Table 1: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under TIC for $\alpha = 0.7$

n	ω	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	0.7	0.68956	-0.01044	0.89256	3.06796	3.65544	90.50	94.60
	0.9	0.69320	-0.00680	0.88792	3.05156	3.63590	90.10	94.90
	1	0.69332	-0.00668	0.88724	3.05103	3.63527	90.30	94.70
50	0.7	0.72788	0.02787	0.54091	2.36777	2.82117	89.30	94.00
	0.9	0.72757	0.02757	0.53554	2.35531	2.80633	89.30	94.60
	1	0.72752	0.02752	0.53430	2.35489	2.80582	89.30	94.50
100	0.7	0.68429	-0.01571	0.26091	1.66761	1.98694	90.70	95.10
	0.9	0.68612	-0.01389	0.25426	1.65878	1.97642	90.80	95.20
	1	0.68623	-0.01377	0.25412	1.65852	1.97610	90.60	95.20

Table 2: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under TIC for $\alpha = 0.9$

n	ω	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	0.7	0.92065	0.02065	0.90903	3.08059	3.67050	90.50	94.50
	0.9	0.92709	0.02709	0.90163	3.06554	3.65256	90.50	94.50
	1	0.92765	0.02765	0.90054	3.06508	3.65201	90.40	94.70
50	0.7	0.89296	-0.00704	0.55493	2.37521	2.83004	89.40	95.00
	0.9	0.89514	-0.00486	0.55210	2.36333	2.81588	89.70	94.70
	1	0.89530	-0.00470	0.55187	2.36296	2.81544	89.70	94.80
100	0.7	0.90718	0.00717	0.27201	1.67430	1.99491	89.60	94.50
	0.9	0.90467	0.00467	0.26735	1.66588	1.98487	89.90	94.30
	1	0.90453	0.00453	0.26675	1.66560	1.98455	89.70	94.60

Table 3: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under TIC for $\alpha=1.5$

n	ω	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	0.7	1.52681	0.02681	0.90315	3.13147	3.73111	90.80	96.00
	0.9	1.52659	0.02659	0.89480	3.11776	3.71478	90.90	95.70
	1	1.52634	0.02634	0.89337	3.11728	3.71421	90.80	95.70
50	0.7	1.53042	0.03042	0.56505	2.41654	2.87928	89.00	94.80
	0.9	1.53219	0.03219	0.55705	2.40624	2.86700	89.50	95.20
	1	1.53166	0.03166	0.55695	2.40589	2.86659	89.50	95.10
100	0.7	1.48933	-0.01067	0.25590	1.70038	2.02598	90.80	96.30
	0.9	1.49232	-0.00768	0.25389	1.69325	2.01749	90.80	96.40
	1	1.49241	-0.00759	0.25386	1.69304	2.01724	90.60	96.40

Table 4: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under TIC for $\alpha=2.5$

n	ω	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	0.7	2.52467	0.02467	0.98968	3.27512	3.90227	91.10	95.90
	0.9	2.52726	0.02726	0.98375	3.26460	3.88973	90.50	95.80
	1	2.52709	0.02709	0.98262	3.26423	3.88929	90.60	95.70
50	0.7	2.56655	0.06655	0.62182	2.53274	3.01774	90.20	94.60
	0.9	2.56530	0.06530	0.61909	2.52431	3.00769	90.20	94.90
	1	2.56515	0.06515	0.61907	2.52406	3.00739	90.30	94.90
100	0.7	2.51467	0.01467	0.31188	1.78016	2.12104	89.10	94.40
	0.9	2.51533	0.01533	0.31028	1.77433	2.11409	89.30	94.30
	1	2.51531	0.01531	0.30986	1.77415	2.11389	89.30	94.30

Table 5: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under THC for $\alpha=0.7$

n	x_c	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	70%	0.73730	0.03730	0.95738	3.13165	3.73133	89.80	95.20
	90%	0.72034	0.02034	0.91137	3.05820	3.64381	90.10	95.30
	100%	0.71588	0.01588	0.90440	3.05278	3.63735	90.50	95.30
50	70%	0.70248	0.00248	0.58703	2.40858	2.86979	88.70	94.60
	90%	0.68772	-0.01228	0.55498	2.35681	2.80812	88.80	94.60
	100%	0.68579	-0.01422	0.55222	2.35375	2.80447	89.00	94.50
100	70%	0.71623	0.01623	0.24198	1.69386	2.01821	90.70	96.70
	90%	0.71143	0.01143	0.22932	1.66038	1.97833	91.50	96.50
	100%	0.71119	0.01119	0.22864	1.65870	1.97633	91.70	96.50

