
Research article

A New Odd Reparameterized Exponential Transformed-X Family of Distributions with Applications to Public Health Data

Gabriel O. Orji¹, Harrison O. Etaga², Ehab M. Almetwally³, Chinyere P. Igbokwe⁴, Obioma Chukwudi Aguwa⁵, Okechukwu J. Obulezi^{6,*}

¹ Department of Statistics, Faculty of Physical Sciences, Nnamdi Azikiwe University, Awka 420110, Nigeria

² Department of Statistics, Faculty of Physical Sciences, Nnamdi Azikiwe University, Awka 420110, Nigeria

³ College of Science, Imam Mohammad Ibn Saud Islamic University (IMSIU), Riyadh 11432, Saudi Arabia

⁴ Department of Statistics, School of Computer Science and Engineering, Lovely Professional University, Phagwara, Punjab, 144411, India.

⁵ Department of Computer Science and Engineering, Faculty of Arts and Science, Edge Hill University, Ormskirk L39 4QP, United Kingdom

⁶ Department of Statistics, Faculty of Physical Sciences, Nnamdi Azikiwe University, Awka 420110, Nigeria

* **Correspondence:** oj.obulezi@unizik.edu.ng

ARTICLE INFO

Keywords:

New odd reparameterized exponential transformed-X
Moment
Mean residual life function
COVID-19
HIV/AIDS
Weibull distribution

Mathematics Subject Classification:

62F15, 62F10, 60E05

Important Dates:

Received: 2 May 2025
Revised: 20 June 2025
Accepted: 27 June 2025
Online: 30 June 2025



Copyright © 2025 by the authors. Published under Creative Commons Attribution (CC BY) license.

ABSTRACT

In this study, we designed a new family of distributions called new odd reparameterized exponential transformed-X family of distributions. This was achieved by utilizing an exponential distribution with a constant scale parameter as the transformer and then the odd function of the baseline distribution as the generalizer. The new family was used to extend the classical Weibull distribution. We further studied the characteristics of the new extended Weibull distribution which include the quantile function, moment, moment generating function, mean residual life function, order statistic, entropy and extropy. Again, the parameters were estimated using both non-Bayesian and Bayesian approaches. A comprehensive Monte Carlo simulation was conducted under four different scenarios. The proposed distribution was fitted to HIV/AIDS and COVID-19 data and then compared with the baseline distribution (Weibull) and other related models. The new distribution commands superior fit with a probability value of 0.9959 and 0.7086 in the two datasets respectively.

1. Introduction

Modeling complex phenomena and life events is a cornerstone of modern scientific inquiry across diverse fields, including environmental science, engineering, biological studies, and economics. The inherent unpredictability and variability within these phenomena present significant analytical challenges. To accurately capture and explain these random or unpredictable realities, researchers increasingly rely on robust statistical models that can effectively incorporate outcome uncertainty. The ability of such models to provide a nuanced understanding of risk, reliability, and event likelihood is paramount for informed decision-making and predictive analysis.

Over the past three decades, the statistical literature has seen a surge of interest in developing more flexible and versatile tools for this purpose. A key innovation in this domain has been the construction of families of distributions and generators for these families. These advancements allow for the creation of new, more adaptable statistical models that can better fit real-world data, particularly when traditional distributions fall short in capturing complex behaviors like skewed data, varying hazard rates, or heavy tails. Among these methodological advancements, the T-X generator, introduced by [8], stands out as a highly influential framework for synthesizing novel distributions. Various families of distributions have been proposed in statistical literature, each offering unique properties for modeling diverse datasets. These include the Logistic-X family by [33], the Odd-Perks-G family by [20], and the Kumaraswamy-G family by [17]. Further contributions encompass the T-X generator of families of continuous distributions by [8], the Burr-Hatke-G family by [38], and the Poisson generated family of distributions by [29]. Other notable families are the Gompertz-G family by [4], the Power Lindley-G family by [23], and the T-normal family by [9]. The Odd Lindley-G family by [22], the Topp-Leone Family of Distributions by [2], and the Odd Chen-G family by [11] also represent significant advancements. Additionally, the Sine generalized family of distributions by [30], the Type I half-logistic family of distributions by [16], ordered families of distributions by [27], and the Type II general inverse exponential family of distributions by [25] collectively showcase the ongoing innovation in constructing flexible statistical models.

Building upon these general frameworks, researchers have also developed specific distributions tailored to particular applications. For example, within these broad families, [28] introduced the Logistic exponentiated-exponential distribution, while [19] proposed the Odd-Perks-Lomax distribution. The Kumaraswamy Pareto distribution was detailed by [12], and Gompertz-power series distributions were explored by [24]. More recently, [10] presented the Topp-Leone Burr-Hatke-exponential distribution, and [3] developed the Exponentiated power Lindley power series class of distributions. Other specific models include the Odd Lindley-Pareto distribution by [34], the Kumaraswamy quasi Lindley distribution by [21], and the Type II Topp Leone power Lomax distribution by [1]. Furthermore, the half-logistic inverse Rayleigh distribution was examined by [5], and the Topp-Leone modified Weibull distribution by [7], all contributing to a richer toolkit for statistical analysis.

The development of the New Odd Reparametrized Exponential Transformed-X (NORET-X) family of distributions is driven by the critical need for a more robust and flexible approach to modeling complex hazard phenomena. Many real-world events, like certain environmental risks, exhibit intricate patterns in their occurrence and intensity that traditional statistical distributions struggle to capture. Specifically, existing models often lack the adaptability to accurately represent hazard functions and the unique tail behaviors associated with extreme or rare events. The NORET-X family addresses these limitations by building upon the well-understood exponential distribution. It introduces a reparametrization combined with an odd func-

tion transformation, significantly enhancing its flexibility for hazard-based applications. This innovative modification allows the ORET-X family to more precisely account for both moderate and extreme event levels, leading to improved accuracy in modeling hazard rates and their associated risks. Moreover, by integrating both Bayesian and non-Bayesian estimation methods, the NORET-X family achieves greater modeling precision and improves the stability of parameter estimates across diverse sample sizes. This new distribution family therefore offers significant promise across various fields, providing a sophisticated statistical tool for examining and predicting a wide range of complex hazards.

This article is structured to guide the reader through the development and application of the proposed statistical framework. Section 2 details the development of the NORET-X family of distributions. Building upon this, Section 3 focuses specifically on the NORET-Weibull distribution. The key characteristics of this new distribution are then explored in Section 4. Section 5 outlines the estimation methodologies used, followed by Section 6, which presents the results from simulation studies. The practical application of the framework is demonstrated in Section 7, and finally, Section 8 provides concluding remarks.

2. Construction of the NORET-X Family of Distributions

In this study, we construct a new distribution from the exponential distribution by parametrization. The probability density function (PDF) of the reparametrized exponential distribution (RE) is $r(t) = \frac{\pi}{2} e^{-\frac{\pi t}{2}}$, $t > 0$ with a new scale weight, $\frac{\pi}{2}$. The cumulative distribution function (CDF) easily follows from the integration of $r(t)$ from zero up to a certain t , hence $R(t) = 1 - e^{-\frac{\pi t}{2}}$. One can then write that $T \sim \text{RE} \left(\frac{\pi}{2} \right)$. Next, we follow the argument of [8], utilizing the odd function $W[G(x; \nabla)] = \frac{G(x; \nabla)}{1 - G(x; \nabla)}$ as the generalizer in order to realized a new family of distributions, where $G(x; \nabla)$ is the CDF of any parent distribution. The CDF and PDF of the new family of distributions are expressed as

$$F(x; \nabla) = \int_0^{\frac{G(x; \nabla)}{1 - G(x; \nabla)}} r(t) dt = 1 - e^{-\frac{\pi G(x; \nabla)}{2[1 - G(x; \nabla)]}}, \quad (2.1)$$

and

$$f(x; \nabla) = \frac{\partial W[G(x; \nabla)]}{\partial x} r\{W[G(x; \nabla)]\} = \frac{\pi g(x; \nabla)}{2(1 - G(x; \nabla))^2} e^{-\frac{\pi G(x; \nabla)}{2[1 - G(x; \nabla)]}}, \quad (2.2)$$

respectively. ∇ is the parameter vector for any parent distribution with CDF $G(x; \nabla)$. The hazard function $h(x; \nabla) = \frac{f(x; \nabla)}{1 - F(x; \nabla)}$ is

$$h(x; \nabla) = \frac{\alpha g(x; \nabla)}{(e - 1)(1 - G(x; \nabla))^2}. \quad (2.3)$$

This new family is referred to as the "New Odd Reparametrized Exponential Transformed-X" (NORET-X) family of distributions. Notice that $r(t)$ is a valid PDF and $\frac{G(x; \nabla)}{1 - G(x; \nabla)}$ satisfies the criteria provided in the $T - X$ generator of family of distributions by [8]. That is, $\int_0^\infty \frac{\pi}{2} e^{-\frac{\pi t}{2}} dt = 1$ and

(a) $W\{G(x; \nabla)\} \in [a, b]$.

(b) W is differentiable and monotonically non-decreasing.

(c) $Z\{G(x; \nabla)\} \rightarrow a$ as $x \rightarrow -\infty$ and $W\{G(x; \nabla)\} \rightarrow b$ as $x \rightarrow \infty$,

where $[a, b]$ is the domain of the random variable T such that $-\infty \leq a < b \leq \infty$.

Introducing $\frac{\pi}{2}$ as a multiplier to the parameter of the exponential distribution effectively scales the parameter down by this constant factor. This change will influence the distribution in some key ways:

1. Rate of Decay: The value $\frac{\pi}{2} \approx 1.5708$ is greater than 1. A larger rate parameter λ in an exponential distribution ($\lambda e^{-\lambda t}$) implies a faster rate of probability decay. Thus, events will, on average, occur more frequently, and the distribution will be more concentrated closer to zero.
2. Mean and Variance: For a standard exponential distribution with parameter λ , the mean is $1/\lambda$ and the variance is $1/\lambda^2$. With $\lambda = \frac{\pi}{2}$, the reparametrized exponential distribution will have:
 - (a) Mean: $E[T] = \frac{1}{\pi/2} = \frac{2}{\pi} \approx 0.6366$.
 - (b) Variance: $Var[T] = \frac{1}{(\pi/2)^2} = \frac{4}{\pi^2} \approx 0.4053$.

This means the distribution will be more peaked and have less spread compared to an exponential distribution with a rate parameter of, for example, 1.

3. Hazard Rate: The hazard rate of the RE distribution is constant, given by $h(t) = \frac{\pi}{2}$. This reflects the memoryless property of the exponential distribution, but the specific value of $\frac{\pi}{2}$ dictates the instantaneous failure rate.
4. Quantiles: All quantiles of the distribution will be compressed towards the lower end of the support due to the higher rate parameter. For example, the median $M = \frac{\ln(2)}{\lambda} = \frac{\ln(2)}{\pi/2} = \frac{2\ln(2)}{\pi} \approx 0.4413$.

These characteristics of the reparametrized exponential distribution $r(t)$ will, in turn, propagate through the construction of the NORET-X family, affecting its overall shape, moments, and hazard behavior. The choice of $\frac{\pi}{2}$ as the fixed rate parameter provides a specific baseline behavior for the T-X generator, influencing the resulting flexibility and fitting capabilities of the NORET-X family.

