

See discussions, stats, and author profiles for this publication at: <https://www.researchgate.net/publication/393631983>

Bayesian and Non-Bayesian Analysis of the Novel Unit Inverse Exponentiated Lomax Distribution Using Progressive Censoring Schemes with Optimal Scheme and Data Application

Article in Computational Journal of Mathematical and Statistical Sciences · July 2025

DOI: 10.21608/cjmss.2025.374277.1151

CITATIONS

0

READS

2

4 authors, including:



Ehab M. Almetwally

Delta University for Science and Technology

267 PUBLICATIONS 3,936 CITATIONS

SEE PROFILE



Research article

Bayesian and Non-Bayesian Analysis of the Novel Unit Inverse Exponentiated Lomax Distribution Using Progressive Censoring Schemes with Optimal Scheme and Data Application

Mohammed Elgarhy^{1,*}, Gaber Sallam Salem Abdalla², Amal S. Hassan³ and Ehab M. Almetwally⁴

¹ Department of Basic Sciences, Higher Institute for Administrative Sciences, Belbeis 44621, Egypt; dr.moelgarhy@gmail.com

² Department of Insurance and Risk Management, Faculty of Business, Imam Mohammad Ibn Saud Islamic University (IMSIU), Riyadh 11432, Saudi Arabia; jssabdullah@imamu.edu.sa

³ Faculty of Graduate Studies for Statistical Research, Cairo University, Giza 12613, Egypt; amal52_soliman@cu.edu.eg

⁴ Department of Mathematics and Statistics, Faculty of Science, Imam Mohammad Ibn Saud Islamic University (IMSIU), Riyadh, 11432, Saudi Arabia; emalmetwally@imamu.edu.sa

* **Correspondence:** dr.moelgarhy@gmail.com

Abstract: This article presents a novel unit distribution called the unit inverse exponentiated Lomax distribution. Some of its key characteristics are carefully examined. The distribution parameters are estimated using both traditional and Bayesian approaches, taking into account the progressive Type II censoring schemes. Using the symmetric loss function and the Markov chain Monte Carlo technique, the Bayesian methodology is investigated to determine the point and credible interval estimates of parameters. In order to select the best progressive censoring scheme, a number of optimization criteria are taken into consideration. According to the selected criteria measures, the simulations showed that the Bayesian estimates outperform the maximum likelihood estimates in terms of accuracy, indicating better parameter estimation precision. Additionally, most coverage probabilities were high, at about 95 %. The novel unit distribution flexibility is validated using a three-real data set from different domains.

Keywords: Inverse exponentiated Lomax distribution; Bayesian estimation; optimal design.

Mathematics Subject Classification: 62F15, 62K05, 62N02.

Received: 11 April 2025; Revised: 4 July 2025; Accepted: 7 July 2025; Online: 11 July 2025.



Copyright: © 2024 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license.

1. Introduction

Lifetime data modeling and statistical analysis are crucial in many applicable fields, including engineering, finance, medicinal research, and insurance. As a result, in these fields, several lifespan distributions have been introduced. In particular, modeling datasets restricted to the interval $(0, 1)$ has become more popular in the last several years. This method has proven to be very helpful in addressing the success and failure rates of products in a variety of industries. Its versatility in managing these kinds of probabilistic models has led to the emergence of several unit distributions that are confined inside the interval $(0, 1)$. Additionally, areas like banking, actuarial, and medical fields are realizing more and more how important these kinds of distributions are. The beta distribution (BD) is one of the most popular distributions that is employed on these kinds of datasets due to its adaptability. The BD has a disadvantage, though, in that it is not always suitable for real-world scenarios, such as hydrological data. Reference [1] presented the Kumaraswamy distribution (KumD), which is a distribution close to the BD. In addition to the KumD, plenty of other unit distributions have been developed to attempt to better suit the constantly increasing amount of data-sets on the unit interval that may result from various complicated events. Fortunately, statisticians have recently become more interested in putting out distributions that are determined by the unit interval that corresponds to any continuous distribution. Readers may examine, unit-Birnbaum-Saunders distribution [2], unit-Lindley distribution [3], unit-inverse Gaussian distribution [4], unit-Weibull distribution (UWD) [5], unit-Gompertz distribution (UGD) [6], unit-exponentiated half-logistic distribution (UEHLD) [7], unit Burr-III distribution [8], unit Burr-XII distribution [9], unit-improved second-degree Lindley distribution [10], unit Teissier distribution [11], unit exponentiated Lomax distribution [12], unit half-logistic geometric distribution [13], unit inverse exponentiated Weibull distribution [14], unit generalized half-normal distribution [15], unit-power Burr X distribution (UPBXD) [16], unit half normal distribution [17], unit power Lomax distribution [18], unit Maxwell-Boltzmann distribution [19], and power new power function distribution [20], among others.

One of the key components in many technical and medical disciplines is modeling heavy-tailed data. One of the most significant heavy-tailed alternatives to the gamma, Weibull, and exponential distributions is the Lomax distribution (LD). It has been used in a number of fields, including biological sciences [21], firm size [22], life testing and reliability [23], and analysis of wealth and income data [24, 25]. Despite having many uses, the LD is unable to handle data that have bathtub upside-down failure shapes. In order to improve the flexibility of the LD in representing different kinds of data, a number of extensions have recently been built utilizing several techniques. A number of the most significant LD extensions are Marshall-Olkin LD [26], exponentiated-Lomax distribution (ELD) [27], McDonald- LD [28], Weibull-LD [29], weighted power LD [30], gamma-LD [31], exponentiated Weibull-Lomax distribution [32], type II half logistic-LD [33], Nadarajah-Haghighi-LD [34], and modified Kies-LD [35], among others.

For the purpose of this study, the ELD is required. The probability density function (PDF) and cumulative distribution function of the ELD with scale parameter $a > 0$, and shape parameters $b > 0$, and $c > 0$, are respectively, given as

$$f(x) = abc(1 + ax)^{-b-1} \left[1 - (1 + ax)^{-b} \right]^{c-1}; \quad x > 0, \quad (1.1)$$

and

$$F(x) = \left[1 - (1 + ax)^{-b}\right]^c; \quad x > 0, \quad (1.2)$$

where a is scale parameter and b, c are two shape parameters. Some important information about the ELD are as follows: For $a = 1$, the PDF (1.1) reduces to the exponentiated Pareto distribution, for $c = 1$, the PDF (1.1) provides LD.

The development of more flexible models through the use of various techniques, including the inverse transformation technique, has therefore gained increasing attention. Along with increasing the distribution's flexibility, this strategy keeps the number of parameters constant. The recently created inverted ELD (IELD) by Reference [36] is the subject of this discussion. Let $Y = \frac{1}{X}$, where $X \sim \text{ELD}(a, b, c)$, the PDF and CDF of the IELD with scale parameter $a > 0$, and shape parameters $b > 0$, and $c > 0$, are respectively, given as

$$f(y) = \frac{abc}{y^2} \left(1 + \frac{a}{y}\right)^{-b-1} \left[1 - \left(1 + \frac{a}{y}\right)^{-b}\right]^{c-1}; \quad y > 0, \quad (1.3)$$

and

$$F(y) = 1 - \left[1 - \left(1 + \frac{a}{y}\right)^{-b}\right]^c; \quad y > 0. \quad (1.4)$$

The IELD displays an increasing, decreasing, reversed j-shaped and inverted bathtub-shaped hazard rate function (HRF), which is widespread in most real-world systems and highly valuable in survival analysis. The IELD can be regarded as a fitting model for positive data with an extended right tail, presuming a smooth increasing HRF and being applicable to a variety of domains.

This paper's goal is to create novel flexible three-parameter distribution for modeling data given in the interval $(0, 1)$. We adopted the transformation $Z = \frac{Y}{1+Y}$, where Y holds the IELD, and utilized the IELD to create a unit IELD (UIELD). We are motivated to introduce the UIELD for the following reasons:

1. The UIELD reveals an increasing, decreasing, j-shaped, and U-shaped HRF, hence, its functional flexibility is helpful for modeling skewed data seen in a variety of domains.
2. Several pivotal statistical attributes of the UIELD are ascertained, comprising moments, quantile function (QF), incomplete moments (IM), probability-weighted moments (PWMs), uncertainty measures, and stress-strength (SS) reliability.
3. The maximum likelihood estimates (MLEs) as well as the approximate confidence intervals (ACIs) are derived using progressively type-II (PTII) censored sampling. The Bayesian estimates (BEs) are then computed using the squared error loss function (SEL) and the Markov chain Monte Carlo (MCMC) technique. The highest posterior density (HPD) credible intervals are also computed. Three optimization criteria are employed to select the most suitable scheme.
4. It is important to note that evaluating the various point and interval estimates conceptually is more difficult when analyzing their efficacy. As a result, a simulation study is provided to achieve this objective.
5. Three separate real-world datasets were used for the empirical validation of our UIELD framework, yielding strong proof of its statistical performance and practical applicability. Our suggested framework's superior fitting ability was consistently shown across all assessed metrics in comparison to six well-known models (UWD, KumD, BD, UEHLD, UGD, and UPBXD).

Here is a summary of the paper's organizational framework. Section 2 presents the UIELD's mathematical description. Its mathematical properties are examined in Section 3. Section 4 looks at MLEs, BEs, ACIs, HPD credible intervals, the best censoring methods, in addition to the simulation research. In Section 5, we talk about applying it to real-world datasets. Section 6 serves as the article's conclusion.

2. Description of the New Model

In this section, the UIELD defined on the interval (0,1) with scale parameter a and shape parameters, b and c is introduced. For this purpose, assume that $Z = \frac{Y}{1+Y}$ where $Y \sim \text{IELD}(a, b, c)$, hence the CDF of the $Z \sim \text{UIELD}(a, b, c)$, is obtained as follows:

$$\begin{aligned} G(z; \zeta) &= P(Z \leq z) = P\left(\frac{Y}{1+Y} \leq z\right) = P\left(Y \leq \frac{z}{1-z}\right) \\ &= F_Y\left(\frac{z}{1-z}\right) = 1 - \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^c; \quad 0 < z < 1. \end{aligned} \quad (2.1)$$

For $z \leq 0$, we easily set $G(z; \zeta) = 0$, and for $z > 1$ we set $G(z; \zeta) = 1$. Note that $\zeta \equiv (a, b, c)$, is set of parameters that satisfies $a > 0, b > 0$, and $c > 0$.

$$g(z; \zeta) = \frac{abc}{z^2} \left(1 + \frac{a(1-z)}{z}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^{c-1}; \quad 0 < z < 1, \quad (2.2)$$

and $g(z; \zeta) = 0$, for $g(z; \zeta) \notin (0, 1)$. For $c = 1$, the PDF (1.3) reduces to the unit ILD (UILD) as a new unit model. The survival function (SF) and the HRF, for $z \in (0, 1)$, are as follows

$$\bar{G}(z; \zeta) = \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^c,$$

and

$$h(z; \zeta) = \frac{abc}{z^2} \left(1 + \frac{a(1-z)}{z}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^{c-1}.$$

The density and HRF plots of the UIELD are given in Figures 1 and 2 for specific values of parameters. It can be observed from Figure 1 that the PDF of the UIELD can be quite flexible, exhibiting a variety of asymmetric shapes: left-skewed, reversed J-shaped, U-shaped, or even a "decreasing-increasing-decreasing" pattern. In contrast, the HRF in Figure 2 is thankfully simpler adopting J-shaped, U-shaped, increasing, or decreasing forms.

3. Some Fundamental Characteristics

Here, we establish some basic mathematical characteristics of the UIELD, such as PWMs, moments and IMs, QF, some entropy measures and S-S reliability parameter.