Table 6: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under THIC for $\alpha=0.9$

n	x_c	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	70%	0.91347	0.01347	1.01390	3.15057	3.75387	89.40	94.80
	90%	0.89901	-0.00099	0.93671	3.07016	3.65806	88.80	95.20
	100%	0.89593	-0.00407	0.92817	3.06403	3.65076	88.80	95.00
50	70%	0.94048	0.04048	0.55752	2.42465	2.88895	89.80	96.20
	90%	0.93610	0.03610	0.52492	2.36777	2.82118	90.10	95.50
	100%	0.93488	0.03488	0.52356	2.36416	2.81687	90.30	95.30
100	70%	0.91980	0.01980	0.30361	1.70556	2.03215	88.50	93.20
	90%	0.91073	0.01073	0.29132	1.66822	1.98766	88.10	93.30
	100%	0.91082	0.01082	0.29090	1.66630	1.98538	88.20	93.20

Table 7: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under THIC for $\alpha=1.5$

n	x_c	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	70%	1.52682	0.02682	1.00526	3.23055	3.84916	89.60	95.20
	90%	1.50604	0.00604	0.91806	3.12432	3.72260	89.60	95.10
	100%	1.50193	0.00193	0.90391	3.11505	3.71155	89.70	95.00
50	70%	1.55677	0.05677	0.61338	2.48997	2.96678	89.60	94.70
	90%	1.54111	0.04111	0.56068	2.41205	2.87393	90.90	95.00
	100%	1.54119	0.04119	0.55399	2.40661	2.86745	90.80	95.10
100	70%	1.52335	0.02335	0.27733	1.74772	2.08239	91.10	94.80
	90%	1.51580	0.01580	0.25723	1.69744	2.02249	90.50	95.10
	100%	1.51381	0.01381	0.25469	1.69432	2.01876	90.80	94.80

Table 8: ML estimates, Biases, MSE, Average length and Coverage probability of RTL distribution under THIC for $\alpha=2.5$

n	x_c	ML	Bias	MSE	Average length		Coverage probability	
					90%	95%	90%	95%
30	70%	2.52723	0.02723	1.12656	3.44250	4.10170	90.40	95.70
	90%	2.50465	0.00465	1.00251	3.27711	3.90464	90.40	95.50
	100%	2.50401	0.00401	0.99109	3.26053	3.88489	90.70	95.70
50	70%	2.56260	0.06260	0.76494	2.65612	3.16474	88.30	93.40
	90%	2.54010	0.04010	0.68577	2.53302	3.01806	88.90	93.50
	100%	2.53098	0.03098	0.66971	2.52072	3.00341	89.00	93.80
100	70%	2.53804	0.03804	0.30333	1.86151	2.21797	91.40	96.20
	90%	2.52891	0.02891	0.27135	1.78086	2.12187	91.10	96.60
	100%	2.53033	0.03033	0.27003	1.77487	2.11474	91.30	96.60

From Tables 1-8 we conclude that

- As the sample size n increases the MSE of ML estimates decrease.
- As the termination time ω increases, the MSE of estimates decreases.
- Based on Tables 1-4, we can see that as the sample size n increases the MSE of estimates decreases and also that as the censoring level time ω increases, the MSE of estimates decreases .
- Based on Tables 1 to 8, the coverage probability is very close to the intended significance level for all values of n and α .
- The average length of confidence intervals for the unknown Ps decrease as n increases.

5. Summary and Conclusion

In this paper, we introduce a new model called the truncated Lomax. We investigate several structural properties of the new distribution, expressions for the ordinary moments, generating function and order statistics. The model parameters are estimated by the maximum likelihood method in case of complete and censored samples. A simulation study reveals that the estimates of the model have desirable properties such as, (i) the maximum likelihood estimates are not too far from the true parameter values; (ii) the biases and mean square errors of estimates in case of complete sample are smaller than the corresponding in censored samples; and (iii) the biases and the mean square values decrease as the sample size increases.

References

- Abid, S.H. (2016). Properties of doubly-truncated Fréchet distribution. *American Journal of Applied Mathematics and Statistics*, 4(1), 9-15.
- Atkinson, A., and Harrison, A. (1978). *Distribution of Personal Wealth in Britain*. Cambridge University Press, Cambridge.
- Genç, A.I. (2017). Truncated inverted generalized exponential distribution and its properties. *Communications in Statistics-Simulation and Computation*, 46(6), 4654-4670.
- Hassan A. S. and Al-Ghamdi, A. (2009) Optimum step stress accelerated life testing for Lomax distribution. *Journal of Applied Sciences Research*, 5, 2153–2164.
- Hassan, A. S., Elgarhy, M., Nassr, S. G., Ahmad, Z., and Alrajhi, S. (2019 a). Truncated Weibull Fréchet Distribution: Statistical Inference and Applications. *Journal of Computational and Theoretical Nanoscience*, 16(1), 1-9.
- Hassan, A. S., Sabry, M.A., Elsehetry, A. (2019 b) Truncated power Lomax distribution with application to flood data. *Journal of Statistics Applications & Probability*; forthcoming.
- Harris, C. M. (1968). The Pareto distribution as a queue service discipline. *Operations Research*, 16 (2), 307–313.