3. The NORET-Weibull (NORETW) Distribution

The Weibull distribution by [37] is a versatile continuous probability distribution widely used in fields such as reliability engineering, survival analysis, extreme value theory, and material science, due to its ability to model various failure rate behaviors. Its PDF and CDF are respectively given as

$$g(x; \beta, \lambda) = \alpha \beta x^{\alpha-1} e^{-\beta x^\alpha}$$

$$\text{and } G(x; \beta, \lambda) = 1 - e^{-\beta x^\alpha}; \quad x \geq 0,$$

where $\alpha > 0$ is the shape parameter and $\beta > 0$ is the scale parameter. The NORETW distribution CDF and PDF are obtained by plugging back the $G(x; \alpha, \beta)$ and $g(x; \alpha, \beta)$ of Weibull distribution into equations (2.1) and (2.2). Thus:

$$F(x; \alpha, \beta) = 1 - e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)}; \quad x > 0, \quad (3.1)$$

and

$$f(x; \alpha, \beta) = \frac{\pi\alpha\beta}{2} x^{\alpha-1} e^{\beta x^\alpha} e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)}, \quad (3.2)$$

where $\beta > 0$ is the shape parameter and $\alpha > 0$ and $\lambda > 0$ are the scale parameters.

The hazard function of the NORETW distribution is

$$h(x; \alpha, \beta) = \frac{\pi\alpha\beta}{2} x^{\alpha-1} e^{\beta x^\alpha}. \quad (3.3)$$

The NORETW hazard function exhibits flexible shapes depending on the value of α :

1. If $\alpha = 1$: Constant hazard rate, starting at $\frac{\pi\beta}{2}$ and remaining constant.
2. If $\alpha > 1$: Increasing hazard rate, starting at 0 and increasing towards ∞ . This represents an increasing failure rate (IFR).
3. If $0 < \alpha < 1$: Decreasing hazard rate, starting at ∞ and decreasing towards ∞ (though the rate of decrease slows before it eventually increases due to exponential dominance, which means it will pass through a minimum if it's considered for a very broad range, but the overall trend for $x \rightarrow \infty$ is still increasing). This initial part represents a decreasing failure rate (DFR).

The term $e^{\beta x^\alpha}$ ensures that, for all $\alpha > 0, \beta > 0$, the hazard function will eventually increase to infinity as x becomes very large, making it suitable for modeling situations where aging or wear-out effects eventually lead to higher failure probabilities.

Figure (1) contains the PDF of the NORETW distribution with majorly positively skewed shape and reversed bathtub shape. Figures (2) are the plots of the hazard rate, which is traditionally used as a mortality indicator in survival analysis. The plots represent strictly decreasing, strictly increasing, L-shape and J-shape respectively.

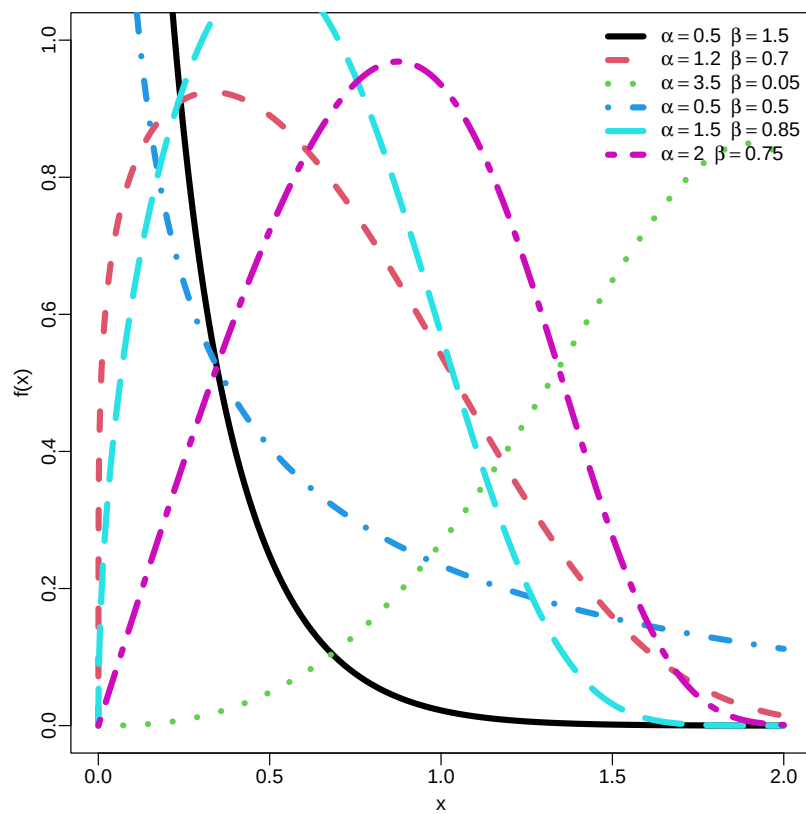


Figure 1. $f(x; \alpha, \beta, \lambda)$ of NORETW

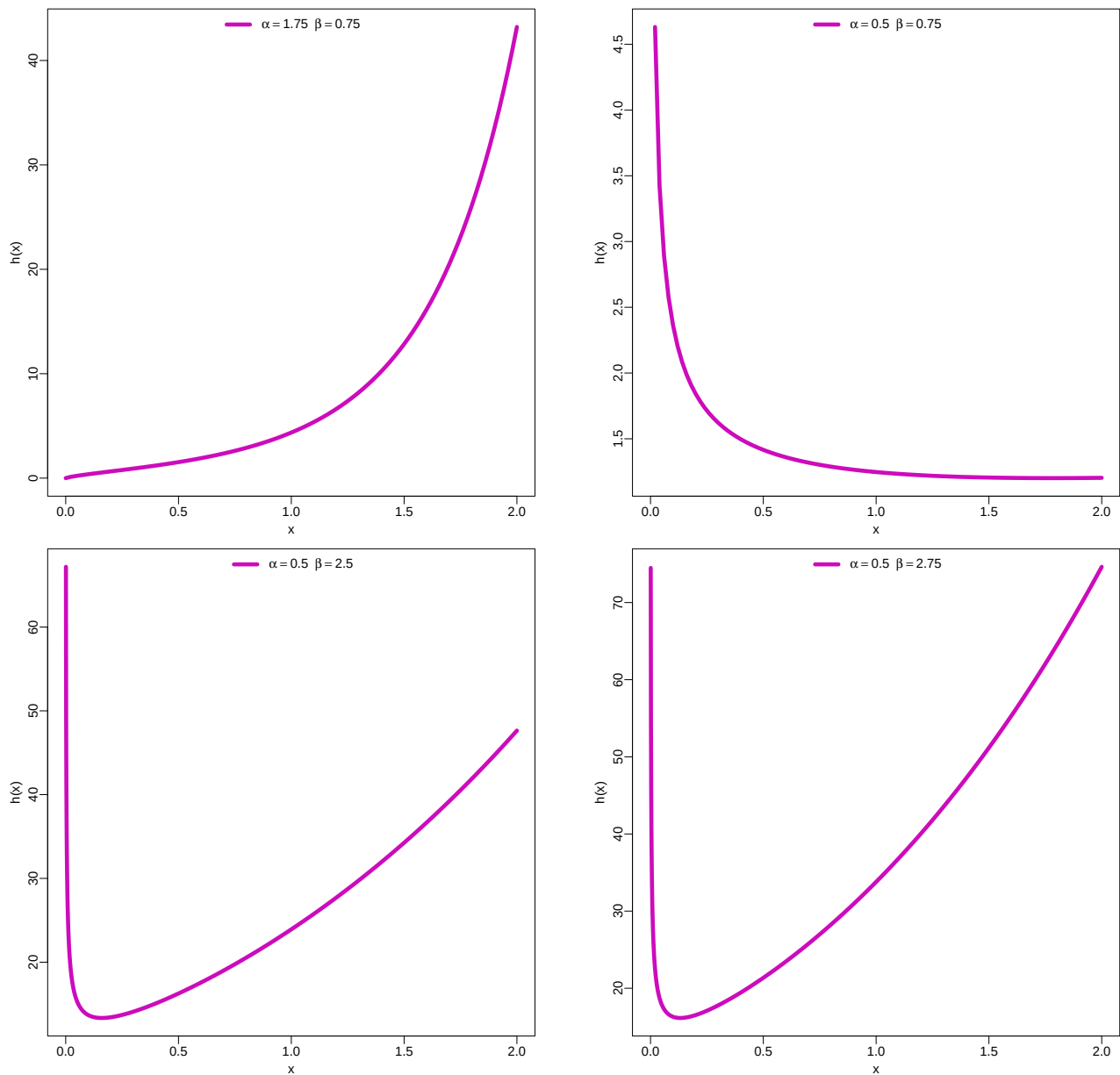


Figure 2. hazard function

4. Characteristics

In this section, we study the characteristics of the proposed NORETW distribution which include the quantile function, moment, moment generating function, mean residual life function, order statistic, entropy and extropy.

4.1. Quantile Function

The quantile function is derived from the CDF by inversion method. So let $p \sim U(0, 1)$, then replacing $F(x; \alpha, \beta)$ in Eq. (3.1) yields

$$Q(q; \alpha, \beta) = \left[\frac{1}{\beta} \ln \left(1 - \frac{2}{\pi} \ln(1 - q) \right) \right]^{1/\alpha}; \quad \text{for } 0 < p < 1. \quad (4.1)$$

4.2. Moment

The r th crude moment of the random variable X that assumes the NORETW distribution with parameters α and β is given as follows

$$\mu'_r = \mathbb{E}(X^r) = \int_0^\infty x^r f(x) dx = \frac{\pi\alpha\beta}{2} \int_0^\infty x^{\alpha+r-1} e^{\beta x^\alpha} e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)} dx$$

Using the generalized power series and binomial expansions

$$e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)} = \sum_{i=0}^{\infty} \frac{(-1)^i}{i!} \left(\frac{\pi}{2}\right)^i (e^{\beta x^\alpha} - 1)^i.$$

$$\text{Also } (e^{\beta x^\alpha} - 1)^i = \sum_{j=0}^i (-1)^j \binom{i}{j} e^{(i-j)\beta x^\alpha},$$

$$\text{so that } e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)} = \sum_{i=0}^{\infty} \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^i e^{(i-j)\beta x^\alpha}.$$

$$\text{Therefore } \mu'_r = \frac{\pi\alpha\beta}{2} \sum_{i=0}^{\infty} \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^i \int_0^\infty x^{\alpha+r-1} e^{-(j-i-1)\beta x^\alpha} dx$$

Let $z = (j - i - 1)\beta x^\alpha \implies x = \left[\frac{z}{(j-i-1)\beta} \right]^{\frac{1}{\alpha}}; dx = \frac{z^{\frac{1}{\alpha}-1}}{\alpha[(j-i-1)\beta]} dz$. So that,

$$\mu'_r = \frac{\pi\beta}{2} \sum_{i=0}^{\infty} \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^i \int_0^\infty \frac{z^{\frac{r}{\alpha}} e^{-z}}{[(j-i-1)\beta]^{\frac{r+1}{\alpha}}} dz$$

The final series form of the r th crude moment of X is

$$\mu'_r = \beta^{\frac{\alpha-r-1}{\alpha}} \sum_{i=0}^{\infty} \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^{i+1} (j-i-1)^{-\left(\frac{r+1}{\alpha}\right)} \Gamma\left[\frac{r}{\alpha} + 1\right]; \quad r = 1, 2, \dots \quad (4.2)$$

4.3. Moment generating function

The moment generating function (mgf), where it exists is expressed as $M_X(t) = \mathbb{E}(e^{tX})$ and it is given as

$$\begin{aligned}
 M_X(t) &= \int_0^{\infty} e^{tx} \frac{\pi\alpha\beta}{2} x^{\alpha-1} e^{\beta x^\alpha} e^{-\frac{\pi}{2}(e^{\beta x^\alpha}-1)} dx \\
 &= \beta^{\frac{\alpha-k-1}{\alpha}} \sum_{i=0}^{\infty} \sum_{j=0}^i \sum_{k=0}^{\infty} \frac{(-1)^{i+j}}{i!k!} t^k \binom{i}{j} \left(\frac{\pi}{2}\right)^{i+1} (j-i-1)^{-\left(\frac{k+1}{\alpha}\right)} \Gamma\left[\frac{k}{\alpha} + 1\right].
 \end{aligned}$$