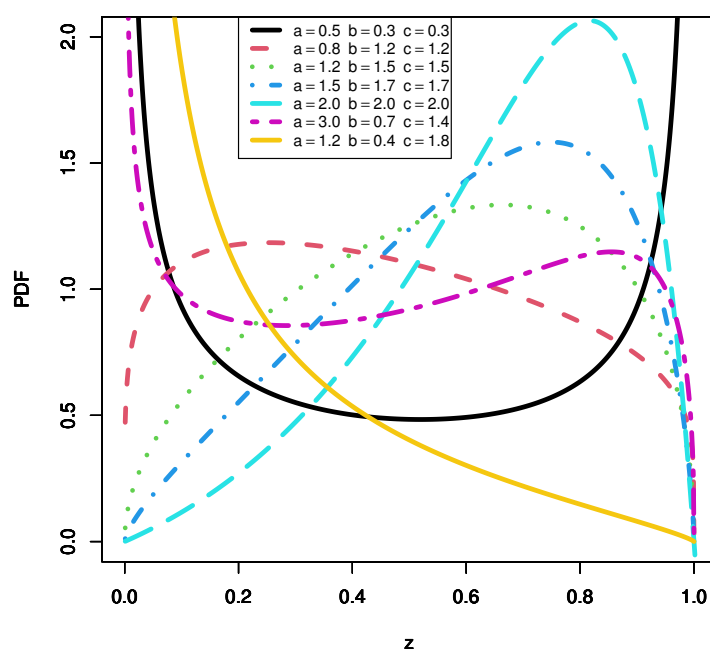


Figure 1. Plots of the PDF for the UIELD

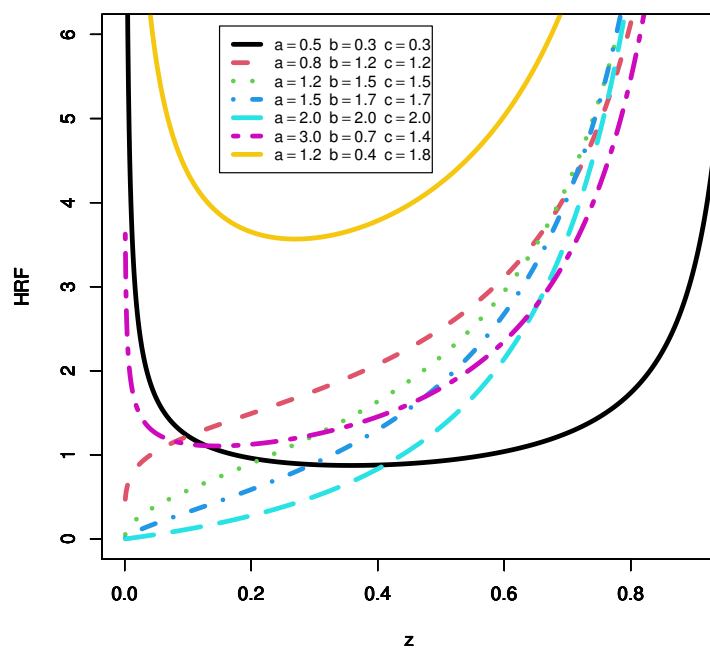


Figure 2. Plots of the HRF for the UIELD

3.1. Quantile Function

For $p \in (0, 1)$, the QF of Z is obtained by inverting CDF (2.2) as follows:

$$p = 1 - \left[1 - \left(1 + \frac{a(1 - Q(p))}{Q(p)} \right)^{-b} \right]^c,$$

that supplies;

$$Q(p) = [1 + A(p, c, b, a)]^{-1}, \quad A(p, c, b, a) = \frac{1}{a} \left[\left(1 - (1 - p)^{1/c} \right)^{-1/b} - 1 \right]. \quad (3.1)$$

In particular, the first quartile, say Q_1 , is produced by setting $p = 0.25$ in (3.1), the second quartile or median, say Q_2 is produced by setting $p = 0.5$ and the third quartile is produced by setting $p = 0.75$ in (3.1).

3.2. Moments Measures

The main characteristics of a distribution, such as skewness, kurtosis, dispersion, and central tendency, may be studied using moments. The moments of the UIELD can be obtained as an infinite power series as mentioned below. The m th moment of the UIELD is obtained by using PDF (2.2) as follows:

$$\mu'_m = abc \int_0^1 z^{m-2} \left(1 + \frac{a(1-z)}{z} \right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z)}{z} \right)^{-b} \right]^{c-1} dz. \quad (3.2)$$

Suppose that $y = \left(1 + \frac{a(1-z)}{z} \right)^{-b}$ then the m th moment of (3.2) will be as follows

$$\mu'_m = abc \int_0^1 \left(1 - \frac{1}{a}(1-y)^{-1/b} \right)^{-m} [1-y]^{c-1} dy. \quad (3.3)$$

Using the following expansion

$$(1-k)^{-d} = \sum_{i_1=0}^{\infty} \frac{\Gamma(d+i_1)k^{i_1}}{\Gamma(d)i_1!}, \quad |k| < 1,$$

and binomial expansion in Equation (3.3), we have

$$\mu'_m = \sum_{i_1=0}^{\infty} \sum_{i_2=0}^{i_1} \binom{i_2}{i_1} \frac{(-1)^{i_2} \Gamma(m+i_1)c}{\Gamma(m)(i_1!)a^{i_1-1}} B\left(1 - \frac{i_2}{b}, c\right), \quad i_2 < b,$$

where $B(., .)$ is the beta function (BF). Furthermore, the m th central moment of Z , is given by

$$\mu_m = E(Z - \mu'_1)^m = \sum_{l=0}^m (-1)^l \binom{m}{l} (\mu'_1)^l \mu'_{m-l}.$$

The 3D plots of mean, variance, coefficient of skewness and coefficient of kurtosis for the UIELD are provided in Figures 3 and 4 for different values of shape parameter c .

Furthermore, the m th moment of the UIELD is obtained by using PDF (2.2) as follows:

$$\vartheta_m(t) = abc \int_0^t z^{m-2} \left(1 + \frac{a(1-z)}{z}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^{c-1} dz.$$

Using binomial expansions and after some simplification the m th incomplete moment of the UIELD is as follows:

$$\vartheta_m(t) = \sum_{i_1=0}^{\infty} \sum_{i_2=0}^{i_1} \binom{i_2}{i_1} \frac{(-1)^{i_2} \Gamma(m+i_1)c}{\Gamma(m)(i_1!)a^{i_1-1}} B\left(1 - \frac{i_2}{b}, c, \left(1 + \frac{a(1-t)}{t}\right)^{-b}\right),$$

where $B(.,.,x)$ is the incomplete BF. Perhaps the prominent uses of the first incomplete moment are the Lorenz curve, represented by $L(t) = \vartheta_1(t)/\mu'_1$ and Bonferroni curves, represented by $B(t) = L(t)/F(t)$. These curves are very useful in the fields of economics, demography, insurance, engineering, and medicine.

The UIELD's m th inverse moment is obtained using PDF (2.2) in the manner shown below:

$$E(Z^{-m}) = abc \int_0^1 z^{-m-2} \left(1 + \frac{a(1-z)}{z}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^{c-1} dz.$$

The m th inverse moment of the UIELD is as follows after simplifying and applying binomial expansions

$$E(Z^{-m}) = \sum_{i_1=0}^{\infty} \sum_{i_2=0}^{i_1} \binom{i_2}{i_1} \frac{(-1)^{i_2} \Gamma(m+i_1)c}{\Gamma(m)(i_1!)a^{i_1-1}} B\left(1 + \frac{i_2}{b}, c\right).$$

After simplification and using binomial expansions twice times, the m th inverse moment of the UIELD.

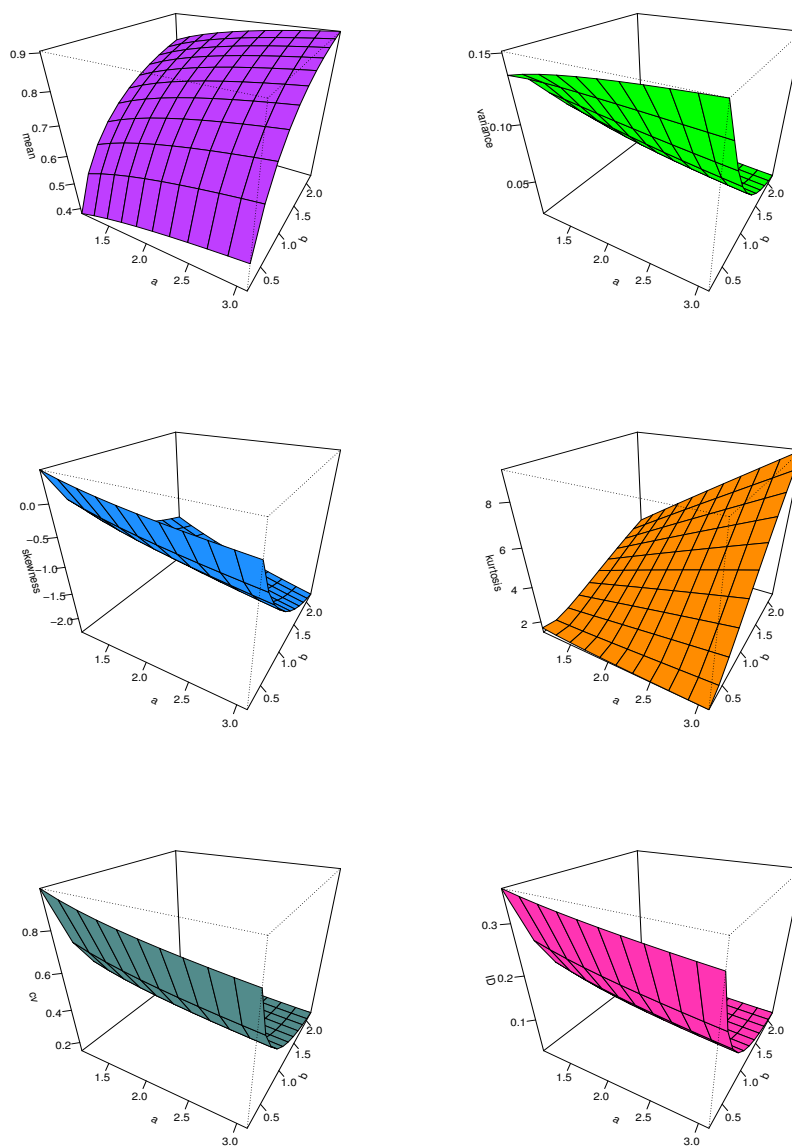


Figure 3. 3D Plots of some measures of moments associated with the UIELD at $c = 0.5$

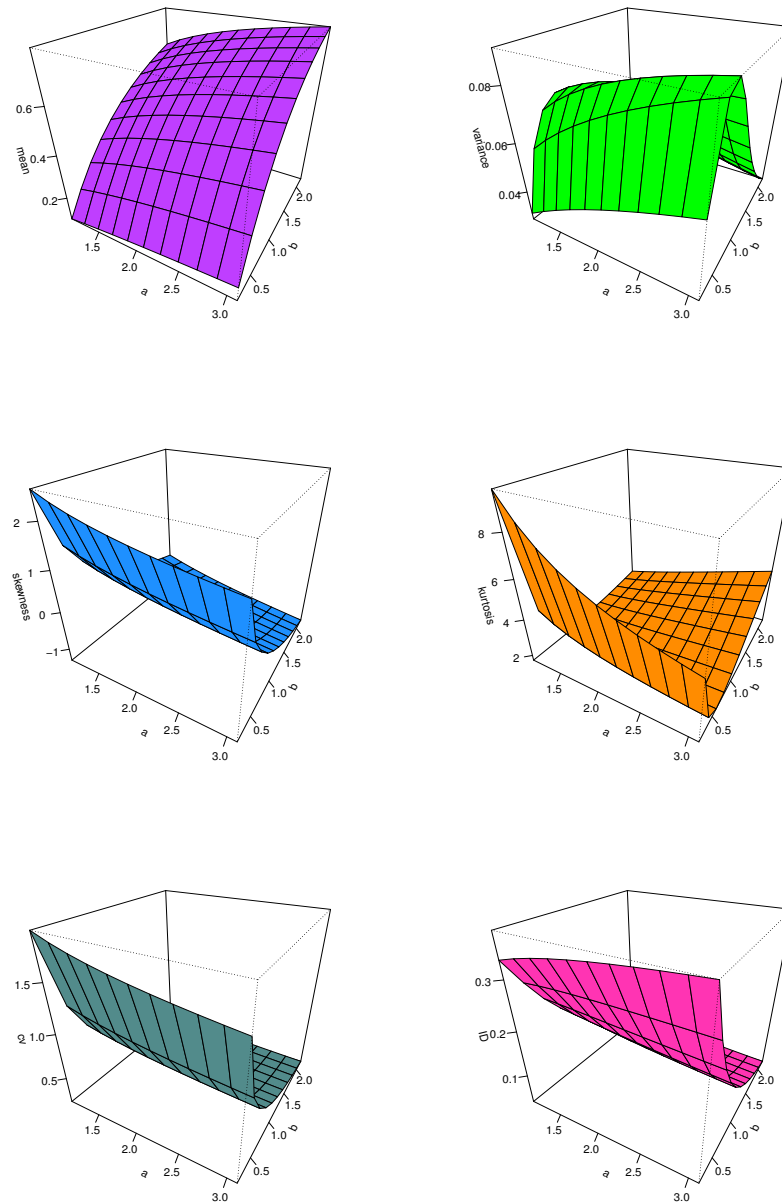


Figure 4. 3D Plots of some measures of moments associated with the UIELD at $c = 1.5$