Jayakumar, K., and Sankaran, K.K. (2017). Generalized exponential truncated negative binomial distribution. *American Journal of Mathematical and Management Sciences*, 36(2), 98-111.

Kantar, Y.M., and Usta, I. (2015). Analysis of the upper-truncated Weibull distribution for wind speed. *Energy Conversion and Management*, 96, 81-88.

Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*. 2nd Edition, Wiley, Hoboken, NJ.

Lomax, K.S. (1954). Business failures: another example of the analysis of failure data. *Journal of the American Statistical Association* 49:847-852.

Thomopoulos, N. T. (2018). Right Truncated Normal. *In Probability Distributions. Springer International Publishing* (85-97), https://doi.org/10.1007/978-3-319-76042-1_6.

Three Parameters Estimation of Log-Logistic Distribution Using Algorithm of Percentile Roots

Doaa Abd El-Shafi Abd El-Rahman^{1,2} Mohammed Mohammed El Genidy³

Abstract This study provided a new algorithm to estimate the three parameters of log-logistic distribution called algorithm of percentile roots, AOPR. The effectiveness of AOPR was tested by the Anderson-Darling test which accepted the results at a high level of significance. The estimated parameters by AOPR were very close to EasyFit program's results. In addition, they were somewhat more accurate than EasyFit program's ones. The analyzed data is a data set of daily global solar exposure for a year in Queensland, Australia.

Keywords: Three-Parameter Log-Logistic Distribution, Algorithm of Percentile Roots, Anderson-Darling Test, Daily Global Solar Exposure

1 Introduction

Data analysis is a wide field that plays a major role in any decision-making process based on a numerical survey. That made statisticians always look forward to the most accurate modeling of data. Out of this, the distribution by which the data modeled should be estimated too accurate to fit the data almost perfectly. The factors that have a real impact on data modeling are the chosen distribution and the parameters estimation method. Some researchers chose three-parameter log-logistic distribution (LLD3p) to model data of various phenomena and fields. Kantam et. al (2007) presented LLD3p in a model of collected life of given items to decide accepting or rejecting the product. Log-normal and log-logistic distributions were analyzed lifetime data in a study by Dey and Kundu (2009). Others used generalized logistic distribution (GLD) in many fields too. Asgharzadeh et al. (2013) analyzed a real data set of log times to breakdown of a fluid insulating in an accelerated test by GLD and maximum likelihood estimation. While, Shin et al. (2012) applied generalized extreme value (GEV) distribution and GLD to data of annual maximum rainfall in Korea. Likewise, GLD was fit to data in the economic sphere, as Huang et al. (2017) performed it on a data set of expected shortfall and value-at-risk in resource sector.

Parameters Estimator also affects substantially on data analysis, that's why many studies have focused on the estimation of the parameters for different distributions. For instance, Ashkar and Mahdi (2006) estimated two-parameter log-logistic distribution by generalized moments and quantiles estimators. Asgharzadeh (2006) presented maximum likelihood (ML) and approximate ML estimators for GLD. Just like Arora et al. (2010) who applied GLD with ML estimation. And another, ML, approximate ML and Bayes estimation methods were used along with Weibull distribution by Valiollahi et al. (2013). Besides, El Genidy (2019) introduced Exponentiated Gumbel Maximum distribution that was estimated by quartiles-moments method to analyze maximum temperature and solar radiation data. Lindley distribution was estimated using the methods of uniformly minimum variance unbiased and ML by Maiti and Mukherjee (2017).

¹Department of Basic Science, Faculty of Computers and Informatics, Suez Canal University, Ismailia, Egypt.

²Corresponding author, E-mail Address: doaa.abdelshafi@gmail.com.

³Department of Mathematics and Computer Science, Faculty of Science, Port Said University, Port Said, Egypt. Email: drmmg2016@yahoo.com.

Moreover, Khan et al. (2015) compared moments, energy pattern factor, ML, and empirical methods as estimators for determination of Weibull parameters to model daily wind speed data in Jiwani town along ten years. The comparison showed that the best results were those that obtained by ML. Another comparison was made by Kaoga et al. (2014) between five numerical estimators. They were the empirical estimation, ML method, modified ML method, the graphical estimator, and energy pattern factor methods. The last one led to the best fit to the analyzed data. Rocha et al. (2012) discussed seven estimation methods that were ML, modified ML, graphical, moments, equivalent Energy, the empirical, and energy pattern factor methods.