4.4. The m th order statistic

The distribution of the m th order statistic for the order sample $X_{(1)}, X_{(2)}, \dots, X_{(m)}$ is given as

$$\begin{aligned}
 f_{X_{(m)}}(x) &= \binom{n}{m} [F(x)]^{m-1} [1 - F(x)]^{n-m} f(x) \\
 &= \binom{n}{m} \left(\frac{\pi\alpha\beta}{2} x^{\alpha-1} e^{\beta x^\alpha}\right) \left[1 - e^{-\frac{\pi}{2}(e^{\beta x^\alpha}-1)}\right]^{m-1} e^{-\frac{\pi}{2}(n-m+1)(e^{\beta x^\alpha}-1)}
 \end{aligned}$$

Using generalized power series and binomial expansions

$$\begin{aligned}
 \left[1 - e^{-\frac{\pi}{2}(e^{\beta x^\alpha}-1)}\right]^{m-1} &= \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} \sum_{k=0}^j \frac{(-1)^{i+j+k}}{j!} \binom{m-1}{i} \binom{j}{k} \left(\frac{\pi i}{2}\right)^j e^{(j-k)\beta x^\alpha} \\
 &= e^{(j-k)\beta x^\alpha} A_{ijk},
 \end{aligned}$$

where

$$A_{ijk} = \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} \sum_{k=0}^j \frac{(-1)^{i+j+k}}{j!} \binom{m-1}{i} \binom{j}{k} \left(\frac{\pi i}{2}\right)^j$$

Similarly,

$$\begin{aligned}
 e^{-\frac{\pi}{2}(n-m+1)(e^{\beta x^\alpha}-1)} &= \sum_{l=0}^{\infty} \sum_{m=0}^l \frac{(-1)^{l+m}}{l!} \binom{l}{m} \left(\frac{\pi}{2}\right)^l (n-m+1)^l e^{(l-m)\beta x^\alpha} \\
 &= A_{lm} e^{(l-m)\beta x^\alpha}
 \end{aligned}$$

Therefore,

$$f_{X_{(m)}}(x) = \binom{n}{m} \frac{\pi\alpha\beta}{2} A_{ijk} A_{lm} x^{\alpha-1} e^{-(k-j-l+m-1)\beta x^\alpha} \quad (4.3)$$

4.5. Mean Residual Life Function

The mean residual life (MRL) is the expected remaining lifetime of an item or individual given that it has already survived up to a certain age x . Suppose a system component has been functioning for x units of time, the MRL function tells on average how much longer it is expected to continue functioning. In the

instance of $X \sim \text{NORETW}(\alpha, \beta)$, the MRL is

$$\begin{aligned} m(x) &= \frac{1}{1-F(x)} \int_x^\infty (1-F(t)) dt = e^{\frac{\pi}{2}(e^{\beta x^\alpha}-1)} \int_x^\infty e^{-\frac{\pi}{2}(e^{\beta t^\alpha}-1)} dt \\ &= e^{\frac{\pi}{2}(e^{\beta x^\alpha}-1)} \sum_{i=0}^\infty \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^i \int_x^\infty e^{-(j-i)\beta t^\alpha} dt \\ &= \frac{e^{\frac{\pi}{2}(e^{\beta x^\alpha}-1)}}{\alpha} \sum_{i=0}^\infty \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi}{2}\right)^i \frac{\Gamma\left[\frac{1}{\alpha}, (j-i)\beta x^\alpha\right]}{[(j-i)\beta]^{\frac{1}{\alpha}}}. \end{aligned}$$

4.6. Entropy

This is the amount of information loss or uncertainty in a system. The common measure is the Rény entropy by [31]. It is defined as

$$R_E = \frac{1}{1-\gamma} \log \left[\int_0^\infty f^\gamma(x) dx \right] = \frac{1}{1-\gamma} \log \left[\int_0^\infty \left(\frac{\pi\alpha\beta}{2}\right)^\gamma x^{\gamma\alpha-\gamma} e^{\gamma\beta x^\alpha} e^{-\frac{\pi\gamma}{2}(e^{\beta x^\alpha}-1)} dx \right],$$

where $\gamma > 0$ and $\gamma \neq 1$. The final form of the Rény entropy is

$$R_E = \frac{1}{1-\gamma} \log \left[\left(\frac{\pi\beta}{2}\right)^\gamma \alpha^{\gamma-1} \sum_{i=0}^\infty \sum_{j=0}^i \frac{(-1)^{i+j}}{i!} \binom{i}{j} \left(\frac{\pi\gamma}{2}\right)^i \{(j-i-\gamma)\beta\}^{-(\gamma-\frac{\gamma}{\alpha}+\frac{1}{\alpha})} \Gamma\left(\gamma-\frac{\gamma}{\alpha}+\frac{1}{\alpha}\right) \right] \quad (4.4)$$

4.7. Extropy

This quantifies the randomness of a distribution in a dual sense. Essentially, extropy measures the expected quadratic distance from the true probability to a perfect prediction, ([26]). Mathematically, extropy is defined thus;

$$J(X) = -\frac{1}{2} \int_0^\infty f^2(x) dx = -\frac{\pi^2\alpha^2\beta^2}{8} e^\pi \int_0^\infty x^{2\alpha-2} e^{2\beta x^\alpha} e^{-\pi e^{\beta x^\alpha}} dx$$

Using generalized power series and binomial expansions, the final form of the extropy is

$$J(X) = -\frac{\pi^2\alpha^2\beta^2}{8} e^\pi \sum_{i=0}^\infty \frac{(-1)^i}{i!} \pi^i [(-2-i)\beta]^{-(2-\frac{1}{\alpha})} \Gamma\left(2-\frac{1}{\alpha}\right) \quad (4.5)$$

Figures (3-6) are the 3-dimensional plots of the Mean, Variance, Skewness and Kurtosis respectively.

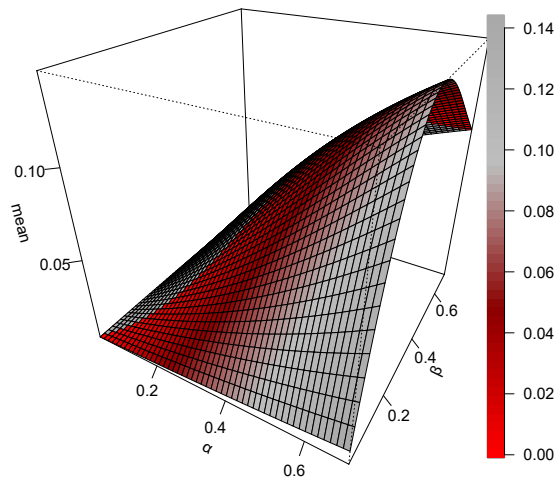


Figure 3. Mean of NORETW

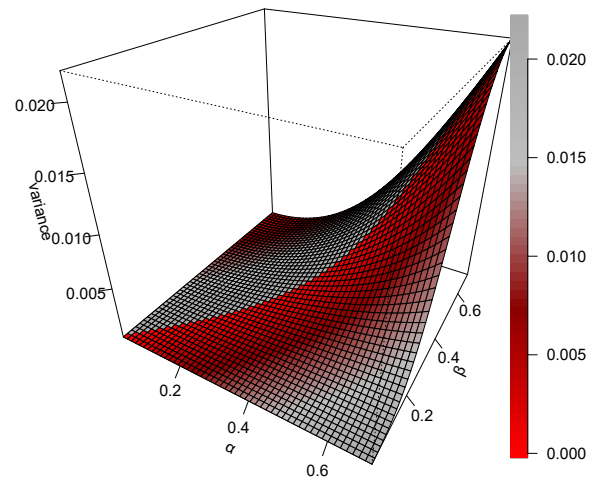


Figure 4. Variance of NORETW

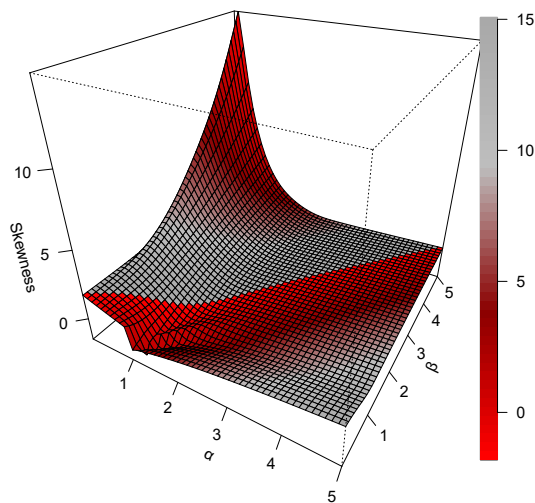


Figure 5. Skewness of NORETW

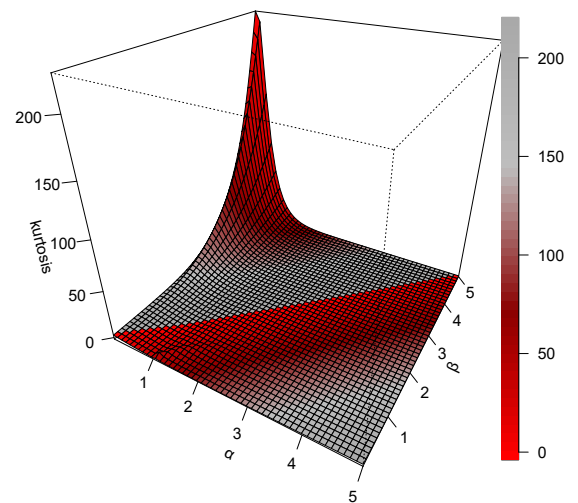


Figure 6. Kurtosis of NORETW

5. Statistical Inference

This section addresses the parameter estimation methods for the NORETW distribution, employing both Bayesian and classical approaches. By utilizing multiple methods, we enable a comprehensive evaluation of their respective efficiencies, particularly for large sample sizes.

5.1. Maximum Likelihood Estimation (MLE)

Suppose $x_1, x_2, \dots, x_n$ represent a sample of n independent observations derived from the NORETW distribution with parameters α , and β . The likelihood function, denoted as $L(X|\Theta)$, along with its corresponding log-likelihood function, serves as the basis for parameter estimation. These functions are expressed as:

$$\begin{aligned} L(\alpha, \beta) &= \prod_{i=1}^n f(x_i; \alpha, \beta) \\ &= \prod_{i=1}^n \left[\frac{\pi\alpha\beta}{2} x_i^{\alpha-1} e^{\beta x_i^\alpha} e^{-\frac{\pi}{2}(e^{\beta x_i^\alpha} - 1)} \right] \\ &= \left(\frac{\pi\alpha\beta}{2} \right)^n \left(\prod_{i=1}^n x_i^{\alpha-1} \right) \left(\prod_{i=1}^n e^{\beta x_i^\alpha} \right) \left(\prod_{i=1}^n e^{-\frac{\pi}{2}(e^{\beta x_i^\alpha} - 1)} \right) \end{aligned}$$

Taking the natural logarithm of the likelihood function simplifies the product into a sum, which is easier to differentiate:

$$\begin{aligned} \ell(\alpha, \beta) &= \ln L(\alpha, \beta) \\ &= \ln \left[\left(\frac{\pi\alpha\beta}{2} \right)^n \prod_{i=1}^n x_i^{\alpha-1} \prod_{i=1}^n e^{\beta x_i^\alpha} \prod_{i=1}^n e^{-\frac{\pi}{2}(e^{\beta x_i^\alpha} - 1)} \right] \\ &= n \ln \left(\frac{\pi\alpha\beta}{2} \right) + \sum_{i=1}^n \ln(x_i^{\alpha-1}) + \sum_{i=1}^n \ln(e^{\beta x_i^\alpha}) + \sum_{i=1}^n \ln \left(e^{-\frac{\pi}{2}(e^{\beta x_i^\alpha} - 1)} \right) \\ &= n(\ln \pi + \ln \alpha + \ln \beta - \ln 2) + (\alpha - 1) \sum_{i=1}^n \ln x_i + \beta \sum_{i=1}^n x_i^\alpha - \frac{\pi}{2} \sum_{i=1}^n (e^{\beta x_i^\alpha} - 1) \end{aligned}$$

To find the MLEs, $\hat{\alpha}$ and $\hat{\beta}$, we take the partial derivatives of the log-likelihood function with respect to each parameter and set them equal to zero.