3.3. Stress-Stress Reliability

Industrial component reliability metrics are widely used, particularly in the field of engineering. A system or product's reliability is determined by how likely it is to operate in the required time frame under standard (or defined) environmental conditions and fulfill its intended purpose. A component's life under random stress (Z_2) and random strength (Z_1) is described by an SS model. When the stress applied to a component exceeds its strength, the component fails instantly, and when $Z_2 > Z_1$, it performs adequately. The $\varsigma = P[Z_2 < Z_1]$ is typically referred to as "SS reliability" in statistical

literature. The SS reliability models are frequently used in mechanical engineering to examine the connection between material strength and applied stress in mechanical systems and components. These models are crucial for risk assessment and design validation because they evaluate the likelihood that a system will function as intended under operating loads without malfunctioning. In medical research, the SS reliability Z_1 represents the treatment effect of an experimental therapy, and Z_2 denotes the corresponding effect in the control group. This framework enables direct quantification of therapeutic superiority in randomized clinical trials.

Let Z_1 and Z_2 are two independent random variables, where $Z_1 \sim \text{UIELD}(z; a, b_1, c_1)$ and $Z_2 \sim \text{UIELD}(z; a, b_2, c_2)$, thus, the following is the determination of the UIELD's SS reliability:

$$S = 1 - \int_0^1 \frac{ab_1c_1}{z^2} \left(1 + \frac{a(1-z)}{z}\right)^{-b_1-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b_1}\right]^{c_1-1} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b_2}\right]^{c_2} dz. \quad (3.4)$$

Through employing the binomial expansions twice times in Equation (3.4), the SS reliability can possibly be expressed as follows.

$$S = 1 - \sum_{j_1, j_2=0}^{\infty} (-1)^{j_1+j_2} \binom{c_1-1}{j_1} \binom{c_2}{j_2} \frac{b_1c_1}{b_1(j_1+1) + b_2j_2}.$$

3.4. Entropy Measures

Entropy, which was first applied in physics, is one of the most widely used methods for evaluating the degree of uncertainty related to a random variable. In several fields, including statistics, chemistry, and biology, entropy measurement is crucial. Less information found in a sample is associated with higher entropy. We examine some information metrics in this context, including Rényi (Ré) entropy [37], Tsallis entropy [38], and Arimoto entropy [39]. For a random variable Z , the Ré entropy is defined by:

$$R(v) = \frac{1}{(1-v)} \log \left(\int_0^1 (f(z))^v dz \right), \quad v > 0, v \neq 1. \quad (3.5)$$

Inserting PDF (2.2) in (3.5), then we have

$$R(v) = \frac{1}{1-v} \log \left(\int_0^1 \frac{(abc)^v}{z^{2v}} \left(1 + \frac{a(1-z)}{z}\right)^{-v(b+1)} \left[1 - \left(1 + \frac{a(1-z)}{z}\right)^{-b}\right]^{v(c-1)} dz \right). \quad (3.6)$$

With the use of binomial expansions in integration (3.6) provides the following expression

$$\begin{aligned} I &= \sum_{k_1, k_2=0}^{\infty} (-1)^{k_1+k_2} \binom{v(c-1)}{k_1} \binom{bk_1 + v(b+1)}{k_2} \int_0^1 z^{-2v-k_2} (1-z)^{k_2} dz \\ &= \sum_{k_1, k_2=0}^{\infty} (-1)^{k_1+k_2} \binom{v(c-1)}{k_1} \binom{bk_1 + v(b+1)}{k_2} (a)^{v+k_2} (bc)^v B(1-2v-k_2, k_2+1), \quad 2v+k_2 < 1. \end{aligned} \quad (3.7)$$

Hence, the Ré entropy of the UIELD is obtained by setting (3.7) in (3.6) as seen below:

$$R(v) = \frac{1}{1-v} \log \left(\sum_{k_1, k_2=0}^{\infty} (-1)^{k_1+k_2} \binom{v(c-1)}{k_1} \binom{bk_1 + v(b+1)}{k_2} (a)^{v+k_2} (bc)^v B(1-2v-k_2, k_2+1) \right).$$

The Tsallis entropy of the Z is given by:

$$T(v) = \frac{1}{(v-1)} \left(1 - \int_0^{\infty} (f(z))^v dz \right), \quad v > 0, \quad v \neq 1. \quad (3.8)$$

Inserting PDF (2.2) and integration (3.7) in (3.8), then the Tsallis entropy of the UIELD is given by:

$$T(v) = \frac{1}{(v-1)} \left(1 - \sum_{k_1, k_2=0}^{\infty} (-1)^{k_1+k_2} \binom{v(c-1)}{k_1} \binom{bk_1 + v(c-1)}{k_2} (a)^{v+k_2} (bc)^v B(1-2v-k_2, k_2+1) \right).$$

The Arimoto entropy of the Z is given by:

$$A(v) = \frac{v}{1-v} \left(\left[\int_0^{\infty} (f(z))^v dz \right]^{1/v} - 1 \right), \quad v > 0, \quad v \neq 1. \quad (3.9)$$

Inserting PDF (2.2) and integration (3.7) in (3.9), then the Arimoto entropy of the UIELD is given by:

$$A(v) = \frac{v}{1-v} \left(\left[\sum_{k_1, k_2=0}^{\infty} (-1)^{k_1+k_2} \binom{v(c-1)}{k_1} \binom{bk_1 + v(c-1)}{k_2} (a)^{v+k_2} (bc)^v B(1-2v-k_2, k_2+1) \right]^{1/v} - 1 \right).$$

4. Progressively Type-II censored scheme

This section discusses different ways data can be censored in experiments, where censoring means not having complete information about all subjects. Two common types are:

Type-I censoring: This stops the experiment after a fixed time, regardless of how many failures occurred.

Type-II censoring: This lets the experiment run until a pre-determined number of failures are observed.

Recently, a more flexible approach called **progressive censoring scheme** has become popular. This method makes better use of resources compared to traditional methods. PTII censoring is an extension of Type-II censoring. We start with n items in a test and aim to observe m failures. When the first failure happens, we randomly remove R_1 of the remaining $n-1$ items. After the second failure, we remove R_2 of the remaining items, and so on. The experiment ends after the ' m th failure, with all remaining items $R_m = n - m - R_1 - R_2 - \dots - R_{m-1}$ being removed. The number of items removed at each stage (R_1, R_2 , etc.) is decided before the experiment starts. This approach is written as $z_{1:m:n}, z_{2:m:n}, \dots, z_{m:m:n}$, where ' z ' represents the observed failure times. Finally, the text clarifies that the total number of items (n) is equal to the number of failures (m) plus the number of items removed at each stage (R_1 to R_m).

This section talks about the statistical properties of PTII censored data. Assume, we have an experiment with n items undergoing a life test. We assume the times it takes for these items to fail follow

a continuous probability distribution. This distribution is characterized by its CDF denoted by $G(z)$ and its PDF denoted by $g(z)$. The text refers to a result by Balakrishnan and Aggarwala [40] which describes the formula for the joint density function of the observed failure times when using PTII censoring (m failures observed with progressive removal of surviving items). This joint density function essentially describes the probability of observing a specific set of failure times under these conditions. The likelihood function

$$L(\zeta) = Q \prod_{i=1}^m f(z_{i:m:n}; \zeta) [1 - F(z_{i:m:n}; \zeta)]^{R_i}, \quad (4.1)$$

where ζ is a vector of parameters, and Q is a fixed value and does not depend on ζ , where $Q = n(n-1-R_1)(n-2-R_1-R_2)\dots(n-m+1-\sum_{j=1}^{m-1} R_j)$.

4.1. Estimation Methods

The ML and Bayesian estimation methods have been discussed for parameters of UIELD based in PTII censoring scheme in this section.

4.1.1. Maximum Likelihood Estimators

In this sub-subsection, the MLEs of the UIELD parameters a, b , and c under PTII censored data are provided. Assume that $z_{1:m:n}, z_{2:m:n}, \dots, z_{m:m:n}$ is a PTII censored sample from a life testing of size m obtained from the UIELD. The likelihood function (LF) is

$$L(\zeta) = Q(abc)^m \prod_{i=1}^m \frac{1}{z_{i:m:n}^2} \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]^{cR_i+c-1}. \quad (4.2)$$

The natural logarithm of the LF for the UIELD distribution, based on a PTII censored sample, is

$$\begin{aligned} \ell(\zeta) = \ln(Q) + m(\ln a + \ln b + \ln c) - 2 \sum_{i=1}^m \ln z_{i:m:n} - (b+1) \sum_{i=1}^m \ln \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right) + \\ \sum_{i=1}^m (cR_i + c - 1) \ln \left[1 - \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]. \end{aligned} \quad (4.3)$$

To find the MLE for the parameters a, b and c in the model, we need to maximize this log-LF in Equation (4.3). Also, we need to solve three nonlinear equations. These equations are derived from the maximized log-LF as follows:

$$\frac{\partial \ln(\zeta)}{\partial a} = \frac{m}{a} - \sum_{i=1}^m \frac{(b+1)z_{i:m:n}(1-z_{i:m:n})}{z_{i:m:n} + a(1-z_{i:m:n})} + \sum_{i=1}^m \frac{(cR_i + c - 1)b}{\left[1 - \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]} \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b-1} \frac{(1-z_{i:m:n})}{z_{i:m:n}}, \quad (4.4)$$

$$\frac{\partial \ell(\zeta)}{\partial b} = \frac{m}{b} - \sum_{i=1}^m \ln \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right) + b \sum_{i=1}^m (cR_i + c - 1) \frac{\left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b} \ln \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)}{1 - \left(1 + \frac{a(1-z_{i:m:n})}{z_{i:m:n}}\right)^{-b}}, \quad (4.5)$$

and

$$\frac{\partial \ell(\zeta)}{\partial c} = \frac{m}{c} + \sum_{i=1}^m (R_i + 1) \ln \left[1 - \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}} \right)^{-b} \right]. \quad (4.6)$$

Equation (4.3) can be explicitly maximized in R using the 'maxlike' function, which implements either the Nelder-Mead (NM) or Newton-Raphson (NR) maximization procedures for MLE calculations. This process addresses the non-linear log-likelihood equations obtained by differentiating Equation (4.3) with respect to the distribution parameters $\zeta = (a, b, c)$ and setting the result to zero.

4.2. Bayesian Estimation

As a powerful and useful alternative to traditional techniques, the Bayesian approach has gained significant attention in statistical analysis over the past few decades. Bayesian methods for parameter estimation have been successfully applied across a wide range of fields, including physics, the food chain, epidemiology, environmental science, COVID-19, and econometrics. In this subsection, Bayes estimates of the unknown parameters are derived using a SEL, and the corresponding HPD credible confidence interval is obtained. It is assumed that the various parameters of the UIELD are independent and have prior distributions that are both informative and non-informative.

We now examine the BEs for the ζ function under the SEL function using PTII sampling. The BEs of the parameters in the ζ BEs vector are constructed based on the posterior distributions given the data. The proposed method is briefly described below. The SEL function for the assumed prior distribution is then minimized using the posterior mean as follows:

$$\tilde{\zeta} = E(\zeta|Z),$$

where $\tilde{\zeta}^{(j)}$ represents the BEs vector for ζ at iteration (j). The choice of the loss function in Bayesian estimation depends on the specific problem and the data assumptions. The SEL function is often preferred for several reasons:

Firstly, SEL provides mathematical convenience by yielding simpler mathematical expressions and computational algorithms. Its optimization properties make it easier to find MLEs, either analytically or numerically.