Based on the foregoing, the accurate estimation of parameters is very important to get the best model of any data set by a chosen probability distribution. This research introduced a new parameters estimation method, called algorithm of percentile roots (AOPR), that is commensurate to calculate the shape, scale and location parameters of LLD3p. This estimator could combine percentile equations with the measures of central tendency, that made the estimation more accurate than EasyFit program estimation. The data set that was used to apply LLD3p with AOPR was daily global solar exposure in Queensland throughout one year. AOPR's results were checked by Anderson-Darling test and they were significant.

2 Materials and Methods

This study applied a new parameter estimator called AOPR with LLD3p. The algorithm proved its efficiency compared to EasyFit program as clarified by Anderson-Darling test.

2.1 The Analyzed Data set

The data set that used to apply the new algorithm was on the daily global solar exposure in Queensland for a year, (2015) (Figure 1.). It was actual and reliable data, as it had been recorded and published by the Bureau of Meteorology in Australian government. It was posted online in <http://www.bom.gov.au/climate/data/index.shtml>.

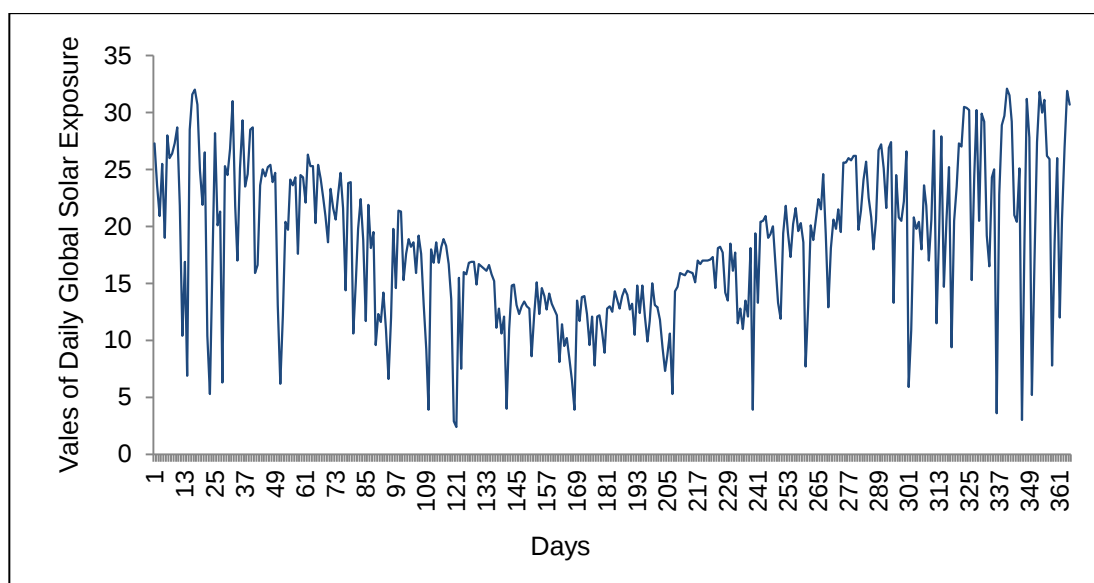


Figure 1. Data set of Daily Global Solar Exposure during 2015 at Queensland, Australia.

2.2 Three-Parameter Log-Logistic Distribution

Let X be a random variable that presented a data set of the daily global solar exposure in Queensland throughout (2015). This random variable of the same data set was matched with two-parameter Weibull distribution in a research by El Genidy and Abd El-Rahman (2019). But here, X has been studied to have LLD3p with the following form of probability density function (PDF):

$$f(x) = \frac{\left(1 + \frac{\beta}{\sigma}(x-\mu)\right)^{-\left(\frac{1}{\beta}+1\right)}}{\sigma \left(1 + \left(1 + \frac{\beta}{\sigma}(x-\mu)\right)^{\left(\frac{1}{\beta}\right)}\right)^2} \quad ; \sigma > 0, \beta > 0 \text{ and } x > a \text{ S.T } a = \mu - \frac{\sigma}{\beta} \quad (1)$$

Such that β is the shape parameter, σ is the scale parameter and μ is the location parameter.