Partial Derivative with respect to β

$$\begin{aligned} \frac{\partial \ell}{\partial \beta} &= \frac{\partial}{\partial \beta} \left[n \ln \beta + \beta \sum_{i=1}^n x_i^\alpha - \frac{\pi}{2} \sum_{i=1}^n (e^{\beta x_i^\alpha} - 1) \right] \\ &= \frac{n}{\beta} + \sum_{i=1}^n x_i^\alpha - \frac{\pi}{2} \sum_{i=1}^n (e^{\beta x_i^\alpha} \cdot x_i^\alpha) \end{aligned}$$

Setting $\frac{\partial \ell}{\partial \beta} = 0$:

$$\frac{n}{\hat{\beta}} + \sum_{i=1}^n x_i^{\hat{\alpha}} - \frac{\pi}{2} \sum_{i=1}^n x_i^{\hat{\alpha}} e^{\hat{\beta} x_i^{\hat{\alpha}}} = 0 \quad (5.1)$$

Partial Derivative with respect to α

$$\frac{\partial \ell}{\partial \alpha} = \frac{\partial}{\partial \alpha} \left[n \ln \alpha + (\alpha - 1) \sum_{i=1}^n \ln x_i + \beta \sum_{i=1}^n x_i^\alpha - \frac{\pi}{2} \sum_{i=1}^n (e^{\beta x_i^\alpha} - 1) \right]$$

Recall the derivative rules: $\frac{\partial}{\partial \alpha}(x_i^\alpha) = x_i^\alpha \ln x_i$ and $\frac{\partial}{\partial \alpha}(e^{cx_i^\alpha}) = e^{cx_i^\alpha} \cdot cx_i^\alpha \ln x_i$. Applying these:

$$\frac{\partial \ell}{\partial \alpha} = \frac{n}{\alpha} + \sum_{i=1}^n \ln x_i + \beta \sum_{i=1}^n (x_i^\alpha \ln x_i) - \frac{\pi}{2} \sum_{i=1}^n (e^{\beta x_i^\alpha} \cdot \beta x_i^\alpha \ln x_i)$$

Setting $\frac{\partial \ell}{\partial \alpha} = 0$:

$$\frac{n}{\hat{\alpha}} + \sum_{i=1}^n \ln x_i + \hat{\beta} \sum_{i=1}^n x_i^{\hat{\alpha}} \ln x_i - \frac{\pi \hat{\beta}}{2} \sum_{i=1}^n x_i^{\hat{\alpha}} \ln x_i e^{\hat{\beta} x_i^{\hat{\alpha}}} = 0 \quad (5.2)$$

The parameters β and α are estimated by solving equations (5.1) and (5.2) simultaneously. Due to the non-linear nature of these equations, numerical methods, such as the Newton-Raphson algorithm, are utilized to iteratively refine the parameter estimates, ensuring the maximization of the likelihood function.

5.2. Least Squares Estimation (LSE)

The Least Squares Estimation (LSE) method, first introduced by [32] for parameter estimation in the Beta distribution, provides the foundation for our approach. Building on their methodology, we estimate parameters by minimizing the sum of squared deviations between observed and theoretical values. This technique ensures that the distribution model closely aligns with the data, offering precise parameter estimates. By emphasizing the fit between the theoretical distribution and empirical observations, this method remains a cornerstone for reliable parameter estimation.

$$E[F(x_{j:n}|\alpha, \beta)] = \frac{j}{n+1}; \quad \text{and} \quad V[F(x_{j:n}|\alpha, \beta)] = \frac{j(n-j+1)}{(n+1)^2(n+2)}.$$

The parameters α , and β are estimated using the Least Squares Estimation (LSE) method, producing the estimates $\hat{\alpha}_{\text{LSE}}$ and $\hat{\beta}_{\text{LSE}}$. These estimates are derived by minimizing the function $L(\alpha, \beta)$ with respect to the parameters. This approach identifies parameter values that minimize the squared differences between observed and theoretical values, ensuring the model achieves the best possible fit to the data.

$$L(\alpha, \beta) = \arg \min_{(\alpha, \beta)} \sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2.$$

The parameters α , and β are estimated by solving a system of non-linear equations that arise from the minimization process. These equations are constructed by equating the partial derivatives of the least squares objective function $L(\alpha, \beta)$ with respect to α , and β to zero. The solutions to this system provide the Least Squares Estimators, denoted as $\hat{\alpha}_{\text{LSE}}$, and $\hat{\beta}_{\text{LSE}}$.

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_1(x_{j:n}|\alpha, \beta) = 0, \quad (5.3)$$

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_2(x_{j:n}|\alpha, \beta) = 0 \quad (5.4)$$

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_3(x_{j:n}|\alpha, \beta) = 0 \quad (5.5)$$

Where

$$\Delta_1(x_{j:n}|\alpha, \beta) = \frac{\pi}{2} x^\alpha e^{\beta x^\alpha} e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)} \quad (5.6)$$

and

$$\Delta_2(x_{j:n}|\alpha, \beta) = \frac{\pi}{2} \beta x^\alpha (\ln x) e^{\beta x^\alpha} e^{-\frac{\pi}{2}(e^{\beta x^\alpha} - 1)}, \quad (5.7)$$

The expressions in equations (5.6), and (5.7) are derived by taking the partial derivatives of the cumulative distribution function (CDF) of the ORET-L distribution, as defined in equation (3.1), with respect to the parameters α , and β . These derivatives form the basis for the equations used in parameter estimation, enabling an assessment of how variations in each parameter influence the CDF. This process ultimately facilitates the computation of the Least Squares Estimates for the parameters.

5.3. Weighted Least Squares Estimation (WLSE)

The Weighted Least Squares Estimation (WLSE) method is employed for parameter estimation in the ORET-L distribution, where the goal is to estimate the parameters α , and β . By minimizing the weighted least squares function $W(\alpha, \beta)$, the estimates $\hat{\alpha}_{WLSE}$, and $\hat{\beta}_{WLSE}$ are obtained. Unlike standard least squares, the WLSE incorporates weights w_j to account for the varying precision of different observations. This adjustment enhances the accuracy of the parameter estimates by assigning greater weight to more reliable or relevant data points.

The weighted least squares function is given by:

$$W(\alpha, \beta) = \arg \min_{(\alpha, \beta)} \sum_{j=1}^n w_j \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \quad (5.8)$$

where the weight w_j for each observation is defined as:

$$w_j = \frac{(n+1)^2 (n+2)}{j(n-j+1)}.$$

The parameter estimates are derived by solving the following system of equations, where the weighted sums of squared deviations are equated to zero:

$$\sum_{j=1}^n w_j \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_1(x_{j:n}|\alpha, \beta) = 0, \quad (5.9)$$

$$\sum_{j=1}^n w_j \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_2(x_{j:n}|\alpha, \beta) = 0, \quad (5.10)$$

$$\sum_{j=1}^n w_j \left[F(x_{j:n}|\alpha, \beta) - \frac{j}{n+1} \right]^2 \Delta_3(x_{j:n}|\alpha, \beta) = 0. \quad (5.11)$$

Here, $\Delta_1(x_{j:n}|\alpha, \beta)$, $\Delta_2(x_{j:n}|\alpha, \beta)$, and $\Delta_3(x_{j:n}|\alpha, \beta)$ are the partial derivatives of the cumulative distribution function (CDF) of the ORET-L distribution with respect to the parameters α , and β , respectively, as defined in equations (5.6), and (5.7). These derivatives are central to the computation of the weighted least squares estimates, as they describe how each parameter influences the CDF.

5.4. Maximum Product of Spacing Estimation (MPSE)

The Maximum Product of Spacing Estimation (MPSE) method, introduced by [15], provides an alternative to the traditional maximum likelihood estimation by focusing on the Kullback-Leibler information criterion rather than maximizing the likelihood function directly. This method is particularly useful for ordered data, as it maximizes the product of the spacings between the cumulative distribution function (CDF) values of ordered observations. By minimizing the discrepancy between the observed and predicted distributions, the MPSE method tends to provide robust and reliable parameter estimates, especially in situations where the maximum likelihood estimation might be problematic.

The product of spacings function $P_s(X|\alpha, \beta)$ is given by:

$$P_s(X|\alpha, \beta) = \left[\prod_{j=1}^{n+1} D_k(x_{j:n}|\alpha, \beta) \right]^{\frac{1}{n+1}},$$

where $D_k(x_{j:n}|\alpha, \beta)$ represents the spacing between the CDF values of consecutive observations, defined as:

$$D_k(x_{j:n}|\alpha, \beta) = F(x_j|\alpha, \beta) - F(x_{j-1}|\alpha, \beta), \quad j = 1, 2, 3, \dots, n.$$

To estimate the parameters, we maximize the following function $M_s(\alpha, \beta)$, which is the logarithm of the product of spacings:

$$M_s(\alpha, \beta) = \frac{1}{n+1} \sum_{j=1}^{n+1} \ln(D_k(x_{j:n}|\alpha, \beta)). \quad (5.12)$$

The parameter estimates are obtained by differentiating the function $M_s(\alpha, \beta)$ with respect to α , and β , and solving the resulting system of nonlinear equations:

$$\frac{\partial M_s(\alpha, \beta)}{\partial \alpha} = 0, \quad \frac{\partial M_s(\alpha, \beta)}{\partial \beta} = 0.$$

Solving this system yields the parameter estimates that maximize the product of spacings function and satisfy the given conditions.

5.5. Cramér-von Mises Estimation (CVME)

The Cramér-von Mises estimation method is used to estimate the parameters α , and β of the NORETW distribution. The estimates, denoted as $\hat{\alpha}_{\text{CVME}}$, and $\hat{\beta}_{\text{CVME}}$ are derived by minimizing the Cramér-von Mises criterion function $W(\alpha, \beta)$ with respect to each parameter. This method aims to minimize the squared differences between the empirical distribution function and the theoretical cumulative distribution function, providing a robust approach for parameter estimation that better captures the model's fit to the observed data.

The Cramér-von Mises criterion function $W(\alpha, \beta)$ is given by:

$$W(\alpha, \beta) = \arg \min_{(\alpha, \beta)} \left\{ \frac{1}{12n} + \sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{2j-1}{2n} \right]^2 \right\}.$$

The estimates are obtained by solving the following set of non-linear equations:

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{2j-1}{2n} \right] \Delta_1(x_{j:n}|\alpha, \beta) = 0, \quad (5.13)$$

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{2j-1}{2n} \right] \Delta_2(x_{j:n}|\alpha, \beta) = 0, \quad (5.14)$$

$$\sum_{j=1}^n \left[F(x_{j:n}|\alpha, \beta) - \frac{2j-1}{2n} \right] \Delta_3(x_{j:n}|\alpha, \beta) = 0, \quad (5.15)$$

where $\Delta_1(x_{j:n} | \alpha, \beta)$, and $\Delta_2(x_{j:n} | \alpha, \beta)$ are defined as outlined in equations (5.6), and (5.7), respectively.

5.6. Anderson-Darling Estimation (ADE)

The Anderson-Darling estimators $\hat{\alpha}_{ADE}$, and $\hat{\beta}_{ADE}$, for the parameters α , and β of the NORETW distribution are obtained by minimizing the function $T(\alpha, \beta)$ with respect to α , and β .