Secondly, when model errors are assumed to follow a normal distribution, SEL aligns with the goal of maximizing the likelihood function. This is because the negative log-likelihood of the normal distribution is proportional to the squared error loss.

Lastly, SEL is robust to outliers in the data, especially when errors are normally distributed. It imposes larger penalties on significant errors compared to smaller ones, which can be beneficial depending on the problem's context. Additionally, using the R software's coda and HDInterval packages, as developed by Plummer and Meredith [42], we can determine the MCMC results and the HPD credible interval for ζ , respectively.

Assume $z_{1:m:n}, z_{2:m:n}, \dots, z_{m:m:n}$ represent associated observations, and $z_{1:m:n}, z_{2:m:n}, \dots, z_{m:m:n}$ denote a random sample of size m drawn from the UIELD based on PTII censored sample with unknown parameters. In the Bayesian analysis, we assume independent gamma priors for the model parameters due to their flexibility and suitability for positive-valued quantities. The gamma distribution is also a conjugate prior for many likelihoods involving scale or rate parameters, which simplifies the derivation

of the posterior distributions. In this work, we adopt relatively non-informative gamma priors with hyperparameters chosen to reflect vague prior knowledge, following the approach commonly used in the literature (see, e.g., [?], [44], [45] and [46]). We assume that these are separate random variables (RVs) following the gamma distribution, with the PDF provided by

$$\Pi(\zeta) \propto a^{\phi_1-1} b^{\phi_2-1} c^{\phi_3-1} e^{-(\omega_1 a + \omega_2 b + \omega_3 c)}. \quad (4.7)$$

To determine the appropriate hyperparameters for the independent joint prior, use the estimates and the variance-covariance matrix from the MLE technique. By equating the mean and variance of the gamma priors, the estimated hyperparameters can be expressed as follows:

$$\begin{aligned} \phi_1 &= \frac{\left[\frac{1}{I} \sum_{j=1}^I \hat{a}^j\right]^2}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{a}^j - \frac{1}{I} \sum_{i=1}^I \hat{a}^j\right]^2}; \quad \phi_2 = \frac{\left[\frac{1}{I} \sum_{j=1}^I \hat{b}^j\right]^2}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{b}^j - \frac{1}{I} \sum_{i=1}^I \hat{b}^j\right]^2}, \quad \phi_3 = \frac{\left[\frac{1}{I} \sum_{j=1}^I \hat{c}^j\right]^2}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{c}^j - \frac{1}{I} \sum_{i=1}^I \hat{c}^j\right]^2}, \\ \omega_1 &= \frac{\frac{1}{I} \sum_{i=1}^I \hat{a}^j}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{a}^j - \frac{1}{I} \sum_{i=1}^I \hat{a}^j\right]^2}; \quad \omega_2 = \frac{\frac{1}{I} \sum_{i=1}^I \hat{b}^j}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{b}^j - \frac{1}{I} \sum_{i=1}^I \hat{b}^j\right]^2}, \quad \omega_3 = \frac{\frac{1}{I} \sum_{i=1}^I \hat{c}^j}{\frac{1}{L-1} \sum_{i=1}^I \left[\hat{c}^j - \frac{1}{I} \sum_{i=1}^I \hat{c}^j\right]^2}, \end{aligned}$$

where I denotes the number of MLE iterations. For more information about these techniques (selecting hyper-parameters), refer to [47].

The parameters' posterior distribution can be stated as follows:

$$\pi^*(\zeta \mid data) = \frac{\Pi(\zeta) L(\zeta \mid data)}{\int_0^\infty \int_0^\infty \int_0^\infty \Pi(\zeta) L(\zeta \mid data) da db dc}. \quad (4.8)$$

Equation (4.9) demonstrates how to formulate the joint posterior proportional to the prior as an equation.

$$\begin{aligned} \pi^*(\zeta \mid data) &\propto a^{m+\phi_1+1} b^{m+\phi_2+1} c^{m+\phi_3+1} e^{-(\omega_1 a + \omega_2 b + \omega_3 c)} \prod_{i=1}^m \frac{1}{z_{i:m:n}^2} \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b-1} \\ &\quad \prod_{i=1}^m \left[1 - \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]^{cR_i + c - 1}. \end{aligned} \quad (4.9)$$

The full conditional distributions for a , b and c can be expressed, up to a proportional constant, as:

$$\begin{aligned} \pi^*(a \mid b, c, data) &\propto a^{m+\phi_1+1} e^{-\omega_1 a} \prod_{i=1}^m \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]^{cR_i + c - 1}, \\ \pi^*(b \mid a, c, data) &\propto b^{m+\phi_2+1} e^{-\omega_2 b} \prod_{i=1}^m \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b-1} \left[1 - \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]^{cR_i + c - 1}, \\ \pi^*(c \mid a, b, data) &\propto c^{m+\phi_3+1} e^{-\omega_3 c} \prod_{i=1}^m \left[1 - \left(1 + \frac{a(1 - z_{i:m:n})}{z_{i:m:n}}\right)^{-b}\right]^{cR_i + c - 1}. \end{aligned} \quad (4.10)$$

We can use both the Metropolis-Hastings algorithm and the Gibbs sampling technique to generate samples from the posterior distributions, since none of these posterior PDFs correspond to a single common distribution. For further insights into this approach, see references [48, 49, 50].

4.3. Method of Optimization

The ideal censoring scheme for gathering data has been a topic of much recent research (Burkschat [51], and Pradhan and Kundu [52]). In progressive censoring, where items are removed at specific points during the experiment, many different combinations $(R_1, R_2, \dots, R_m)$ of removal times are possible given the predetermined number of total items (n) and observed failures (m) .

Before choosing a specific sampling plan, we need a way to determine which progressive censoring approach provides the most informative data about the unknown parameters we're trying to estimate. There are two main challenges:

- **Generating unknown parameter data:** How can we estimate the unknown parameters using a particular progressive censoring scheme?
- **Comparing information content:** How can we compare the value of two different progressive censoring schemes?

To address these challenges for UIELD based on PTII censored sample with different schemes, this paper establishes a set of optimality criteria which listed in Table 1. Table 1 also provides various popular information measures that can be used to identify the best progressive censoring strategy for a given experiment.

Table 1. optimization criteria and methods

Criterion	Method	Target
Optim ₁	Min-trace $[I_{3 \times 3}]^{-1}$	Minimum
Optim ₂	Min-det $[I_{3 \times 3}]^{-1}$	Minimum
Optim ₃	Max-trace $[I_{3 \times 3}]$	Maximum

The text defines three optimality criteria ($Optim_1$, $Optim_2$, and $Optim_3$) to identify the most informative progressive censoring scheme for estimating multiple unknown parameters. For information about optimum censoring plans, the reader can refer to [53, 54, 55, 56, 57, 58].

- $Optim_3$: This criterion maximizes the observed Fisher information, represented by a 3x3 matrix $([I_{3 \times 3}])$. Fisher information essentially measures how much information an experiment provides about the parameters being estimated. Here, a larger value indicates more informative data.
- $Optim_1$ and $Optim_2$: These criteria both aim to minimize the amount of uncertainty in the estimates. They achieve this by minimizing the determinant and trace (sum of diagonal elements) of the inverse of the Fisher information matrix $([I_{3 \times 3}]^{-1})$. A smaller value suggests less uncertainty.

4.3.1. Fisher information matrix

A core concept in statistics, the Fisher information matrix quantifies the amount of information data holds about an unknown parameter. It is utilized to understand the behavior of maximum-likelihood estimates as they approach infinity and to determine the variance of an estimator. The inverse of the

Fisher information matrix serves as an estimator for the asymptotic covariance matrix. To compute the Fisher information matrix, one uses the expected values of the negative second-partial and mixed-partial derivatives of the log-likelihood function with respect to

$$I_{3 \times 3} = \begin{bmatrix} \kappa_{11} & \kappa_{12} & \kappa_{13} \\ \kappa_{21} & \kappa_{22} & \kappa_{23} \\ \kappa_{31} & \kappa_{32} & \kappa_{33} \end{bmatrix}_{\zeta=\hat{\zeta}}$$

where $\kappa_{11} = \frac{\partial^2 \ell(x; \zeta)}{\partial a^2}$, $\kappa_{21} = \kappa_{12} = \frac{\partial^2 \ell(x; \zeta)}{\partial a \partial b}$, $\kappa_{31} = \kappa_{13} = \frac{\partial^2 \ell(x; \zeta)}{\partial a \partial c}$, $\kappa_{22} = \frac{\partial^2 \ell(x; \zeta)}{\partial b^2}$, $\kappa_{32} = \kappa_{23} = \frac{\partial^2 \ell(x; \zeta)}{\partial b \partial c}$, and $\kappa_{33} = \frac{\partial^2 \ell(x; \zeta)}{\partial c^2}$.

4.4. Simulation

This section explains how researchers employed computer simulations, specifically Monte Carlo methods, to identify unknown parameters of a distribution known as UIELD, using the PTII censoring scheme. We compared two estimation techniques: the widely used ML and Bayesian approaches. To evaluate these methods, they analyzed bias, the accuracy of the estimates through average mean squared errors (MSEs), and the reliability of the estimates using the length of confidence intervals (LCI), specifically, LCI of ACI, denoted as LACI, and LCI of Bayesian credible interval, denoted as LCCI. Additionally, they assessed the frequency with which the CIs included the true value, known as coverage probabilities (CP), using CI techniques.

The simulations were conducted for various combinations of predetermined settings (n, m) and different schemes. Ultimately, the researchers used the simulation results to estimate the UIELD parameters from data samples generated using the PTII censoring scheme. The entire estimation process adhered to a specific sequence of steps carried out through computer simulations, as detailed below:

- Initial sample values: We start by setting initial values for the variables n , and m . These variables likely represent key parameters or settings for the simulation as: $n=100$, and 200 , & for $m=50$, 70 , 150 , and 170 .
- Additional initializations: We also establish initial values for the UIELD parameters in the following scenarios:
Case 1: $a = 1.6; b = 1.2; c = 1.3$ & Case 2: $a = 0.6; b = 0.8; c = 1.3$ & Case 3: $a = 0.6; b = 0.8; c = 0.5$.
- Involves generating a specific sample: We generate a sample known as a PTII sample, which likely refers to a specific type of data point utilized in the simulation. This process operates under the assumption that the number of units removed after each failure follows a fixed pattern (fixed scheme) as follows:
Scheme 1: $R_1 = R_m = \frac{n-m}{2}$, and $R_i = 0; i = 2, \dots, m-1$.
Scheme 2: $R_m = n-m$, and $R_i = 0; i = 1, 2, \dots, m-1$.
Scheme 3: $R_1 = n-m$, and $R_i = 0; i = 2, \dots, m$.

The following conclusions can be drawn from Tables 2, 3, 4, and 5:

- The estimates for the unidentified UIELD parameters a , b , and c are outstanding based on the least MSE, bias, LACI, LCCI, and CP values.

- As n (or m) increases, all estimates perform as expected. When $n - m$ decreases, all estimates maintain consistent performance.
- The MCMC estimates, which use SEL function, outperform the MLE in terms of the smallest MSE, Bias, and LCI values due to the gamma prior for the parameters a , b , and c .
- Scheme 2 has better results comparing to schemes 1 and 3, where this has smaller values of $Optim_1$, $Optim_2$, and $Optim_3$ for most cases.
- For intervals estimation, the LCCI gets small values compared to LACI. Most of the coverage probabilities were high, ranging around 95%.