Since the cumulative distribution function (CDF) is given by:

$$F(x) = \int_a^x f(y) dy \quad (2)$$

Then, it would be:

$$F(x) = \left(1 + \left(1 + \frac{\beta}{\sigma}(x-\mu)\right)^{\left(\frac{1}{\beta}\right)}\right)^{-1} \quad (3)$$

The mean is the expected value of x , and it is obtained by the equation:

$$M = E(x) = \int_a^{\infty} x f(x) dx \quad (4)$$

Which implies:

$$M = \mu + \frac{\sigma}{\beta} (\pi \beta \csc(\pi \beta) - 1) \quad (5)$$

And the median is the value of x that satisfies:

$$F(x) = 0.5 \quad (6)$$

$$\left(1 + \left(1 + \frac{\beta}{\sigma}(x-\mu)\right)^{\left(\frac{1}{\beta}\right)}\right)^{-1} = 0.5$$

$$\left(1 + \frac{\beta}{\sigma}(x-\mu)\right)^{\left(\frac{1}{\beta}\right)} = 1$$

$$\frac{\beta}{\sigma}(x-\mu) = 0$$

And $\frac{\beta}{\sigma}$ couldn't be zero, which means the median would be:

$$\text{Median} = x = \mu \quad (7)$$

2.3 Algorithm of Percentile Roots

The algorithm, which was created in this research, AOPR, depended on solving three equations together to get approximated vales for the three parameters of LLD3p. The equations were CDF y_i at any point x_i , the mean and the median. This technique combined the percentile roots with some important common measures of central tendency, such as mean and median.

By using Eq.7, the median of the used data set gave the value of the location parameter μ , as μ =the median. From Eq.5, the used data set's mean was used to simplify it as follows:

$$\sigma = \frac{\beta(M-\mu)}{(\pi\beta \csc(\pi\beta)-1)} \quad (8)$$

Then, Eq.5 was used with an arbitrary percentile equation of the used data set at any point x_i , $y_i=F(x_i)$, to get the value of the shape parameter β_i as follows:

$$y_i = \left(1 + \left(1 + \frac{(x_i-\mu)(\pi\beta_i \csc(\pi\beta_i)-1)}{M-\mu} \right)^{\left(\frac{-1}{\beta_i}\right)} \right)^{-1} \quad (9)$$

Finally, the scale parameter σ_i could be calculated by Eq.5 after obtaining β_i .

This operation would be iterated arbitrary number of times N , such that $i=1,2,3 \dots N$, so that N values of the shape and the scale parameters would be obtained. The median could be calculated for the N values of each parameter to get the final estimation.

2.4 Algorithm Statement

Input:

N the number of percentile equations, $X[0..N-1]$ array of x values, $Y[0..N-1]$ array of y values, M the mean, μ the median

Output:

A the median of β values, B the median of σ values

Procedure:

$N \leftarrow \text{read}()$

for $i \leftarrow 0$ to $N-1$ do

{

$X[i] \leftarrow \text{read}()$, $Y[i] \leftarrow \text{read}()$, $M \leftarrow \text{read}()$, $\mu \leftarrow \text{read}()$

```

FindRoot $\left(\left(1+\left(1+\frac{(x-18.1)(\pi\beta \csc(\pi\beta)-1)}{0.2}\right)^{\left(\frac{-1}{\beta}\right)}\right)^{-1}==y,\{\beta,0.01\}\right)$ 


$$\sigma=\frac{\beta(M-\mu)}{(\pi\beta \csc(\pi\beta)-1)}$$


J[i] <-  $\beta$ 
K[i] <-  $\sigma$ 

}

A = median(J)
B = median(K)
print A, B

```

2.5 Goodness-of-Fit Test

The fitness of AOPR's results with the used data set was measured by Anderson-Darling test. It's defined as:

$$A^2 = - \sum_{i=1}^n \left(\frac{2i-1}{n} \left[\ln(F(x_i)) + \ln(1-F(x_{n+1-i})) \right] \right) - n \quad (10)$$

2.6 Software

The programs carried out in this study were EasyFit professional version 5.5 and Mathematica 8 version 8.0.1. The first was released in February (2010) and available from MathWave Technologies, <http://www.mathwave.com>. And the else was released in March 2011 and available from Wolfram Mathematica, <http://www.wolfram.com>.