The function $T(\alpha, \beta)$ is given by:

$$T(\alpha, \beta) = \arg \min_{(\alpha, \beta)} \sum_{j=1}^n (2j-1) \left\{ \ln \left(F(x_{j:n}|\alpha, \beta) \right) + \ln \left[1 - F(x_{n+1-j:n}|\alpha, \beta) \right] \right\}$$

The estimates are derived by solving the following set of non-linear equations:

$$\sum_{j=1}^n (2j-1) \left[\frac{\Delta_1(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} - \frac{\Delta_1(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.16)$$

$$\sum_{j=1}^n (2j-1) \left[\frac{\Delta_2(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} - \frac{\Delta_2(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.17)$$

$$\sum_{j=1}^n (2j-1) \left[\frac{\Delta_3(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} - \frac{\Delta_3(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.18)$$

where $\Delta_1(x_{j:n} | \alpha, \beta)$, and $\Delta_2(x_{j:n} | \alpha, \beta)$ are defined as specified in equations (5.6), and (5.7), respectively.

5.7. Right-Tailed Anderson-Darling Estimation (RTADE)

The estimates $\hat{\alpha}_{RTADE}$, and $\hat{\beta}_{RTADE}$ for the parameters α , and β of the NORETW distribution, derived using the Right-Tailed Anderson-Darling method, are obtained by minimizing the function $T_r(\alpha, \beta)$ with respect to α , and β .

The function $T_r(\alpha, \beta)$ is given by:

$$T_r(\alpha, \beta) = \arg \min_{(\alpha, \beta)} \left\{ \frac{n}{2} - 2 \sum_{j=1}^n F(x_{j:n}|\alpha, \beta) - \frac{1}{n} \sum_{j=1}^n (2j-1) \ln \left[1 - F(x_{n+1-j:n}|\alpha, \beta) \right] \right\}.$$

The estimates are derived by solving the following set of non-linear equations:

$$-2 \sum_{j=1}^n \frac{\Delta_1(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} + \frac{1}{n} \sum_{j=1}^n (2j-1) \left[\frac{\Delta_1(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.19)$$

$$-2 \sum_{j=1}^n \frac{\Delta_2(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} + \frac{1}{n} \sum_{j=1}^n (2j-1) \left[\frac{\Delta_2(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.20)$$

$$-2 \sum_{j=1}^n \frac{\Delta_3(x_{j:n}|\alpha, \beta)}{F(x_{j:n}|\alpha, \beta)} + \frac{1}{n} \sum_{j=1}^n (2j-1) \left[\frac{\Delta_3(x_{n+1-j:n}|\alpha, \beta)}{1 - F(x_{n+1-j:n}|\alpha, \beta)} \right] = 0, \quad (5.21)$$

where $\Delta_1(x_{j:n} | \alpha, \beta)$, and $\Delta_2(x_{j:n} | \alpha, \beta)$ are defined as detailed in equations (5.6), and (5.7), respectively. The estimates presented in equations (5.1), (5.2), (5.3), (5.4), (5.5), (5.9), (5.10), (5.11), (5.12), (5.13), (5.14), (5.15), (5.16), (5.17), (5.18), (5.19), (5.20), (5.21) were obtained using the *optim()* function in R, which applies the Newton-Raphson iterative method.

5.8. Bayesian Estimation

This section focuses on deriving the Bayesian estimates (BE) for the unknown parameters of the ORET-L distribution. In Bayesian parameter estimation, various loss functions can be used, such as the squared error, LINEX, and generalized entropy loss functions. For parameter estimation, we assume that independent gamma priors are assigned to α , and β , and the corresponding probability density functions (PDFs) form the prior distribution for the NORETW model.

$$\left. \begin{aligned} \pi_1(\alpha) &\propto \alpha^{p_1-1} e^{-q_1\alpha}, & \alpha > 0, p_1 > 0, q_1 > 0 \\ \pi_2(\beta) &\propto \beta^{p_2-1} e^{-q_2\beta}, & \beta > 0, p_2 > 0, q_2 > 0 \end{aligned} \right\} \quad (5.22)$$

The hyperparameters p_j and q_j , for $j = 1, 2$, are chosen based on prior knowledge of the parameters. The joint prior distribution for $\Theta = (\alpha, \beta)$ is expressed as follows:

$$\left. \begin{aligned} \pi(\Theta) &= \pi_1(\alpha)\pi_2(\beta) \\ \pi(\Theta) &\propto \alpha^{p_1-1}\beta^{p_2-1}e^{-q_1\alpha-q_2\beta} \end{aligned} \right\} \quad (5.23)$$

The posterior probability distribution, conditional on the observed data $X = (x_1, x_2, \dots, x_n)$, is represented by the following expression:

$$\pi(\Theta | X) = \frac{\pi(\Theta)\iota(\Theta)}{\int_{\Theta} \pi(\Theta)\iota(\Theta) d\Theta}.$$

For any function $\vartheta(\Theta)$, the Bayes estimator under the squared error loss (SEL) criterion is given by:

$$\hat{\Theta}_{BESEL} = E[\vartheta(\Theta)|x] = \int_{\Theta} \vartheta(\Theta)\pi(\Theta|x)d\Theta. \quad (5.24)$$

The squared error loss (SEL) treats underestimations and overestimations symmetrically due to its symmetric loss function. However, in real-world applications, either overestimating or underestimating may have significantly different consequences. Therefore, the LINEX loss function can be a better alternative to SEL, as described by

$$(\vartheta(\Theta), \hat{\vartheta}(\Theta)) = e^{\{\hat{\vartheta}(\Theta) - \vartheta(\Theta)\}} - \kappa(\hat{\vartheta}(\Theta) - \vartheta(\Theta)) - 1,$$

where $\kappa \neq 0$ is a shape parameter. If $\kappa > 1$, it indicates that overestimations are more crucial than underestimations, while $\kappa < 0$ indicates the opposite. As κ approaches zero, the loss function converges to the standard squared error (SE) loss. More details on this can be found in [36] and [18]. The Bayes estimator (BE) of $\vartheta(\Theta)$ under this loss function is given by:

$$\hat{\Theta}_{\text{BE.LINEX}} = \mathbb{E} \left[e^{-\kappa \vartheta(\Theta)} \mid x \right] = -\frac{1}{\kappa} \log \left(\int_{\Theta} e^{-\kappa \vartheta(\Theta)} \pi(\Theta \mid x) d\Theta \right). \quad (5.25)$$

We also introduce the Generalized Entropy Loss (GEL) function, as defined by [14], which is given by:

$$(\vartheta(\Theta), \hat{\vartheta}(\Theta)) = \left(\frac{\hat{\vartheta}(\Theta)}{\vartheta(\Theta)} \right)^{-\beta} - \beta \log \left(\frac{\hat{\vartheta}(\Theta)}{\vartheta(\Theta)} \right) - 1.$$

The shape parameter $\beta \neq 0$ indicates the deviation from symmetry. When $\beta > 0$, overestimations are more critical than underestimations, while $\beta < 0$ suggests the reverse. In this case, the Bayes estimator under the GEL is computed as:

$$\hat{\Theta}_{\text{BE.GEL}} = \mathbb{E} \left[(\vartheta(\Theta))^{-\beta} \mid x \right]^{-1/\beta} = \left(\int_{\Theta} (\vartheta(\Theta))^{-\beta} \pi(\Theta \mid x) d\Theta \right)^{-1/\beta}. \quad (5.26)$$

As the estimates derived from equations (5.24), (5.25), and (5.26) do not have closed-form expressions, posterior samples must be generated and Bayes estimates computed using the Markov Chain Monte Carlo (MCMC) method (refer to [13] and [35] for further details on MCMC). In MCMC, an initial batch of samples, of size N , is drawn from the posterior distribution, and a "burn-in" period is applied, discarding the first ϑ_b samples. The remaining samples are used to calculate the Bayes estimates. Using MCMC for SEL, LINEX, and GEL loss functions, the Bayes estimates $\Theta^{(j)} = (\alpha^{(j)}, \beta^{(j)})$ are computed as follows:

$$\hat{\Theta}_{\text{BE.SEL}} = \frac{1}{N - \vartheta_b} \sum_{j=\vartheta_b}^N \Theta(j), \quad (5.27)$$

$$\hat{\Theta}_{\text{BE.LINEX}} = -\frac{1}{\kappa} \log \left(\frac{1}{N - \vartheta_b} \sum_{j=\vartheta_b}^N e^{-\kappa \Theta(j)} \right), \quad (5.28)$$

and

$$\hat{\Theta}_{\text{BE.GEL}} = \left(\frac{1}{N - \vartheta_b} \sum_{j=\vartheta_b}^N \Theta(j)^{-\beta} \right)^{-\frac{1}{\beta}}, \quad (5.29)$$

where ϑ_b denotes the number of burn-in samples.

6. Simulation

This study applies a variety of non-Bayesian estimation methods for comparison, including Maximum Likelihood (MLE), Maximum Product of Spacings (MPSE), Least Squares (LSE), Weighted Least Squares (WLSE), Cramér-von Mises (CVME), Anderson-Darling (ADE), and Right-Tail Anderson-Darling

(RTADE). For each method, parameter estimates were obtained through 10,000 bootstrap samples, evaluated across four distinct sample sizes. In addition to these non-Bayesian methods, the Markov Chain Monte Carlo (MCMC) technique was employed to generate posterior samples for Bayesian estimation, as no closed-form solutions could be derived directly from equations (5.25), and (5.26).

To enhance the robustness of the Bayesian estimations, a bootstrap procedure was integrated with MCMC. For each bootstrap sample, parameter estimates $\Theta^{(j)} = (\alpha^{(j)}, \beta^{(j)})$ were calculated, with an initial burn-in period applied to the MCMC chain to ensure it had reached convergence. This process was repeated for each sample size ($n = 25, 75, 150, 200$), resulting in 10,000 bootstrap iterations. Bayesian estimates were subsequently derived using three distinct loss functions: squared error loss (SEL), LINEX loss, and generalized entropy loss (GEL).

Case I: $\alpha = 1.75, \beta = 2.75$

Case II: $\alpha = 2.0, \beta = 1.25$

Case III: $\alpha = 2.75, \beta = 2.0$

Case IV: $\alpha = 2.50, \beta = 1.75$

Table 1. Bias and RMSE for Simulation Case I

Type	Method	Estimators	$n = 25$		$n = 75$		$n = 150$		$n = 200$	
			Bias	RMSE	Bias	RMSE	Bias	RMSE	Bias	RMSE
Non-Bayesian	MLE	$\hat{\alpha}$	0.11286	0.12090	0.05086	0.03474	0.02234	0.01469	0.01820	0.01103
		$\hat{\beta}$	0.46940	1.43545	0.15870	0.23078	0.06830	0.08480	0.05783	0.06505
	MPSE	$\hat{\alpha}$	0.09569	0.09438	0.04157	0.03060	0.03243	0.01434	0.02576	0.01085
		$\hat{\beta}$	0.16934	0.60087	0.10011	0.16233	0.08140	0.07392	0.06215	0.05774
	LS	$\hat{\alpha}$	0.00923	0.14895	0.01450	0.04730	0.00465	0.02174	0.00469	0.01615
		$\hat{\beta}$	0.22856	2.11097	0.07249	0.36435	0.02801	0.15136	0.02773	0.11719
	WLSE	$\hat{\alpha}$	0.02049	0.12790	0.02572	0.04022	0.01118	0.01774	0.01009	0.01310
		$\hat{\beta}$	0.22946	1.79682	0.09283	0.28474	0.04108	0.11422	0.03801	0.08614
	CvME	$\hat{\alpha}$	0.12041	0.19056	0.04951	0.05187	0.02191	0.02272	0.01759	0.01673
		$\hat{\beta}$	0.67634	4.06532	0.18677	0.44669	0.08217	0.16747	0.06804	0.12691
	ADE	$\hat{\alpha}$	0.03304	0.11210	0.02415	0.03710	0.00972	0.01733	0.00819	0.01266
		$\hat{\beta}$	0.22243	1.11151	0.08574	0.24304	0.03705	0.10827	0.03254	0.08102
	RTADE	$\hat{\alpha}$	0.06279	0.13519	0.03359	0.04351	0.01490	0.01984	0.01158	0.01437
		$\hat{\beta}$	0.30856	1.54896	0.10932	0.26772	0.04872	0.11165	0.03983	0.08178
	SEL	$\hat{\alpha}$	0.13323	0.09209	0.27787	0.10536	0.35740	0.14250	0.38120	0.15745
		$\hat{\beta}$	0.60523	1.00746	1.04966	1.36206	1.32139	1.90225	1.40639	2.11357
Bayesian	LINEX ₁	$\hat{\alpha}$	0.14608	0.09713	0.28449	0.10940	0.36156	0.14560	0.38455	0.16010
		$\hat{\beta}$	0.72662	1.31288	1.12359	1.55534	1.37314	2.05432	1.44973	2.24683
	LINEX ₂	$\hat{\alpha}$	0.12062	0.08758	0.27129	0.10144	0.35326	0.13946	0.37786	0.15484
		$\hat{\beta}$	0.49914	0.79733	0.98056	1.19749	1.27202	1.76405	1.36483	1.99065
	GEL ₁	$\hat{\alpha}$	0.12654	0.09017	0.27462	0.10350	0.35544	0.14108	0.37963	0.15624
		$\hat{\beta}$	0.57426	0.95349	1.03150	1.31985	1.30924	1.86857	1.39638	2.08424
	GEL ₂	$\hat{\alpha}$	0.11315	0.08667	0.26813	0.09986	0.35151	0.13826	0.37650	0.15383
		$\hat{\beta}$	0.51376	0.85718	0.99546	1.23880	1.28504	1.80263	1.37646	2.02656