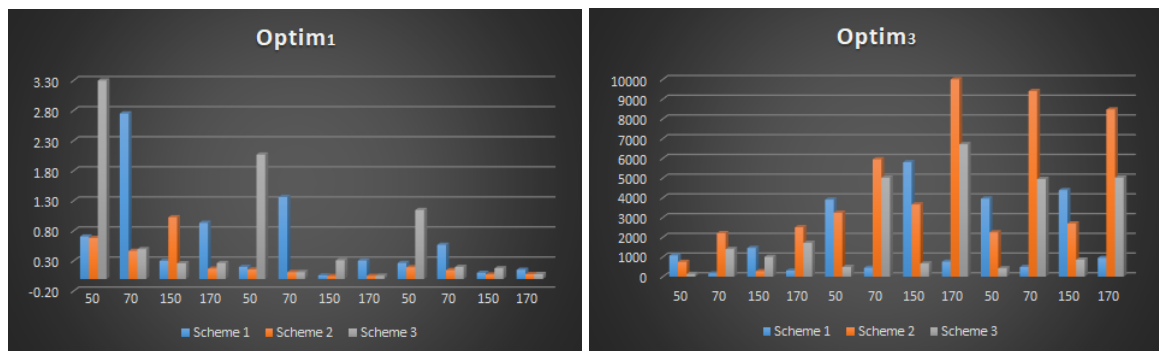


Figure 5. Par chart of optimal criteria

Figure 5 discussed bar chart of optimal criteria for $Optim_1$, and $Optim_3$. Scheme 2 has better results compared to schemes 1 and 3, where this has smaller values of $Optim_1$, $Optim_2$, and $Optim_3$ for most cases.

Table 2. ML and Bayesian estimation for parameters: $a = 1.6; b = 1.2; c = 1.3$

n	scheme	m		ML				Bayesian			
				Bias	MSE	LACI	CP	Bias	MSE	LCCI	CP
100	1	50	a	-1.2427	1.6052	0.9674	94.58%	0.5348	0.8274	0.9247	94.70%
			b	0.3916	0.4024	1.9574	95.99%	0.3608	0.3943	1.0015	94.90%
			c	-1.0376	1.0807	0.2478	95.89%	-0.1258	0.0202	0.1195	94.90%
		70	a	-1.2183	1.5404	0.9294	96.10%	0.2386	0.1678	0.8051	94.90%
			b	0.3559	0.3568	1.8156	96.40%	0.1105	0.0371	0.6062	95.90%
			c	-0.8850	0.7908	0.2342	95.97%	-0.1157	0.0201	0.0865	95.20%
	2	50	a	-0.0834	0.9818	3.8723	96.60%	0.3244	0.7060	2.7350	94.90%
			b	0.4598	0.7259	2.8133	93.60%	0.1825	0.1183	0.9706	94.70%
			c	0.0857	0.2058	1.7470	94.90%	0.2095	0.1800	1.2258	94.80%
		70	a	0.0357	0.9207	3.3071	96.85%	0.1326	0.5940	2.3741	95.15%
			b	0.2697	0.3230	1.9618	94.00%	0.0932	0.0458	0.6907	95.49%
			c	0.0691	0.1592	1.5410	95.10%	0.1599	0.1074	0.9544	95.35%
	3	50	a	-0.9641	1.0249	1.2121	97.70%	0.7087	0.9777	1.0782	95.70%
			b	-0.2619	0.2262	1.5571	95.00%	0.2416	0.2081	1.1691	94.60%
			c	-1.1355	1.2905	0.9128	93.40%	0.2037	1.0426	0.9027	94.90%
		70	a	-0.6297	1.0071	0.6705	98.46%	0.3766	0.3940	0.5853	95.15%
			b	0.2388	0.2136	1.3804	96.50%	0.2800	0.1165	0.7814	95.25%
			c	-1.0315	1.0670	0.1162	95.70%	-0.1267	0.0200	0.0834	95.49%
200	1	150	a	-1.1268	1.0626	0.5208	94.59%	0.1259	0.0277	0.3683	94.59%
			b	0.3586	0.3952	1.5289	95.29%	0.0671	0.0140	0.3608	94.70%
			c	-0.8720	0.7643	0.2469	95.19%	-0.1160	0.0201	0.0616	94.90%
		170	a	-1.1320	0.9323	0.8035	94.70%	0.0879	0.0210	0.3928	95.20%
			b	0.3150	0.3042	1.6006	95.40%	-0.0603	0.0071	0.3268	95.30%
			c	-0.7216	0.5280	0.2335	96.90%	-0.1017	0.0184	0.0808	95.20%
	2	150	a	0.0538	0.5944	3.0179	94.38%	0.2239	0.2624	1.5106	94.78%
			b	0.1101	0.0923	1.1111	94.78%	0.0396	0.0168	0.4342	94.58%
			c	0.0442	0.0743	1.0558	95.98%	0.0950	0.0441	0.6218	95.80%
		170	a	0.0296	0.5753	2.9520	94.80%	0.1300	0.2046	1.3694	95.80%
			b	0.0763	0.0694	0.9894	95.20%	0.0272	0.0120	0.3906	95.38%
			c	0.0434	0.0744	1.0563	96.20%	0.0757	0.0336	0.5838	95.98%
	3	150	a	-0.9308	0.8716	0.3028	93.60%	0.2254	0.0805	0.2439	95.10%
			b	0.2557	0.1442	1.1013	95.00%	0.2064	0.0537	0.4069	94.80%
			c	-1.0170	1.0361	0.1166	94.60%	-0.1305	0.0172	0.0527	94.80%
		170	a	-0.6028	0.8096	0.2935	95.20%	0.1313	0.0300	0.2359	95.80%
			b	0.1565	0.1185	1.0351	96.00%	0.0748	0.0141	0.3590	94.99%
			c	-0.8891	0.7941	0.1023	96.20%	-0.1058	0.0151	0.0535	95.80%

Table 3. ML and Bayesian estimation for parameters: $a = 0.6; b = 0.8; c = 1.3$

n	scheme	m		ML				Bayesian			
				Bias	MSE	LACI	CP	Bias	MSE	LCCI	CP
100	1	50	a	-0.5191	0.2729	0.2316	95.19%	0.1069	0.0367	0.2137	93.60%
			b	0.2264	0.1291	1.0951	96.35%	0.2840	0.1033	0.5770	94.80%
			c	-1.0687	1.1452	0.2176	95.74%	-0.1037	0.0118	0.0912	94.80%
		70	a	-0.4897	0.2447	0.2174	95.42%	0.0690	0.0135	0.1700	94.80%
			b	0.2158	0.1144	1.0925	96.80%	0.0889	0.0173	0.3684	95.38%
			c	-0.9224	0.8576	0.2032	95.80%	-0.0914	0.0218	0.1298	95.58%
	2	50	a	0.0767	0.3959	2.4507	95.25%	0.1409	0.1358	1.1376	94.80%
			b	0.2276	0.1748	1.3762	95.04%	0.0967	0.0279	0.4889	94.11%
			c	0.1754	0.4437	2.5218	95.45%	0.2216	0.1793	1.2727	94.80%
		70	a	0.0612	0.3147	2.1527	96.60%	0.1247	0.1196	1.0198	95.28%
			b	0.1197	0.0688	0.9157	95.94%	0.0502	0.0125	0.3505	94.38%
			c	0.1504	0.3183	2.1337	95.92%	0.1910	0.1567	1.0661	95.80%
	3	50	a	-0.5302	0.2847	0.2342	95.56%	0.1544	0.0730	0.2149	95.40%
			b	0.0413	0.0471	0.8357	94.92%	0.0406	0.0353	0.7831	94.80%
			c	-1.1656	1.3595	0.1244	95.95%	-0.0674	0.0110	0.0528	94.80%
		70	a	-0.5237	0.2763	0.1780	96.60%	0.0797	0.0166	0.1206	95.80%
			b	0.0320	0.0293	0.8092	95.60%	0.0229	0.0268	0.4548	95.80%
			c	-1.0616	1.1300	0.0922	96.80%	-0.0511	0.0092	0.0496	95.38%
200	1	150	a	-0.5015	0.2533	0.1687	95.17%	0.0369	0.0031	0.1128	89.00%
			b	0.2032	0.0917	0.7650	95.53%	0.0536	0.0063	0.2287	94.90%
			c	-0.9105	0.8328	0.2038	94.89%	-0.0945	0.0100	0.0818	95.00%
		170	a	-0.4561	0.2118	0.1520	95.69%	0.0277	0.0019	0.1089	94.90%
			b	0.1976	0.0732	0.7250	95.69%	0.0044	0.0025	0.1904	95.19%
			c	-0.7731	0.6057	0.1930	95.49%	-0.0911	0.0091	0.0811	95.80%
	2	150	a	0.0293	0.0981	1.2231	95.56%	0.1053	0.0540	0.6697	94.90%
			b	0.0574	0.0217	0.5320	95.70%	0.0219	0.0043	0.2456	94.90%
			c	0.0442	0.0862	1.1387	95.49%	0.1045	0.0494	0.7150	94.90%
		170	a	0.0232	0.0910	1.1242	95.67%	0.0910	0.0454	0.6036	95.90%
			b	0.0495	0.0197	0.5153	95.97%	0.0165	0.0034	0.2243	95.90%
			c	0.0394	0.0836	1.1236	96.19%	0.0882	0.0405	0.6349	95.90%
	3	150	a	-0.5036	0.2767	0.0872	95.43%	0.0381	0.0261	0.0810	95.00%
			b	0.0319	0.0401	0.6517	95.64%	0.0402	0.0343	0.3072	94.60%
			c	-1.0498	1.1035	0.1145	94.57%	-0.0511	0.0091	0.0401	95.40%
		170	a	-0.4507	0.2582	0.0715	95.59%	0.0375	0.0028	0.0591	96.70%
			b	0.0224	0.0184	0.6072	96.39%	0.0206	0.0070	0.2295	94.90%
			c	-0.9314	0.8714	0.0825	96.39%	-0.0401	0.0015	0.0392	96.70%

Table 4. ML and Bayesian estimation for parameters: $a = 0.6; b = 0.8; c = 0.5$

n	scheme	m		ML				Bayesian			
				Bias	MSE	LACI	CP	Bias	MSE	LCCI	CP
100	1	50	a	-0.4674	0.2248	0.3134	95.03%	0.1013	0.0401	0.2464	94.90%
			b	0.1668	0.1309	1.2593	94.80%	0.2032	0.0831	0.7134	94.90%
			c	-0.3708	0.1381	0.0981	95.84%	-0.0557	0.0036	0.0778	91.50%
		70	a	-0.4353	0.1973	0.2935	95.33%	0.0460	0.0108	0.2715	95.69%
			b	0.1594	0.1191	1.1192	95.53%	0.0630	0.0217	0.4756	95.79%
			c	-0.3042	0.0943	0.0892	95.82%	-0.0520	0.0034	0.0683	96.70%
	2	50	a	0.0977	0.5219	2.8075	95.93%	0.1222	0.1089	0.9911	94.79%
			b	0.3298	0.3591	1.9625	94.01%	0.1429	0.0701	0.7677	94.50%
			c	0.0457	0.0236	0.5757	95.71%	0.0460	0.0090	0.2874	94.60%
		70	a	0.0804	0.3078	2.1529	95.99%	0.1069	0.0849	0.8628	94.93%
			b	0.2171	0.1992	1.5296	94.61%	0.0733	0.0326	0.6006	94.95%
			c	0.0334	0.0137	0.4405	95.85%	0.0319	0.0049	0.2270	94.91%
	3	50	a	-0.4633	0.2195	0.2707	94.87%	0.2614	0.2061	0.2461	94.12%
			b	-0.1074	0.0543	0.8107	95.04%	0.0940	0.0508	0.7860	95.26%
			c	-0.4226	0.1788	0.0522	94.64%	-0.0383	0.0263	0.0461	94.68%
		70	a	-0.4278	0.2032	0.2380	95.36%	0.0660	0.0136	0.2072	95.29%
			b	0.0996	0.0479	0.8073	95.97%	0.0814	0.0441	0.5339	95.36%
			c	-0.3729	0.1399	0.0501	95.60%	-0.0355	0.0037	0.0460	95.45%
200	1	150	a	-0.4443	0.2020	0.2673	94.95%	0.0287	0.0036	0.1525	94.36%
			b	0.1517	0.1214	0.9451	95.60%	0.0286	0.0078	0.3183	94.39%
			c	-0.2802	0.0789	0.0761	93.63%	-0.0467	0.0035	0.0728	95.00%
		170	a	-0.3617	0.1400	0.2532	96.27%	0.0197	0.0035	0.1420	95.19%
			b	0.1481	0.0750	0.8080	95.86%	0.0089	0.0061	0.3006	96.29%
			c	-0.2129	0.0464	0.0718	95.54%	-0.0459	0.0031	0.0608	95.94%
	2	150	a	0.0506	0.1080	1.2736	94.73%	0.0838	0.0416	0.5959	95.12%
			b	0.0858	0.0536	0.8431	94.83%	0.0362	0.0094	0.3509	95.11%
			c	0.0202	0.0052	0.2707	94.94%	0.0206	0.0018	0.1437	95.22%
		170	a	0.0468	0.0978	1.2079	95.77%	0.0800	0.0379	0.5703	95.89%
			b	0.0656	0.0397	0.7379	95.67%	0.0278	0.0080	0.3131	95.79%
			c	0.0201	0.0049	0.2610	95.36%	0.0191	0.0016	0.1314	95.90%
	3	150	a	-0.4189	0.2040	0.1317	95.01%	0.0416	0.0351	0.1208	94.90%
			b	0.0916	0.0507	0.7891	95.44%	0.0811	0.0199	0.3375	94.90%
			c	-0.3512	0.1235	0.0521	94.46%	-0.0316	0.0044	0.0461	96.10%
		170	a	-0.4048	0.2030	0.1209	96.96%	0.0299	0.0029	0.1149	95.90%
			b	0.0912	0.0285	0.7822	95.67%	0.0303	0.0073	0.3040	95.70%
			c	-0.2924	0.0861	0.0494	95.73%	-0.0306	0.0031	0.0446	96.80%