3 Results

3.1 Estimation method

Calculating the measures of central tendency of the cited data set of daily global solar exposure in Queensland along 365 days. Therefore, from Eq. 5,7 and 8, the mean and the median were as follows:

$$M = \mu + \sigma \beta \pi \csc(\pi \beta) - 1 \quad M = \mu + \frac{\sigma}{\beta} (\pi \beta \csc(\pi \beta) - 1) = 18.3 \quad (11)$$

$$\text{Median} = \mu = 18.1 \quad (12)$$

$$y_i = \left(1 + \left(1 + \frac{(x-18.1)(\pi\beta_i \csc(\pi\beta_i)-1)}{0.2} \right)^{\left(\frac{-1}{\beta_i}\right)} \right)^{-1} \quad (13)$$

By solving Eq.13 using Mathematica program at percentile equations $y_1=0.005$, $y_2=0.01$, $y_3=0.015$, $y_4=0.02$, $y_5=0.025$..., $y_{99}=0.995$ and there corresponding values $x_1, x_2, x_3, \dots, x_{99}$ as follows:

$$\text{FindRoot} \left[\left(1 + \left(1 + \frac{(x-18.1)(\pi\beta \csc(\pi\beta)-1)}{0.2} \right)^{\left(\frac{-1}{\beta}\right)} \right)^{-1} == y, \{\beta, 0.01\} \right]$$

Hence, 197 values of the parameter β could be obtained, excluding $y_{100}=0.5$ as it's the median equation, and then their medians were calculated as shown in Table1.

Table 1. Numerical values of the shape and the scale parameters obtained by AOPR.

x_i	y_i	β_i	σ_i
2.982	0.005	0.03845	3.15403
3.792	0.01	0.03594	3.37521
3.9	0.015	0.0334	3.6331
4.336	0.02	0.03227	3.76033
5.3	0.025	0.03275	3.70469
5.852	0.03	0.03258	3.7238
.	.	.	.
.	.	.	.
.	.	.	.
.	.	.	.
30.516	0.97	0.03622	3.34898
30.7	0.975	0.03786	3.20369
31.072	0.98	0.03935	3.08195
31.362	0.985	0.04184	2.89754
31.672	0.99	0.0457	2.65217
31.918	0.995	0.05369	2.25525
Median		0.02862	4.24075

Therefore, AOPR estimated the parameters as 18.1, 0.02862 and 4.24075 for the location, the shape and the scale parameters, respectively. Formula of the CDF of LLD3p became:

$$F(x) = \left(1 + \left(1 + \frac{0.02862}{4.24075} (x-18.1) \right)^{\left(\frac{-1}{0.02862}\right)} \right)^{-1} \quad (14)$$

3.2 Checking the validity of AOPR

Anderson-Darling test was performed in this research to measure the validity of AOPR's results, as it's considered one of the most effective statistical tests. Applying Eq. 10, the test statistic value of AOPR's results was $A^2=1.32095$ which was less than the EasyFit program's one, $A^2=1.4144$ (Table 2.). This means AOPR gave an estimated LLD3p that was closer to the actual data than the other one produced by EasyFit program.

Table 2. *Estimated parameters by AOPR and EasyFit program and their and Anderson-Darling test values.*

	EasyFit Program	AOPR
μ	18.21	18.1
β	0.01374	0.02954
σ	3.85913	4.1085
A^2	1.4144	1.32095

So, AOPR proved its ability of estimating accurate parameters for LLD3p as well as the other traditional known estimation methods.

4 Discussion

This study is concerned about finding a new estimator of LLD3p. This estimator, AOPR, is characterized by its ease of use, accuracy and its fitness with any three-parameter distribution. According to Anderson-Darling test, AOPR produced significant results when it was used to model real data of daily global solar exposure with LLD3p. Moreover, the results were slightly more accurate than results obtained by EasyFit program when using the same data, distribution and statistical test.

In the previous studies, many types of logistic distributions had been performed to fit real data in various fields with different estimators of parameters. Due to its relative simplicity and flexibility. GLD was performed side by side with Weibull, gamma, and log-normal distributions with ML estimation to analyze data of light, deep and rapid eye movement sleep stages by Saat and Jemain (2009). Besides, Gupta and Kundu (2010) discussed some properties of the proportional reversed hazard logistic distributions and the skew logistic distribution with proposing data of the measured strength in grade point average for single carbon fibers of a specific length. Ahmed (2018) modeled a data set of the monthly closing values of indices for the entire composite index using GLD with two different estimation methods. In another study by Kantam et. al (2007), LLD3p was estimated to model the life of a specific product in order to accepting or rejecting it. Also, LLD3p was also able to represent well solar exposure data in the present study.

Turning to the role of parameters estimation method, several estimators are considered good generators of parameters, as ML method, method of moments, least squares method and the methods depends on percentiles. According to Soukissian and Tsalis (2015), data of extreme wind speed is analyzed by applying GEV distribution with nine different estimation methods such as moments, ML, ordinary entropy and quantile least squares methods. El Genidy and Ali (2016) applied method of moments with GEV distribution on data of pollutants ozone amount. Further,

Khan and King (2013) estimated the parameters of modified Weibull distribution in which the estimation of parameters was done by ML estimation as well as least squares estimation.