Table 2. Bias and RMSE for Simulation Case II

Type	Method	Estimators	$n = 25$		$n = 75$		$n = 150$		$n = 200$	
			Bias	RMSE	Bias	RMSE	Bias	RMSE	Bias	RMSE
Non-Bayesian	MLE	$\hat{\alpha}$	0.11645	0.15185	0.03858	0.04096	0.02173	0.01783	0.01520	0.01450
		$\hat{\beta}$	0.09333	0.07373	0.03277	0.01547	0.01740	0.00687	0.01062	0.00493
	MPSE	$\hat{\alpha}$	0.12075	0.12465	0.06631	0.03997	0.04038	0.01797	0.03509	0.01482
		$\hat{\beta}$	0.05863	0.04330	0.03106	0.01263	0.01956	0.00617	0.01904	0.00473
	LSE	$\hat{\alpha}$	0.00313	0.19250	0.00039	0.05846	0.00688	0.02629	0.00173	0.02192
		$\hat{\beta}$	0.02538	0.13496	0.01058	0.02649	0.00004	0.01169	0.00369	0.00907
	WLSE	$\hat{\alpha}$	0.01128	0.16635	0.01268	0.04935	0.00482	0.02134	0.00740	0.01784
		$\hat{\beta}$	0.03085	0.10938	0.01736	0.02098	0.00691	0.00880	0.00645	0.00674
	CvME	$\hat{\alpha}$	0.12472	0.24354	0.04022	0.06297	0.01264	0.02704	0.01645	0.02259
		$\hat{\beta}$	0.13248	0.22792	0.04031	0.03086	0.01419	0.01248	0.01440	0.00961
	ADE	$\hat{\alpha}$	0.02417	0.14188	0.00998	0.04520	0.00170	0.02024	0.00517	0.01726
		$\hat{\beta}$	0.03293	0.06902	0.01567	0.01791	0.00509	0.00817	0.00512	0.00642
	RTADE	$\hat{\alpha}$	0.06763	0.16648	0.02390	0.05273	0.00929	0.02333	0.00782	0.01850
		$\hat{\beta}$	0.04862	0.06533	0.02104	0.01741	0.00789	0.00776	0.00579	0.00585
	SEL	$\hat{\alpha}$	0.35072	0.26583	0.45614	0.25243	0.49288	0.26456	0.50156	0.26892
		$\hat{\beta}$	0.30069	0.17086	0.35121	0.14268	0.37132	0.14711	0.37633	0.14884
Bayesian	LINEX ₁	$\hat{\alpha}$	0.37732	0.29101	0.46853	0.26480	0.49966	0.27152	0.50674	0.27429
		$\hat{\beta}$	0.31463	0.18364	0.35730	0.14746	0.37460	0.14966	0.37883	0.15079
	LINEX ₂	$\hat{\alpha}$	0.32496	0.24332	0.44391	0.24058	0.48614	0.25774	0.49640	0.26364
		$\hat{\beta}$	0.28714	0.15914	0.34517	0.13803	0.36807	0.14459	0.37384	0.14690
	GEL ₁	$\hat{\alpha}$	0.33980	0.25741	0.45116	0.24775	0.49018	0.26186	0.49950	0.26683
		$\hat{\beta}$	0.29228	0.16460	0.34748	0.13989	0.36932	0.14558	0.37480	0.14766
	GEL ₂	$\hat{\alpha}$	0.31801	0.24154	0.44121	0.23854	0.48477	0.25651	0.49537	0.26268
		$\hat{\beta}$	0.27552	0.15277	0.34001	0.13443	0.36532	0.14256	0.37175	0.14533

Table 3. Bias and RMSE for Simulation Case III

Type	Method	Estimators	$n = 25$		$n = 75$		$n = 150$		$n = 200$	
			Bias	RMSE	Bias	RMSE	Bias	RMSE	Bias	RMSE
Non-Bayesian	MLE	$\hat{\alpha}$	0.24776	0.27309	0.49295	0.31881	0.60192	0.40077	0.63135	0.43000
		$\hat{\beta}$	0.44987	0.49668	0.69957	0.58604	0.82513	0.73309	0.86033	0.78417
	MPSE	$\hat{\alpha}$	0.28484	0.30138	0.51214	0.33985	0.61335	0.41517	0.64039	0.44187
		$\hat{\beta}$	0.50143	0.58479	0.72862	0.63436	0.84334	0.76560	0.87497	0.81112
	LSE	$\hat{\alpha}$	0.21185	0.24929	0.47400	0.29890	0.59058	0.38677	0.62236	0.41839
		$\hat{\beta}$	0.40214	0.42617	0.67159	0.54196	0.80735	0.70217	0.84599	0.75830
	WLSE	$\hat{\alpha}$	0.23573	0.26643	0.48710	0.31284	0.59854	0.39663	0.62869	0.42660
		$\hat{\beta}$	0.43124	0.47328	0.68932	0.57039	0.81887	0.72235	0.85534	0.77527
	CvME	$\hat{\alpha}$	0.21168	0.25424	0.47537	0.30113	0.59175	0.38843	0.62337	0.41983
		$\hat{\beta}$	0.39455	0.43044	0.66893	0.54009	0.80637	0.70120	0.84539	0.75772
	ADE	$\hat{\alpha}$	0.17371	0.27966	0.05478	0.07627	0.03400	0.03270	0.01962	0.02868
		$\hat{\beta}$	0.25676	0.43877	0.07526	0.06944	0.03471	0.02799	0.02663	0.02358
	RTADE	$\hat{\alpha}$	0.15460	0.21989	0.08913	0.07400	0.05207	0.03239	0.04924	0.02932
		$\hat{\beta}$	0.11614	0.20789	0.07618	0.05516	0.05384	0.02587	0.04409	0.02225
	SEL	$\hat{\alpha}$	0.01855	0.36353	0.00645	0.10368	0.00003	0.04694	0.00902	0.04075
		$\hat{\beta}$	0.15115	1.22840	0.01672	0.10767	0.00099	0.04416	0.00180	0.03952
Bayesian	LINEX ₁	$\hat{\alpha}$	0.04328	0.33520	0.01205	0.08761	0.01360	0.03773	0.00227	0.03346
		$\hat{\beta}$	0.16318	1.03493	0.03219	0.08365	0.01248	0.03294	0.00908	0.02968
	LINEX ₂	$\hat{\alpha}$	0.19345	0.46699	0.04809	0.11121	0.02691	0.04882	0.01110	0.04152
		$\hat{\beta}$	0.41760	2.15216	0.08399	0.12860	0.03114	0.04781	0.02220	0.04177
	GEL ₁	$\hat{\alpha}$	0.05889	0.25941	0.01104	0.08131	0.01101	0.03671	0.00020	0.03310
		$\hat{\beta}$	0.13066	0.39240	0.03104	0.07461	0.01014	0.03155	0.00646	0.02865
	GEL ₂	$\hat{\alpha}$	0.10700	0.33262	0.02580	0.09380	0.02276	0.04403	0.00721	0.03582
		$\hat{\beta}$	0.17787	0.48808	0.04363	0.07865	0.02028	0.03404	0.01206	0.02768

Table 4. Bias and RMSE for Simulation Case IV

Type	Method	Estimators	$n = 25$		$n = 75$		$n = 150$		$n = 200$	
			Bias	RMSE	Bias	RMSE	Bias	RMSE	Bias	RMSE
Non-Bayesian	MLE	$\hat{\alpha}$	0.29255	0.27689	0.48993	0.30538	0.57208	0.35999	0.59317	0.37838
		$\hat{\beta}$	0.42497	0.39496	0.59463	0.41639	0.67430	0.48737	0.69532	0.51027
	MPSE	$\hat{\alpha}$	0.32596	0.30522	0.50691	0.32379	0.58194	0.37177	0.60088	0.38787
		$\hat{\beta}$	0.46136	0.44973	0.61389	0.44308	0.68577	0.50396	0.70438	0.52362
	LSE	$\hat{\alpha}$	0.26016	0.25238	0.47316	0.28788	0.56231	0.34852	0.58550	0.36908
		$\hat{\beta}$	0.39068	0.34885	0.57590	0.39152	0.66302	0.47138	0.68638	0.49731
	WLSE	$\hat{\alpha}$	0.28092	0.26937	0.48432	0.29969	0.56890	0.35630	0.59069	0.37540
		$\hat{\beta}$	0.40994	0.37776	0.58674	0.40624	0.66967	0.48090	0.69168	0.50506
	CvME	$\hat{\alpha}$	0.25770	0.25539	0.47308	0.28852	0.56252	0.34897	0.58573	0.36948
		$\hat{\beta}$	0.38022	0.34587	0.57102	0.38650	0.66044	0.46815	0.68442	0.49475
	ADE	$\hat{\alpha}$	0.17839	0.26294	0.05276	0.06468	0.01949	0.02742	0.01732	0.02247
		$\hat{\beta}$	0.20538	0.27381	0.05651	0.04425	0.01949	0.01917	0.01848	0.01336
	RTADE	$\hat{\alpha}$	0.12053	0.19621	0.07833	0.06193	0.05816	0.02888	0.04529	0.02315
		$\hat{\beta}$	0.08812	0.13443	0.06317	0.03628	0.04989	0.01892	0.03744	0.01304
	SEL	$\hat{\alpha}$	0.01613	0.31099	0.00003	0.09083	0.00585	0.04275	0.00339	0.03528
		$\hat{\beta}$	0.08035	0.49334	0.01767	0.07823	0.00030	0.03501	0.00564	0.02602
Bayesian	LINEX ₁	$\hat{\alpha}$	0.04480	0.28370	0.01693	0.07603	0.00436	0.03381	0.00954	0.02830
		$\hat{\beta}$	0.10237	0.45374	0.02829	0.05853	0.00756	0.02581	0.01076	0.01879
	LINEX ₂	$\hat{\alpha}$	0.17312	0.39320	0.04975	0.09797	0.01868	0.04412	0.02179	0.03637
		$\hat{\beta}$	0.27708	0.84317	0.07200	0.09307	0.02625	0.03775	0.02505	0.02775
	GEL ₁	$\hat{\alpha}$	0.05700	0.23637	0.01476	0.07063	0.00069	0.03260	0.00624	0.02712
		$\hat{\beta}$	0.08586	0.22001	0.02501	0.05055	0.00426	0.02412	0.00760	0.01721
	GEL ₂	$\hat{\alpha}$	0.10345	0.28786	0.02187	0.07676	0.00433	0.03591	0.01465	0.03060
		$\hat{\beta}$	0.12193	0.27038	0.02848	0.04800	0.00574	0.02213	0.01356	0.01721

Tables (1-4) examine the performance of various non-Bayesian and Bayesian estimators for parameters α and β across four distinct simulation scenarios (Cases I, II, III, and IV), evaluated by their Bias and Root Mean Squared Error (RMSE) at increasing sample sizes ($n = 25, 75, 150, 200$).