Table 5. Optimal Criteria values for different schemes and different cases

Scheme			1			2			3		
Case	n	m	$Optim_1$	$Optim_2$	$Optim_3$	$Optim_1$	$Optim_2$	$Optim_3$	$Optim_1$	$Optim_2$	$Optim_3$
1	100	50	0.70708	0.0000198	1075.33768	0.68243	0.00002054	752.19778	3.35684	0.00520617	140.90745
		70	2.75237	0.0021850	159.53737	0.46881	0.00000603	2206.53088	0.49681	0.00000204	1395.97675
	200	150	0.30134	0.0000015	1448.79984	1.02631	0.00015720	278.21680	0.26005	0.00000287	1000.85472
		170	0.93343	0.0001147	309.92700	0.17239	0.00000040	2503.06675	0.26570	0.00000091	1713.76366
2	100	50	0.20070	0.0000007	3906.37209	0.15764	0.000000676	3238.66376	2.06844	0.001772065	486.64064
		70	1.36299	0.0005915	448.83964	0.11385	0.000000116	5951.49782	0.11411	0.000000139	5015.63199
	200	150	0.06104	0.000000040	5800.39856	0.04940	0.000000101	3657.09628	0.31107	0.000014857	665.17986
		170	0.30643	0.0000125	754.21788	0.05160	0.000000009	10008.14262	0.05397	0.000000025	6710.58635
3	100	50	0.26701	0.000000459	3956.63613	0.19587	0.000000557	2247.74623	1.14685	0.000222097	422.91207
		70	0.56782	0.000034327	493.71070	0.14927	0.000000124	9405.73409	0.20043	0.000000148	4947.81148
	200	150	0.10303	0.000000053	4384.02412	0.07431	0.000000115	2697.28313	0.17722	0.000002390	862.96628
		170	0.15196	0.000001647	945.28273	0.08089	0.000000014	8474.58016	0.08174	0.000000033	5016.80489

5. Applications

We use traditional criteria values to compare fitted models, such as the Information Criterion (IC) of Akaike (ICA), IC of consistent Akaike (ICCA), IC of Bayesian (ICB), IC of Hannan-Quinn (ICHQ), Anderson-Darling Statistic (ADS), Cramer-von-Mises Statistic (CVMS), Kolmogorov-Smirnov Distance (KSD), p-value of Kolmogorov-Smirnov (KSP), and Standard Error (StEr). Our main statistical objective is to apply a fitting approach model to examine three real datasets relevant to different fields. In this context, we compare the fit of the proposed UEHLD with that of the UWD, KumD, BD, UEHLD, UG, and UPBXD.

Data I: Based on Griffiths et al. [59], the data set represents the proportion of income spent on food for each household in the sample. The unit variable represents a numerical variable containing the proportion of income spent on food for each of the 38 households in the sample.** This variable likely falls between 0 (no income spent on food) and 1 (all income spent on food). These data have been obtained as in Table 6.

Data II: The second dataset includes the failure times of an airplane's air conditioning system (measured in hours), as documented by [60]. This "failure times" dataset is: 12, 120, 11, 23, 261, 87, 7, 120, 14, 62, 71, 11, 14, 47, 225, 71, 246, 21, 42, 20, 5, 3, 14, 11, 16, 90, 1, 16, 52, 95. Once more, we perform a normalization procedure by dividing these values by 265 to obtain data ranging from 0 to 1. In other words, we work with the following dataset in Table 6.

Data III: Third, we analyze data on the number of months it takes for renal dialysis patients to become infected, as reported by [61]. The "times of infection" dataset is: 12.5, 13.5, 3.5, 4.5, 5.5, 6.5, 6.5, 7.5, 3.5, 7.5, 12.5, 3.5, 2.5, 2.5, 7.5, 8.5, 9.5, 10.5, 11.5, 7.5, 14.5, 14.5, 21.5, 25.5, 27.5, 21.5, 22.5, 22.5. We now perform a normalization operation by dividing these data by thirty, resulting in values ranging from 0 to 1. The collected data have been updated as following in Table 6.

Table 6. Observed Data Sets Used in the Analysis

Data I	0.2561	0.2023	0.2911	0.1898	0.1619	0.3683	0.2800	0.2068	0.1605
	0.2281	0.1921	0.2542	0.3016	0.2570	0.2914	0.3625	0.2266	0.3086
	0.3705	0.1075	0.3306	0.2591	0.2502	0.2388	0.4144	0.1783	0.2251
	0.2631	0.3652	0.5612	0.2424	0.3419	0.3486	0.3285	0.3509	0.2354
	0.5140	0.5430							
Data II	0.0189	0.0453	0.0868	0.9849	0.3283	0.0264	0.4528	0.0528	0.2340
	0.1962	0.3585	0.1774	0.8491	0.2679	0.9283	0.0792	0.1585	0.0755
	0.4528	0.0415	0.0113	0.0528	0.0415	0.0528	0.2679	0.0415	0.0604
	0.3396	0.0038	0.0604						
Data III	0.4500	0.4833	0.1167	0.8500	0.1167	0.2500	0.2833	0.3167	0.1167
	0.1500	0.1833	0.2167	0.9167	0.2167	0.2500	0.2500	0.0833	0.0833
	0.2500	0.3500	0.3833	0.4167	0.4167	0.7500	0.4833	0.7167	0.7167
	0.7500								

To compare fitted models, we used the measurements of the criteria listed in Tables 7, 8, and 9, respectively, for each data set. Figures 6, 7, and 8 display the dataset's Total Time on Test (TTT), hazard line, estimated CDF with empirical CDF, estimated PDF with histogram, QQ plots, and P-P

Table 7. MLE and some statistical measures for each distribution: Data I

		Estimates	StEr	KSD	KSP	ICA	ICB	ICCA	ICHQ	CVMS	ADS
UIELD	<i>a</i>	0.0098	0.0046	0.0901	0.8904	-66.4426	-61.5299	-65.7367	-64.6947	0.0409	0.3159
	<i>b</i>	90.2878	42.7320								
	<i>c</i>	6.3824	1.9927								
UPBXD	<i>a</i>	2.7831	0.0420	0.1417	0.3935	-66.3996	-60.0272	-64.2341	-63.1920	0.2171	1.4953
	<i>b</i>	0.2007	0.0226								
	<i>c</i>	3161.6115	524.9722								
UWD	<i>a</i>	0.2320	0.0656	0.0964	0.8384	-65.8687	-61.4594	-65.5259	-64.3703	0.0414	0.3296
	<i>b</i>	4.1286	0.4995								
KumD	<i>a</i>	2.9543	0.3691	0.1237	0.5641	-62.9782	-59.7030	-62.6353	-61.8129	0.1193	0.8627
	<i>b</i>	26.9556	10.8204								
BD	<i>a</i>	6.0716	1.3586	0.1101	0.7056	-66.1693	-61.4177	-65.3500	-64.5276	0.0695	0.5196
	<i>b</i>	14.8221	3.3988								
UGD	<i>a</i>	0.0209	0.0120	0.1337	0.4655	-61.2804	-58.0052	-60.9375	-60.1151	0.0937	0.6499
	<i>b</i>	2.6623	0.3297								
UEHLD	<i>a</i>	3.0100	0.3559	0.1192	0.6103	-63.5101	-60.2349	-63.1672	-62.3448	0.1100	0.7991
	<i>b</i>	14.8214	5.6451								

Table 8. MLE and some statistical measures for each distribution: Data II

		Estimates	StEr	KSD	KSP	ICA	ICB	ICCA	ICHQ	CVMS	ADS
UIELD	<i>a</i>	0.0277	0.0283	0.1091	0.8678	-32.2007	-27.9971	-31.2776	-30.8559	0.0615	0.3620
	<i>b</i>	2.1073	1.3892								
	<i>c</i>	0.6370	0.1724								
UWD	<i>a</i>	0.2787	0.0858	0.1742	0.3228	-26.3847	-23.5823	-25.9402	-25.4882	0.1593	1.0191
	<i>b</i>	1.4562	0.2262								
KumD	<i>a</i>	0.5451	0.1148	0.1879	0.2401	-23.0778	-20.2754	-22.6334	-22.1813	0.2110	1.3470
	<i>b</i>	1.3835	0.3361								
BD	<i>a</i>	0.5141	0.1118	0.1958	0.2003	-22.4926	-19.6902	-22.0482	-21.5961	0.2173	1.3859
	<i>b</i>	1.3430	0.3643								
UGD	<i>a</i>	0.4060	0.2500	0.1146	0.8253	-32.1935	-29.3912	-31.7491	-31.2970	0.0838	0.5213
	<i>b</i>	0.4687	0.1350								
UEHLD	<i>a</i>	0.6894	0.1200	0.1786	0.2940	-26.8532	-24.0508	-26.4088	-25.9567	0.1586	1.0162
	<i>b</i>	1.1676	0.2679								
UPBXD	<i>a</i>	1.3207	0.1042	0.1749	0.3179	-30.1454	-25.5465	-29.9435	28.8345	0.3699	2.2458
	<i>b</i>	0.2360	0.0432								
	<i>c</i>	6.1192	1.2555								

Table 9. MLE and some statistical measures for each distribution: Data III

		Estimates	StEr	KSD	KSP	ICA	ICB	ICCA	ICHQ	CVMS	ADS
UIELD	a	0.0067	0.0059	0.1008	0.9384	-6.6025	-2.6059	-5.6025	-5.3807	0.0339	0.2598
	b	47.9149	42.2413								
	c	1.0242	0.2549								
UWD	a	0.6125	0.1424	0.1240	0.7825	-6.1222	-3.4578	-5.6422	-5.3076	0.0660	0.4416
	b	1.6989	0.2668								
KumD	a	1.2652	0.2544	0.1377	0.6628	-3.3249	-0.6605	-2.8449	-2.5104	0.1136	0.7049
	b	2.0799	0.5714								
BD	a	1.3567	0.3332	0.1412	0.6321	-3.5552	-0.8908	-3.0752	-2.7407	0.1101	0.6859
	b	2.1058	0.5496								
UGD	a	0.5085	0.4500	0.1127	0.8691	-6.0811	-5.4167	-7.6011	-7.2665	0.0360	0.2689
	b	0.7921	0.3578								
UEHLD	a	1.4691	0.2530	0.1265	0.7613	-4.7989	-2.1345	-4.3189	-3.9844	0.0835	0.5455
	b	1.5705	0.3992								
UPBXD	a	1.5323	0.0961	0.1367	0.6718	-4.3124	-1.3158	-3.3124	-3.0906	0.1765	1.0386
	b	0.3018	0.0504								
	c	6.7265	1.3549								

plots for the UIELD for each dataset. These graphical goodness-of-fit methods in Figures 6, 7, and 8 also support the results shown in Tables 7, 8, and 9 for each dataset.