Moreover, Alkawasbeh and Raqab_ (2009) estimated shape and scale parameters of GLD by the methods of ML, moments, least squares, weighted least squares, L-moment. They conducted that the best fit was obtained by least squares estimations. Ahmed (2018) made a comparison between the estimation of GLD extreme quantiles by moments and probability weighted moments methods, which conducted that moments method gave a higher estimation. Later on, El Genidy and Abd El-Rahman (2019) presented nested percentiles algorithm as a higher accurate estimator with Weibull distribution. Many other studies, that made by Ahmed et al. (2012); El Genidy (2012); Agboola et al. (2018) and Wasinrat et al. (2013), dealt with ML method as a good estimator with various distributions.

Overall, ML is considered a common and widely used estimator, even though it is quite hardly applied compared to AOPR. However, AOPR derived satisfactory results. It deals with too many values (200) of actual data, unlike some other methods which deal with limited values just as moments method. Also, AOPR combines the percentiles values and two impactful central tendency measures like mean and median, while nested percentiles algorithm solves only the percentiles equations. As documented, AOPR is an appropriate estimator for any distribution that has three parameters.

Since the goodness of fit test is the main measure of the accuracy of any estimator, there have been many developments on statistical tests. For example, Razali and Yap (2011) had ranked Anderson-Darling test second after Shapiro-Wilk test in a study compared between certain statistical tests. In another study, four goodness-of-fit tests, including Anderson-Darling test, were carried out to fit WD after predicting the parameters by Poudel and Cao (2013). Likewise, Shin et al. (2012) conducted modified Anderson–Darling test to have the highest efficiency with GLD and GEV distribution compared to other traditional tests such as Chi-Square, Kolmogorov–Smirnov, and Cramer Von Mises tests. It was used to fit GEV distribution and GLD with annual maximum rainfall data. Selectively, in this study, Anderson-Darling test was applied for fitting AOPR estimator with the analyzed data set. It has been concluded that AOPR is a suitable estimator which is easy to apply and gives significant results.

Finally, it's recommended to simulate AOPR on other three-parameter probability distributions, such as gamma, log-normal and GEV distributions with other data sets in different fields. Programmers can implement the algorithm to generalize its usage in the statistical analysis.

5 Conclusion

This study suggested a method, called Algorithm of percentile roots, that works as a three parameters estimator. This algorithm proved its ability to compete the other known estimation methods. As resulted from Anderson-Darling test, AOPR's outcomes were quite better suited than software's results such as EasyFit program. AOPR dealt with percentile equations along with some important measures of the central tendency such as mean and median, which gave the results more accuracy and realism. It can be applied to estimate any other three-parameter distributions in relevant fields.

Acknowledgment

The authors are thankful of Australian Bureau of Meteorology for providing online climate data of Queensland, Australia. They also thank contributing researchers in the field of study.

References

- Agboola S., Dikko H. G., Asiribo O. E., 2018. On Exponentiated Skewed Student T Error Distribution on Some Volatility Models: Evidence of Standard and Poor-500 Index Return. *Asian Journal of Applied Sciences*, 11: 38-45.
- Ahmad, M. I., 2018. Estimation of Generalized Logistic Distribution for Financial Data. *Journal of Engineering and Science Research*, 2 (2): 34-37.
- Ahmed A. O. M., Ibrahim N. A., Adam M. B., Arasan J., 2012. Bayesian Survival and Hazard Estimate for Weibull Censored Time Distribution. *Journal of Applied Sciences*, 12: 1313-1317.
- Alkasasbeh, M. R., Raqab, M. Z., 2009. Estimation of the generalized logistic distribution parameters: Comparative study. *Statistical Methodology*, 6: 262-279.
- Arora, S. H., Bhimani, G. C., Patel, M. N., 2010. Some Results on Maximum Likelihood Estimators of Parameters of Generalized Half Logistic Distribution under Type-I Progressive Censoring with Changing Failure Rate. *Int. J. Contemp. Math. Sciences*, 5 (14): 685 – 698.
- Asgharzadeh, A., 2006. Point and Interval Estimation for a Generalized Logistic Distribution Under Progressive Type II Censoring. *Communications in Statistics—Theory and Methods*, 35: 1685–1702.
- Asgharzadeh, A., Valiollahi, R., Raqab, M. Z., 2013. Estimation of the stress–strength reliability for the generalized logistic distribution. *Statistical Methodology*, 15: 73–94
- Ashkar, F., Mahdi, S., 2006. Fitting the log-logistic distribution by generalized moments. *Journal of Hydrology*, 328 (3-4): 694-703.
- Dey, A. K., Kundu, D., 2009. Discriminating Between the Log-Normal and Log-Logistic Distributions. *Communications in Statistics - Theory and Methods*, 39 (2): 280-292.
- El Genidy M. M., 2012. Estimation Some Parameters of K-station Series Model. *Asian Journal of Applied Sciences*, 5: 307-313.
- El Genidy, M. M., 2019. Multiple nonlinear regression of the Markovian arrival process for estimating the daily global solar radiation. *Commun. Stat. Theory*, 5 (3): 93-98.
- El Genidy, M. M., Abd El-Rahman D. A., 2019. A New High Accurate Estimation Method for Evaluating the Daily Solar Energy by Nested Percentiles Algorithm. *Asian J. Sci. Res.*, 12 (4): 480-487.