A fundamental observation across all simulation cases is the expected improvement in estimator performance with increasing sample size. For most methods, both bias and RMSE consistently decrease as n grows, indicating that larger datasets lead to more accurate and precise parameter estimates. Non-Bayesian methods, including MLE, MPSE, LS, WLSE, CvME, ADE, and RTADE, exhibit varying degrees of effectiveness depending on the simulation case. For Cases I and II, methods like ADE, MPSE, and WLSE frequently demonstrate superior performance, characterized by generally lower bias and RMSE for both $\hat{\alpha}$ and $\hat{\beta}$, especially as sample sizes increase. LSE also shows notably low bias for $\hat{\alpha}$ in Case II. This suggests these methods are well-suited under the conditions of Simulation Cases I and II. In contrast, for Cases III and IV, non-Bayesian estimators generally show higher initial bias and RMSE compared to the first two cases. While their performance still improves with increasing sample size, the overall magnitudes of error are greater. Within these cases, ADE and RTADE often emerge as relatively better performers, demonstrating smaller biases and RMSEs at larger sample sizes compared to other non-Bayesian counterparts. Bayesian estimators (SEL, LINEX₁, LINEX₂, GEL₁, GEL₂) show a more nuanced and sometimes counter-intuitive behavior. A significant and unexpected finding in Cases I and II is the tendency for the bias of Bayesian estimators (for both $\hat{\alpha}$ and $\hat{\beta}$) to increase with increasing sample size. This runs contrary to the fundamental

asymptotic properties of good estimators and warrants further investigation into the prior specifications or the specific characteristics of these simulation cases that lead to such behavior. Despite this increasing bias, for Case II, the RMSE for Bayesian estimators generally still decreases with n , implying that their overall error (which includes variance) is still reducing. In a stark contrast to Cases I and II, the Bayesian methods demonstrate robust and often superior performance in Cases III and IV. Here, both the bias and RMSE for $\hat{\alpha}$ and $\hat{\beta}$ consistently decrease as the sample size increases, as expected. For larger sample sizes, several Bayesian estimators, particularly those under SEL, LINEX₁, and GEL₁ loss functions, achieve remarkably low bias and RMSE values, frequently outperforming their non-Bayesian counterparts. This indicates that under the conditions simulated in Cases III and IV, the Bayesian approach, especially with appropriate loss functions, provides highly efficient and accurate estimates. The simulation results highlight that the optimal choice of estimator is highly dependent on the underlying data generation process (represented by the "Simulation Cases"). While increasing sample size generally leads to better estimates, the study reveals a critical divergence in Bayesian estimator performance. Bayesian methods exhibit unexpected increasing bias in Cases I and II, suggesting potential issues in these scenarios, but demonstrate strong performance with decreasing bias and RMSE in Cases III and IV, often surpassing non-Bayesian methods. This underscores the importance of context-specific evaluation and potentially careful prior elicitation when employing Bayesian estimation.

7. Application

The first applications is one mortality rate due to HIV/AIDS in Germany from year 2000 to 2020. The data is contained in <https://platform.who.int/mortality/themes/theme-details/topics/indicator-groups/indicator-group-details/MDB/hiv-aids> and presented in Table 5.

Table 5. Mortality rate due to HIV/AIDS in Germany

0.70570244	0.65946256	0.62801346	0.6143952	0.61453596	0.59540885	0.61190438
0.56040019	0.53945593	0.52641369	0.55652406	0.56615856	0.50050448	0.49723726
0.47911586	0.3270769	0.34298934	0.35461942	0.37625366	0.41652161	0.45295061

The second application is on the 60 days (1 December 2020 to 29 January 2021) COVID-19 mortality rate data of the United Kingdom, extracted from <https://covid19.who.int/>. This data has been investigated by [6]. The data is presented in Table 6 below.

Table 6. 60 days COVID-19 mortality rate data of the United Kingdom

0.1292	0.3805	0.4049	0.2564	0.3091	0.2413	0.1390	0.1127
0.3547	0.3126	0.2991	0.2428	0.2942	0.0807	0.1285	0.2775
0.3311	0.2825	0.2559	0.2756	0.1652	0.1072	0.3383	0.3575
0.2708	0.2649	0.0961	0.1565	0.1580	0.1981	0.4154	0.3990
0.2483	0.1762	0.1760	0.1543	0.3238	0.3771	0.4132	0.4602
0.3523	0.1882	0.1742	0.4033	0.4999	0.3930		

Table 7. Summary Statistical Measures

Statistics	Mortality rate due to HIV/AIDS in Germany	60 days COVID-19 mortality rate of the United Kingdom
n	21	46
Q_1	0.45295	0.17465
Q_3	0.61190	0.35410
Mean	0.52027	0.26903
Median	0.53946	0.27320
Std. Dev.	0.10896	0.10803
Variance	0.01187	0.01167
Range	0.37863	0.41920
IQR	0.15895	0.17945
Skewness	-0.33351	0.06229
Kurtosis	2.11013	2.02746

Table 7 provides a summary of key statistical measures for two distinct mortality datasets: "Mortality rate due to HIV/AIDS in Germany" and "60 days COVID-19 mortality rate of the United Kingdom." These measures offer insights into the central tendency, dispersion, and shape of each dataset.

For the HIV/AIDS mortality rate in Germany, the sample size (n) is 21. The mean mortality rate is approximately 0.520, with a median of 0.539, suggesting a slight left skew or that half the observations are above 0.539. The first quartile (Q_1) is 0.453 and the third quartile (Q_3) is 0.612, indicating the middle 50% of data falls within this range, with an Interquartile Range (IQR) of 0.159. The standard deviation is about 0.109, and the variance is 0.012, showing a relatively small spread. The range of 0.379 further supports this limited variability. A negative skewness of -0.334 suggests a slight left-skewed distribution, meaning the tail on the left side of the distribution is longer or fatter. The kurtosis of 2.110 indicates a platykurtic distribution, meaning it has lighter tails and fewer extreme values compared to a normal distribution (which has a kurtosis of 3).

In contrast, the 60-day COVID-19 mortality rate in the United Kingdom has a larger sample size (n) of 46. The mean mortality rate is approximately 0.269, with a very similar median of 0.273. This suggests a more symmetrical distribution around the center compared to the HIV/AIDS data. The quartiles are $Q_1 = 0.175$ and $Q_3 = 0.354$, with an IQR of 0.179, indicating a spread of the central 50% of data. The standard deviation (0.108) and variance (0.012) are remarkably similar to the HIV/AIDS data, implying similar absolute variability. However, given the lower mean, this suggests a higher relative variability for the COVID-19 data. The range of 0.419 is slightly larger than that for the HIV/AIDS data. A positive skewness of 0.062 indicates a very slight right-skew, close to symmetrical. The kurtosis of 2.027, also platykurtic, suggests lighter tails than a normal distribution, similar to the HIV/AIDS data.

In essence, while both data sets exhibit comparable levels of absolute variability, as shown by their standard deviations, the mortality rate of COVID-19 in the UK generally presents lower average rates and a distribution that is more symmetric around its mean compared to the mortality rate of HIV / AIDS in Germany, which shows a higher average and a slight left skew.

Figures 7 and 8, respectively, demonstrate the characteristics of the two datasets. The images in order of appearance are violin/boxplot, histogram/density plot, empirical/theoretical CDF, empirical/theoretical

survival function plot, hazard rate plot, TTT plot, probability-probability (P-P) plot, quantile-quantile (Q-Q) plot, contour plot of the likelihood profile.

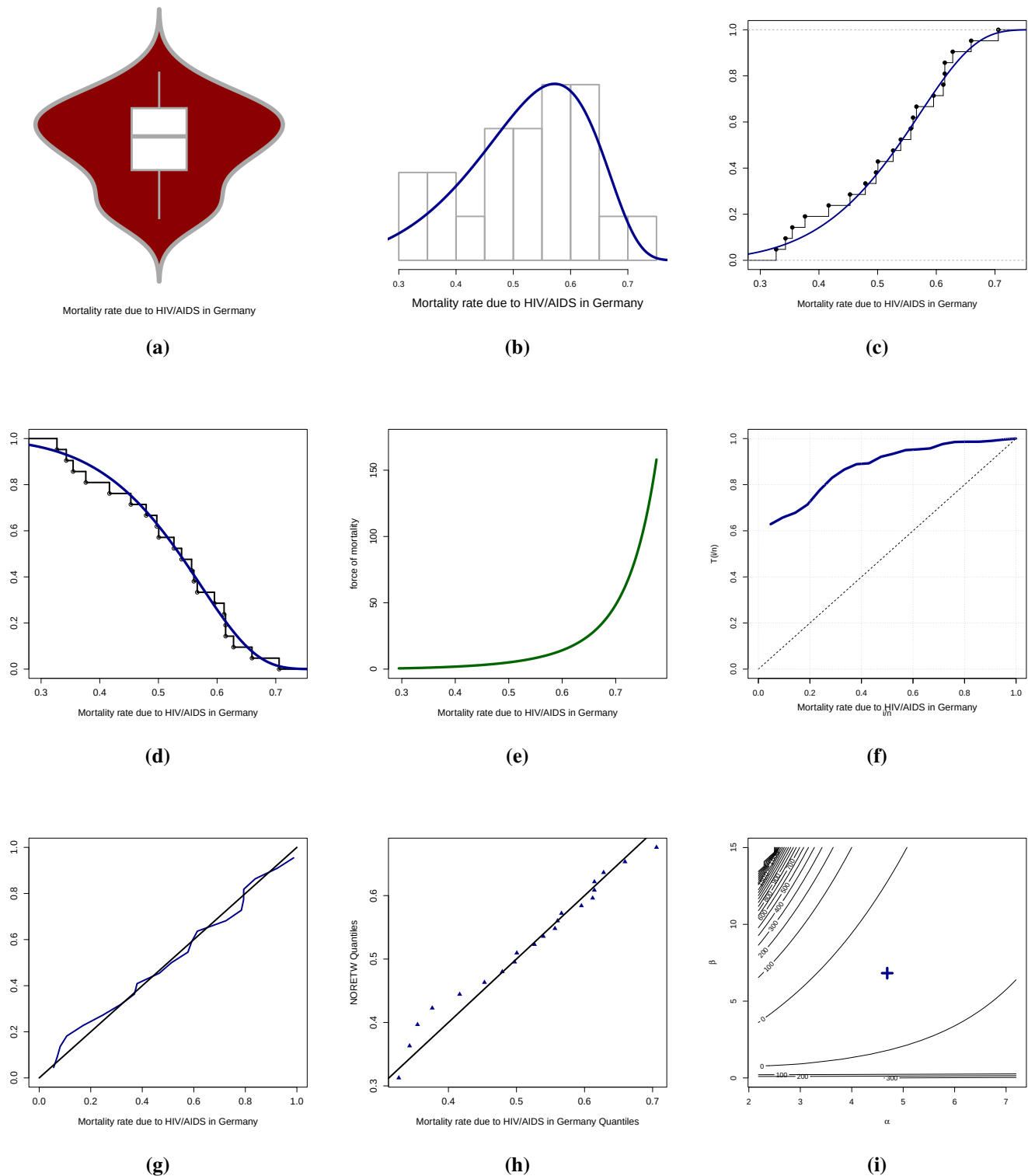


Figure 7. (a) Violin/boxplot, (b) histogram/density plot, (c) Empirical/theoretical CDF, (d) Empirical/theoretical survival function, (e) hazard function, (f) TTT plot, (g) P-P plot, (h) Q-Q plot, (i) contour plot for mortality rate due to HIV/AIDS in Germany.