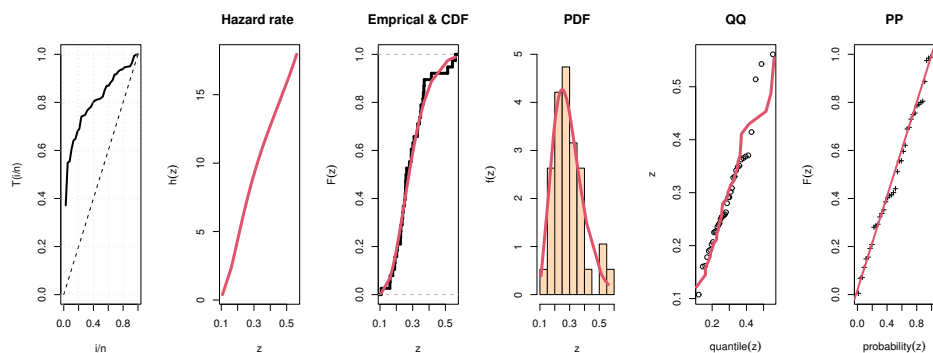
**Figure 6.** Different plot for fitting for UIELD: Data I

Figure 12, 13, and 14 illustrates the profile log-likelihood of the UIELD for each parameter by fixing one parameter and varying the other. Figures 9, 10, and 11 demonstrate that each dataset performs well, as evidenced by the two roots of the parameters representing global maxima. Figures 12, 13, and 14 present contour plots with varying parameters and the log-likelihood of the UIELD, confirming that the estimates have unique points.

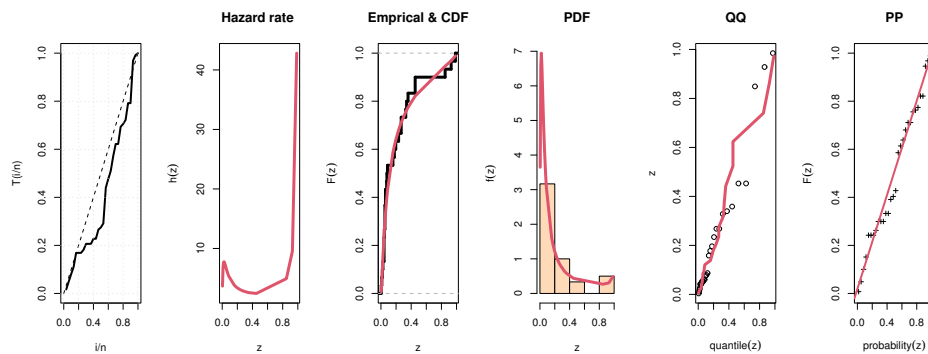


Figure 7. Different plot for fitting for UIELD: Data II

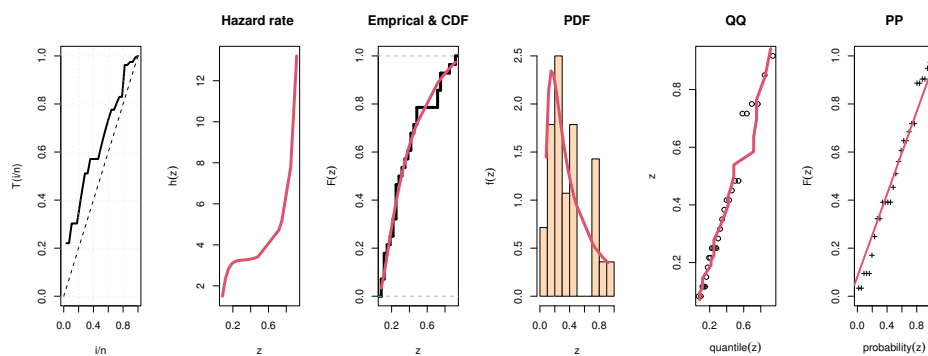


Figure 8. Different plot for fitting for UIELD: Data III

6. Summary and Conclusion

This paper presents a novel probability distribution known as UIELD. It is intended exclusively to represent data that is contained inside the unit interval, or 0 to 1. A variety of shapes, including left-skewed, reversed J-shaped, U-shaped, and even more complex patterns, are available for the density function of the UIELD. Compared to simpler distributions, this enables it to fit a wider range of data. The UIELD's HRF, in contrast to the density function, shows more distinct shapes, such as J-, U-, ascending, and decreasing shapes. The UIELD's mathematical features are explored, including formulas for its moments, incomplete moments, quantile function, and other significant properties. Both conventional and Bayesian techniques for estimating the UIELD's parameters are examined. Estimation procedures are considered under PTII censoring schemes, a technique to gather data when observation times are costly or inconvenient. Simulations are used to assess the effectiveness of point and interval estimates from both techniques based on some measures of precision. In addition, the selection of the optimal progressive censoring scheme is investigated under three different optimization criteria. Simulation research demonstrated that the MCMC approach with the SEL function performed better than the MLE in terms of MSE, bias, and lower LCI values for different estimates of parameters. When considering interval estimates, the LCCI values were significantly lower than the LACI values. Furthermore, most of the coverage probabilities were high, ranging around 95%. This implies that the constructed CIs captured the true parameter values with a probability of 95%. A comparison of the three schemes revealed that Scheme 2 outperformed Schemes 1 and 3, indicating its superiority in

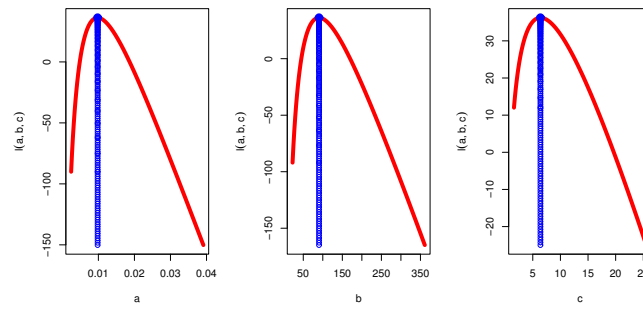


Figure 9. Profile likelihood for parameters of UIELD: Data I

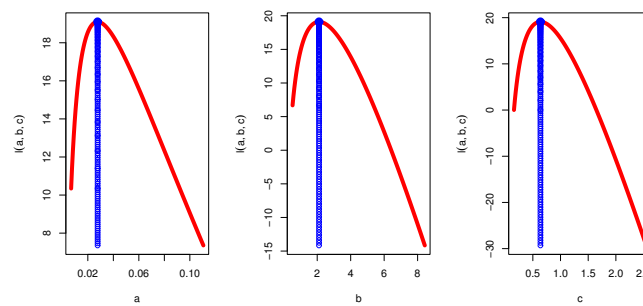


Figure 10. Profile likelihood for parameters of UIELD: Data II

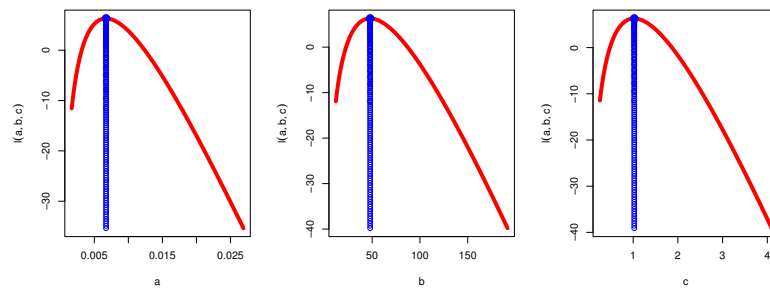


Figure 11. Profile likelihood for parameters of UIELD: Data III

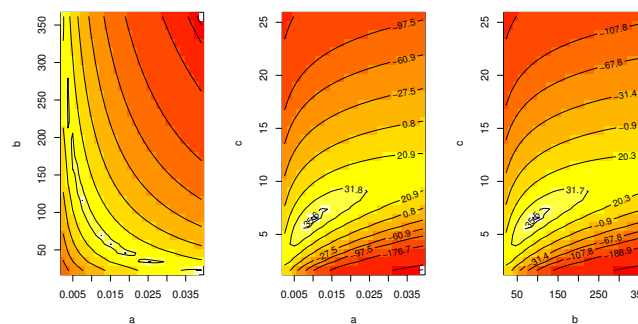


Figure 12. Contour plot for parameters of UIELD: Data I

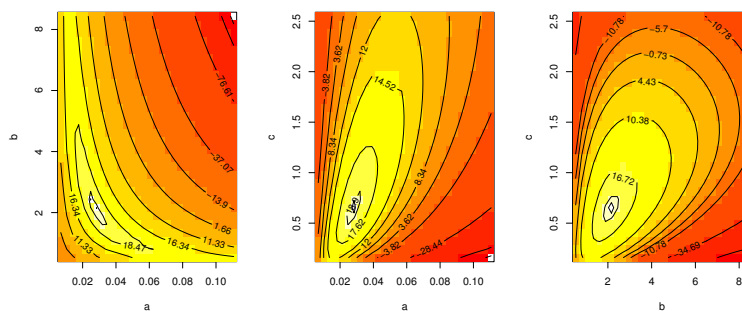


Figure 13. Contour plot for parameters of UIELD: Data II

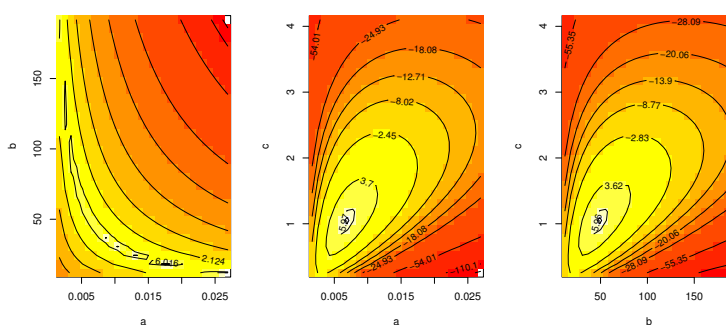


Figure 14. Contour plot for parameters of UIELD: Data III

terms of the proposed optimization criteria. The adaptability of the UIELD is validated by three real data sets, which show that it achieves a close fit compared to alternatives. This study is limited by its use of the MCMC method within the Bayesian estimation framework and symmetric loss function. As future work, the Bayesian estimation study of the proposed model using some asymmetric loss function and Tierney-Kadane approximation method can be discussed under PTII censoring schemes with random removals.

Funding Statement: This work was supported and funded by the Deanship of Scientific Research at Imam Mohammad Ibn Saud Islamic University (IMSIU) (grant number IMSIU-DDRSP2502).

Data Availability: The data that support the findings of this study are available upon request from the corresponding author.

Author contributions: All authors have accepted responsibility for the entire content of this manuscript and approved its submission.