El Genidy, M. M., Ali, A. K., 2016. Modeling the Amount of Pollutants Ozone Using Moments Method and Generalized Extreme Value Distribution. *Asian J. Sci. Res.*, 9 (4): 143-151.

Gupta, R. D, Kundu, D., 2010. Generalized Logistic Distributions. *Journal of Applied Statistical Science*, 18 (1): 51

Huang, C., Pather, V., Hammujuddy, J., Chinghamu, K., 2017. Extreme Risk in Resource Indices and The Generalized Logistic Distribution. *The Journal of Applied Business Research*, 33 (2): 47-60.

Kantam, R. R.L., Rao, G. S., Sriram, B., 2007. An economic reliability test plan: Log-logistic distribution. *Journal of Applied Statistics*, 33 (3): 291-296

Kaoga, D. K., Raidandi, D., Djongyang, N., Doka, S. Y., 2014. Comparison of five numerical methods for estimating Weibull parameters for wind energy applications in the district of Kousseri, Cameroon. *AJSC.*, 3 (1): 73-88.

Khan, J. K., Ahmed, F., Uddin, Z., Iqbal, S. T., Jilani, S. U., Siddiqui, A. A., Aijaz, A., 2015. Determination of Weibull Parameter by Four Numerical Methods and Prediction of Wind Speed in Jiwni (Balochistan). *Journal of Basic & Applied Sciences*. 11: 62-68.

Khan, M. S., King, R., 2013. Transmuted Modified Weibull Distribution: A Generalization of the Modified Weibull Probability Distribution. *European Journal of Pure and Applied Mathematics.*, 6: 66-88.

Maiti, S., Mukherjee, I., 2017. On estimation of the PDF and C.D.F of the Lindley distribution. *Commun. Stat. Simulat.*, 47 (4): 1370-1381.

Poudel, K. P., Cao, Q. V., 2013. Evaluation of Methods to Predict Weibull Parameters for Characterizing Diameter Distributions. *Forest Sci.*, 59 (2): 243–252.

Razali, N. M., Wah, Y. B., 2011. Power comparisons of Shapiro-Wilk, Kolmogorov-Smirnov, Lilliefors and Anderson-Darling tests. *Journal of Statistical Modeling and Analytics.*, 2: 21-33.

Rocha, P. A., Sousa, R. C., Anderade, C. F., Silva, M. E., 2012. Comparison of seven numerical methods for determining Weibull parameters for wind energy generation in the northeast region of Brazil. *Appl. Energ.*, 89 (1): 395-400.

Saat N. Z. M., Jemain A. A., 2009. Modelling the Duration of Hypopnea. *J. Applied Sci.*, 9: 2162-2167.

Shin, H., Jung, Y., Jeong, C., Heo, J. H., 2012. Assessment of modified Anderson–Darling test statistics for the generalized extreme value and generalized logistic distributions. *Stoch. Env. Res. Risk A.*, 26 (1): 105–114.

Soukissian, T. H., Tsalis, C., 2015. The effect of the generalized extreme value distribution parameter estimation methods in extreme wind speed prediction. *Nat. Hazards*, 78 (3): 1777–1809.

Valiollahi, R., Asgharzadeh, A., Raqab, M. Z., 2013. Estimation of $P(Y < X)$ for Weibull Distribution Under Progressive Type-II Censoring. *Commun. Stat. Theory*, 42 (42): 4476-4498.

Wasinrat, S., Bodhisuwan W., Zeephongsekul P., Thongtheeraparp A., 2013. A Mixture of Weibull Hazard Rate with a Power Variance Function Frailty. *J. Applied Sci.*,13: 103-110.



جامعة القاهرة
كلية الدراسات العليا للبحوث الإحصائية

المؤتمر السنوى الرابع والخمسون للإحصاء
وعلوم الحاسب وبحوث العمليات

إحصاء رياضي

11-9 ديسمبر 2019