References

1. Kumaraswamy, P. (1980). A generalized probability density function for double-bounded random processes. *Journal of Hydrology*, 46(1-2), 79-88. [https://doi.org/10.1016/0022-1694\(80\)90036-0](https://doi.org/10.1016/0022-1694(80)90036-0)

2. Mazucheli, J., Menezes, A. F. B., & Dey, S. (2018). The unit-Birnbaum-Saunders distribution with applications. *Chilean Journal of Statistics*, 9, 47-57.
3. Mazucheli, J., Menezes, A. F. B., & Chakraborty, S. (2019). On the one parameter unit-Lindley distribution and its associated regression model for proportion data. *Journal of Applied Statistics*, 46(4), 700-714.
4. Ghitany, M. E., Mazucheli, J., Menezes, A. F. B., & Alqallaf, F. (2019). The unit-inverse Gaussian distribution: A new alternative to two parameter distributions on the unit interval. *Communications in Statistics—Theory and Methods*, 48(14), 3423-3438.
5. Mazucheli, J., Menezes, A. F. B., Fernandes, L. B., de Oliveira, R. P., & Ghitany, M. E. (2020). The unit-Weibull distribution as an alternative to the Kumaraswamy distribution for the modeling of quantiles conditional on covariates. *Journal of Applied Statistics*, 47(5), 954-974.
6. Mazucheli, J., Menezes, A. F., & Dey, S. (2019). Unit-Gompertz distribution with applications. *Statistica (Bologna)*, 79(1), 25-43.
7. Hassan, A. S., Fayomi, A., Algarni, A., & Almetwally, E. M. (2022). Bayesian and non-Bayesian inference for unit-exponentiated half-logistic distribution with data analysis. *Applied Sciences*, 12(21), 11253.
8. Modi, K., & Gill, V. (2020). Unit Burr-III distribution with application. *Journal of Statistical Management Systems*, 23(3), 579-592.
9. Korkmaz, M. C., & Chesneau, C. (2021). On the unit Burr-XII distribution with the quantile regression modeling and applications. *Computational and Applied Mathematics*, 40(1), 29.
10. Altun, E. and Cordeiro, G.M. (2020). The unit-improved second-degree Lindley distribution: inference and regression modeling *Computational Statistics* (2020) 35:259-279 <https://doi.org/10.1007/s00180-019-00921-y>
11. Krishna, A., Maya, R., Chesneau, C., & Irshad, M. R. (2022). The unit Teissier distribution and its applications. *Mathematics and Computers in Simulation*, 27(1), 12.
12. Fayomi, A., Hassan, A.S., and Almetwally, E.A. (2023). Inference and quantile regression for the unit exponentiated Lomax distribution. *PLoS ONE* 18(7): e0288635. <https://doi.org/10.1371/journal.pone.0288635>
13. Ramadan, A.T.; Tolba, A.H.; El-Desouky, B.S. (2022). A Unit Half-Logistic Geometric Distribution and Its Application in Insurance. *Axioms* 2022, 11, 676. <https://doi.org/10.3390/axioms11120676>
14. Hassan, A.S., and Alharbi, R. S. (2023). Different estimation methods for the unit inverse exponentiated Weibull distribution. *Communications for Statistical Applications and Methods*, 30(2), 191-213
15. Mazucheli, J., Korkmaz, M. C., Menezes, A. F., & Leiva, V. (2023). The unit generalized half-normal quantile regression model: Formulation, estimation, diagnostics, and numerical applications. *Soft Computing*, 27(1), 279-295. <https://doi.org/10.1007/s00500-022-07278-3>
16. Fayomi, A.; Hassan, A.S.; Baaqeel, H.M.; Almetwally, E.M. (2023). Bayesian Inference and Data Analysis of the Unit-Power Burr X Distribution. *Axioms* 2023, 12, 297. <https://doi.org/10.3390/axioms12030297>

17. Bakouch, H. S., Nikb, A. S., Asgharzadehb, A., & Salinas, H. S. (2021). A flexible probability model for proportion data: Unit-half-normal distribution. *Communications in Statistics—Case Studies and Data Analysis*, 7(2), 271-288.
18. Hassan, A.S., Khalil, A.M. and Nagy, H.F. (2024). Data Analysis and Classical Estimation Methods of the Bounded Power Lomax Distribution. *Reliability Theory & Applications*, 19(1), 770-789.
19. BiCer, C.; Bakouch, H.S.; Bicer, H.D.; Alomair, G.; Hussain, T.; Almohisen, A. (2024). Unit Maxwell Boltzmann Distribution and Its Application to Concentrations Pollutant Data. *Axioms* 2024, 13, 226. <https://doi.org/10.3390/axioms13040226>
20. Karakaya, K., Rajitha, C. S., Saglam, S., Tashkandy, Y.A., Bakr, M. E., Muse, A.H., Kumar, A. Hussam, E. and Gemeay, A.M. (2024). A new unit distribution: properties, estimation, and regression analysis. *Scientific Reports* 14, 7214 (2024). <https://doi.org/10.1038/s41598-024-57390-7>
21. Holland, O., Golaup, A., & Aghvami, A. (2006). Traffic characteristics of aggregated module downloads for mobile terminal reconfiguration. *IEE Proceedings—Communications*, 153(5), 683-690.
22. Corbellini, A., Crosato, L., Ganugi, P and Mazzoli, M. (2007). Fitting Pareto II distributions on firm size: Statistical methodology and economic puzzles. Paper presented at the International Conference on Applied Stochastic Models and Data Analysis, Chania, Crete
23. Hassan A., & Al-Ghamdi A. (2009). Optimum step stress accelerated life testing for Lomax distribution. *Journal of Applied Sciences Research*, 5, 2153-2164.
24. Harris, C. M. (1968). The Pareto distribution as a queue service discipline. *Operations Research*, 16(2), 307-313.
25. Atkinson, A.B. & Harrison, A.J. (1978). *Distribution of Personal Wealth in Britain*. Cambridge University Press, Cambridge
26. Ghitany, M. E., Al-Awadhi, F.A. & Alkhalfan, L.A. (2007). Marshall Olkin extended Lomax distribution and its application to censored data. *Communications in Statistics-Theory & Methods*, 36, 1855-1866.
27. Abdul-Moniem, I.B.; Abdel-Hameed, H.F. On exponentiated Lomax distribution, *Int. J. Math. Arch.*, 2012, 3, 2144-2150.
28. Lemonte, A.J.; & Cordeiro, G.M. An extended Lomax distribution, *Statistics*, 2013, 47, 800-816.
29. Tahir, M. H., Cordeiro, G. M., Mansoor, M., & Zubair, M. (2015). The Weibull-Lomax distribution: properties and applications. *Hacettepe Journal of Mathematics and Statistics*, 44(2), 461-480.
30. Hassan, A.S., Almetwally, E.M., Khaleel, M.A., & Nagy, H.F. (2021). Weighted Power Lomax Distribution and its Length Biased Version: Properties and Estimation based on Censored Samples. *Pakistan Journal of Statistics and Operation Research*, 17(2), 343-356.
31. Cordeiro, M., Ortega, E. M. M., & Popovic, B. V. (2015). The gamma-Lomax distribution. *Journal of Statistical Computation and Simulation*, 85(2), 305-319.
32. Hassan, A.S.; & Abd-Alla, M. Exponentiated Weibull-Lomax distribution: properties and estimation, *Journal of Data Sciences*, 2018, 16(2), 275-298.
33. Hassan, A.S., Elgarhy, M. & Mohamed, R.E. (2020). Statistical properties and estimation of type II half logistic Lomax distribution. *Thailand Statistician*, 18(3): 290-305.

34. Nagarjuna, V. B. V., Vardhan, R. V., & Chesneau, C. (2022). Nadarajah-Haghighi Lomax distribution and its applications. *Mathematical and Computational Applications*, 27(30), 1-13. <https://doi.org/10.3390/mca27020030>
35. Alsubie, A. (2021). Properties and Applications of the Modified Kies-Lomax Distribution with Estimation Methods. *Journal of Mathematics* Volume 2021, Article ID 1944864, 18 pages <https://doi.org/10.1155/2021/1944864>
36. Hassan, A.S., and Mohamed, R.E (2019). Parameter Estimation of Inverse Exponentiated Lomax distribution with Right Censored Data. *Gazi University Journal of Science*, 32(4), 1370-1386.
37. Rényi, A. (1960). On measures of entropy and information. *Proceedings of the 4th Berkeley Symposium on Mathematical Statistics and Probability*, 1, 47-561.
38. Tsallis, C. (1988). Possible generalization of Boltzmann-Gibbs statistics. *Journal of Statistical Physics*, 52(1-2), 479-487.
39. Arimoto, S. (1971). Information-theoretical considerations on estimation problems. *Information and Control*, 19(3), 181-194.
40. Balakrishnan, N., & Aggarwala, R. (2000). *Progressive censoring: theory, methods, and applications*. Springer Science & Business Media.
41. Meredith, M.; & Kruschke, J. HDInterval: Highest (posterior) density intervals. R Package Version 0.1. 2016. Available online: (accessed on 20 March 2023)
42. Plummer, M.; Best, N.; Cowles, K.; Vines, & K. CODA: convergence diagnosis and output analysis for MCMC. *R news* **2006**, 6, 7-11.
43. El-Saeed, A. R., & Almetwally, E. M. (2024). On Algorithms and Approximations for Progressively Type-I Censoring Schemes. *Statistical Analysis and Data Mining: The ASA Data Science Journal*, 17(6), e11717.
44. Mohamed, A. A., Refaey, R. M., & AL-Dayian, G. R. (2024). Bayesian and E-Bayesian estimation for odd generalized exponential inverted Weibull distribution. *Journal of Business and Environmental Sciences*, 3(2), 275-301.
45. Albadawy, A., Ashour, E., EL-Helbawy, A. A., & AL-Dayian, G. R. (2024). Bayesian estimation and prediction for exponentiated inverted Topp-Leone distribution. *Computational Journal of Mathematical and Statistical Sciences*, 3(1), 33-56.
46. Almetwally, E. M., Jawa, T. M., Sayed-Ahmed, N., Park, C., Zakarya, M., & Dey, S. (2023). Analysis of unit-Weibull based on progressive type-II censored with optimal scheme. *Alexandria Engineering Journal*, 63, 321-338.
47. Dey, S., Singh, S., Tripathi, Y. M., & Asgharzadeh, A. (2016). Estimation and prediction for a progressively censored generalized inverted exponential distribution. *Statistical Methodology*, 32, 185-202.
48. Metropolis, N.; Rosenbluth, A. W.; Rosenbluth, M. N.; Teller, A. H.; Teller, E. Equation of state calculations by fast computing machines. *The j. chem. phys.* **1953**, 21, 1087-1092.
49. Hastings, W. K. Monte Carlo sampling methods using Markov chains and their applications. *Biometrika* **1970**, 57, 97-109.

50. Gelfand, A. E.; Smith, A. F. Sampling-based approaches to calculating marginal densities. *J. the Amer. stat. assoc.* **1990**, 85, 398-409.
51. Burkschat, M. (2008). On optimality of extremal schemes in progressive type II censoring. *Journal of Statistical Planning and Inference*, 138(6), 1647-1659.
52. Pradhan, B., & Kundu, D. (2009). On progressively censored generalized exponential distribution. *Test*, 18, 497-515.
53. Ashour, S.K., El-Sheikh, A.A. & Elshahhat, A. Inferences and Optimal Censoring Schemes for Progressively First-Failure Censored Nadarajah-Haghighi Distribution. *Sankhya A* 84, 885–923 (2022). <https://doi.org/10.1007/s13171-019-00175-2>
54. Dey, S. & Elshahhat, A. (2022). Analysis of Wilson-Hilferty distribution under progressive Type-II censoring. *Quality and Reliability Engineering International*, 38(7), 3771–3796.
55. Nassar, M., & Elshahhat, A. (2023). Estimation procedures and optimal censoring schemes for an improved adaptive progressively type-II censored Weibull distribution. *Journal of Applied Statistics*, 51(9), 1664–1688. <https://doi.org/10.1080/02664763.2023.2230536>
56. Alotaibi, R.; Nassar, M.; & Elshahhat, A. Estimation and Optimal Censoring Plan for a New Unit Log-Log Model via Improved Adaptive Progressively Censored Data. *Axioms* 2024, 13, 152. <https://doi.org/10.3390/axioms13030152>
57. Alotaibi, R., Nassar, M., & Elshahhat, A. (2024). Rainfall data modeling using improved adaptive type-II progressively censored Weibull-exponential samples. *Scientific Reports*, 14, Article 30484. <https://doi.org/10.1038/s41598-024-80529-5>
58. Elshahhat, A., A.H., Egeh, O.M. & Elemar, B.R. (2022) Estimation for Parameters of Life of the Marshall-Olkin Generalized-Exponential Distribution Using Progressive Type-II Censored Data. *Complexity*, Volume 2022, Article ID 8155929, 36 pages <https://doi.org/10.1155/2022/8155929>
59. Griffiths, W.E., Hill, R.C., & Judge, G.G. (1993). *Learning and Practicing Econometrics* New York: John Wiley and Sons.
60. Linhart, H.; & Zucchini, W. *Model Selection*; Wiley: New York, NY, USA, 1986.
61. Klein, J. P., & Moeschberger, M. L. (2006). *Survival analysis: Techniques for censored and truncated data*. Springer.



© 202 by the authors. Disclaimer/Publisher's Note: The content in all publications reflects the views, opinions, and data of the respective individual author(s) and contributor(s), and not those of the scientific association for studies and applied research (SASAR) or the editor(s). SASAR and/or the editor(s) explicitly state that they are not liable for any harm to individuals or property arising from the ideas, methods, instructions, or products mentioned in the content.