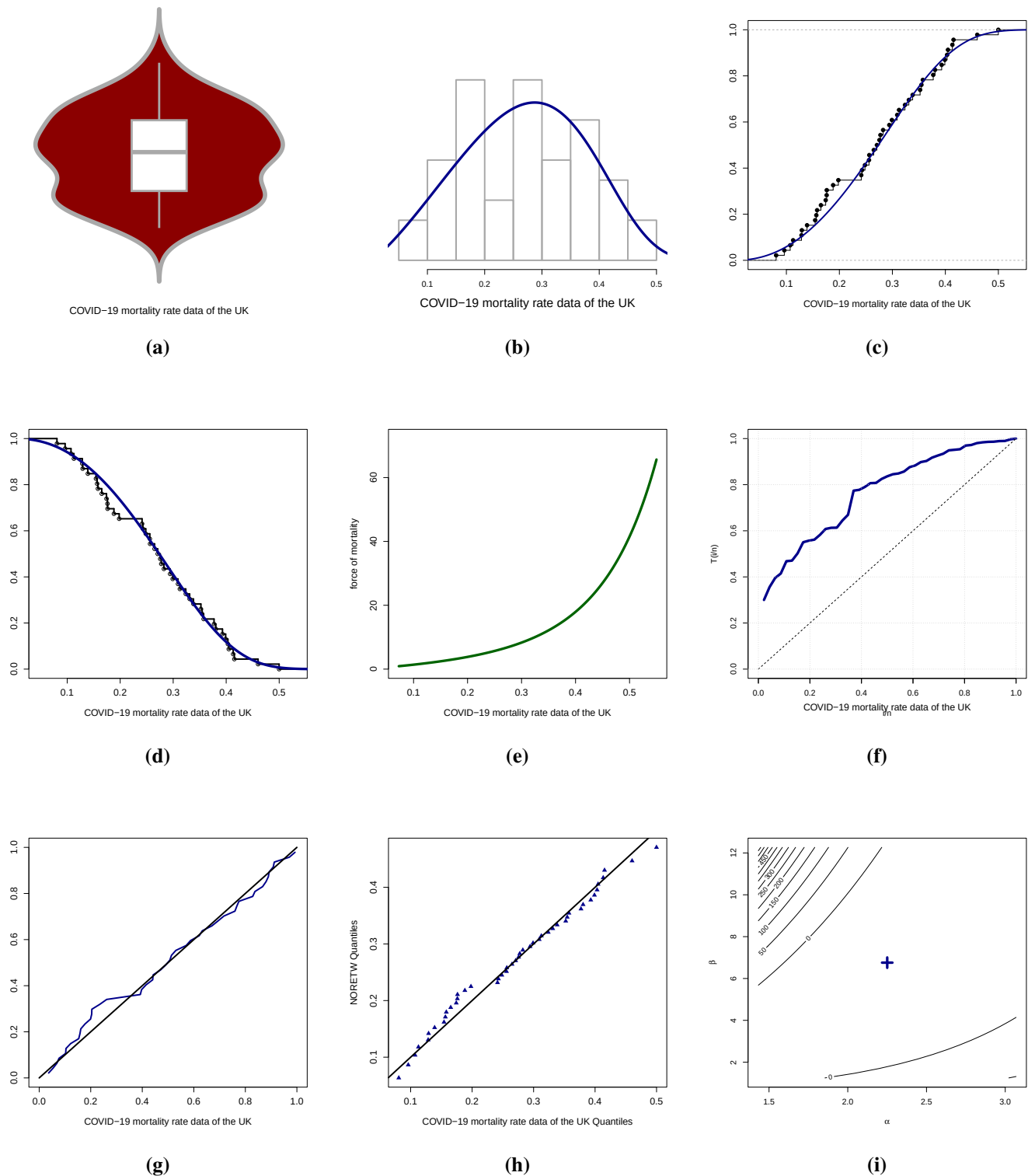


Figure 8. (a) Violin/boxplot, (b) histogram/density plot, (c) Empirical/theoretical CDF, (d) Empirical/theoretical survival function, (e) hazard function, (f) TTT plot, (g) P-P plot, (h) Q-Q plot, (i) contour plot for COVID-19 mortality rate data of the United Kingdom.

Table 8. Model Performance, Goodness of Fit and MLE values

Data	Distribution	LL	AIC	CAIC	BIC	HQIC	W	A	KS	p-value	shape _{MLE} (std. err)	scale _{MLE} (std. err)
HIV/AIDS	NORETW	17.98	-31.9615	-31.2949	-29.8725	-31.5082	0.0218	0.1943	0.0834	0.9959	4.6903(0.8351)	6.8244(2.7316)
	Gumbel	15.4	-26.8009	-26.1342	-24.7118	-26.3475	0.1297	0.7978	0.1441	0.7232	0.4656(0.0243)	0.1049(0.0171)
	TIHTR	10.76	-17.5284	-16.8618	-15.4394	-17.0751	0.0835	0.5403	0.1974	0.3405	0.1875(0.1154)	22.5224(12.5954)
	TIHTE	-2.98	6.6674	7.3341	8.7564	7.1208	0.0913	0.5836	0.3430	0.0104	0.2174(0.058)	10.2038(3.9753)
	Gamma	16.69	-29.3856	-28.7189	-27.2965	-28.9322	0.0826	0.5336	0.1280	0.8391	21.9803(6.7324)	42.2483(13.0890)
	Weibull	17.79	-31.5739	-30.9073	-29.4849	-31.1206	0.0326	0.2473	0.1000	0.9710	0.5630(0.0221)	5.8439(1.0281)
COVID-19	NORETW	39.06	-74.1238	-73.8448	-70.4666	-72.7538	0.0447	0.3157	0.1001	0.7086	2.2506(0.2729)	6.7581(1.8374)
	Gumbel	36.33	-68.6585	-68.3794	-65.0012	-67.2884	0.1381	0.7932	0.1154	0.5350	0.2161(0.0151)	0.0964(0.0110)
	TIHTR	35.25	-70.6186	-70.3395	-66.9613	-69.2486	0.0829	0.5219	0.0917	0.8004	0.4183(0.1967)	36.0350(18.7924)
	TIHTE	16.76	-45.6314	-45.3524	-41.9741	-44.2614	0.1437	0.8146	0.1843	0.0770	0.1847(0.0759)	23.9902(8.9099)
	Gamma	37.35	-70.7007	-70.4216	-67.0434	-69.3307	0.1167	0.6703	0.1093	0.6030	5.5273(1.1195)	20.5445(4.3559)
	Weibull	38.62	-73.2340	-72.9550	-69.5768	-71.8640	0.0651	0.4116	0.1061	0.6400	0.3030(0.0169)	2.7845(0.3306)

Table 8 presents a concise evaluation of various statistical distributions, specifically NORETW, Gumbel, TIHTR, TIHTE, Gamma, and Weibull, in modeling two distinct datasets: HIV/AIDS and COVID-19. The assessment is based on model fit (Log-Likelihood), model selection criteria (AIC, CAIC, BIC, HQIC), and goodness-of-fit tests (W, A, KS statistics with p-values). Additionally, the Maximum Likelihood Estimates (MLEs) for the shape and scale parameters, along with their standard errors, are provided for each distribution.

For the HIV/AIDS data, the NORETW distribution consistently demonstrates the strongest performance. It boasts the highest Log-Likelihood (17.98) and the most favorable (lowest negative) values across all information criteria, indicating superior fit and parsimony. Critically, its exceptionally high p-value (0.9959) from the goodness-of-fit tests confirms an excellent fit to the data, suggesting it's highly suitable for this dataset. The Weibull distribution also shows strong performance for HIV/AIDS data, closely following NORETW across the metrics. Conversely, the TIHTE distribution exhibits a poor fit, characterized by a very low Log-Likelihood and a p-value (0.0104) that indicates a statistically significant departure from the observed data.

Similarly, for the COVID-19 data, the NORETW distribution again emerges as the top performer. It yields the highest Log-Likelihood (39.06) and the lowest information criteria values. Its strong p-value (0.7086) from the goodness-of-fit tests further validates its robust fit to the COVID-19 data. The Weibull distribution also proves to be a strong contender for this dataset. The TIHTE distribution, once more, indicates the poorest fit among the evaluated models, although its performance is slightly better than its fit to the HIV/AIDS data.

In summary, the NORETW distribution consistently provides the best statistical fit for both the HIV/AIDS and COVID-19 datasets, making it the most suitable model among those compared, closely followed by the Weibull distribution.

8. Concluding Remarks

This study successfully introduced the novel odd reparameterized exponential transformed-X (NORET-X) family of distributions, utilizing an exponential transformer and an odd function generalizer. The efficacy of this new framework was thoroughly demonstrated through its application to extend the classical Weibull distribution, yielding the new odd reparameterized exponential transformed-Weibull (NORETW) distribution. A comprehensive characterization of the NORETW distribution was undertaken, encompassing key

statistical properties such as quantile function, moments, moment generating function, mean residual life, order statistics, entropy, and extropy. Parameter estimation for this new distribution was robustly explored using both non-Bayesian and Bayesian methodologies, supported by an extensive Monte Carlo simulation across four diverse scenarios, the results of which highlighted varying performance characteristics across methods and parameters, with sample size generally improving estimation quality. Crucially, when applied to real-world datasets of HIV/AIDS and COVID-19 mortality rates, the NORETW distribution consistently demonstrated a superior fit compared to the baseline Weibull distribution and other related models. This superior performance is evidenced by its commanding probability values of 0.9959 and 0.7086 for the HIV/AIDS and COVID-19 datasets, respectively, indicating excellent goodness-of-fit. Furthermore, the NORETW model exhibited the lowest AIC, CAIC, BIC, and HQIC values in parameter estimation, unequivocally establishing its best performance among the compared models. These findings strongly advocate for the NORETW distribution as a highly promising and flexible model for analyzing complex lifetime data across various fields.

Conflict of Interest

The authors declare there is no existing conflict of interest.

References

1. Al-Marzouki, S., Jamal, F., Chesneau, C., and Elgarhy, M. (2019). Type ii topp leone power lomax distribution with applications. *Mathematics*, 8(1):4.
2. Al-Shomrani, A., Arif, O., Shawky, A., Hanif, S., and Shahbaz, M. Q. (2016). Toppâ€leone family of distributions: Some properties and application. *Pakistan Journal of Statistics and Operation Research*, pages 443–451.
3. Alizadeh, M., Bagheri, S., Samani, E. B., Ghobadi, S., and Nadarajah, S. (2018). Exponentiated power lindley power series class of distributions: Theory and applications. *Communications in Statistics-Simulation and Computation*, 47(9):2499–2531.
4. Alizadeh, M., Cordeiro, G. M., Pinho, L. G. B., and Ghosh, I. (2017). The gompertz-g family of distributions. *Journal of statistical theory and practice*, 11:179–207.
5. Almarashi, A. M., Badr, M. M., Elgarhy, M., Jamal, F., and Chesneau, C. (2020). Statistical inference of the half-logistic inverse rayleigh distribution. *Entropy*, 22(4):449.
6. Almetwally, E. M. (2022). The odd weibull inverse toppâ€leone distribution with applications to covid-19 data. *Annals of Data Science*, 9(1):121–140.
7. Alyami, S. A., Elbatal, I., Alotaibi, N., Almetwally, E. M., Okasha, H. M., and Elgarhy, M. (2022). Toppâ€leone modified weibull model: Theory and applications to medical and engineering data. *Applied Sciences*, 12(20):10431.
8. Alzaatreh, A., Lee, C., and Famoye, F. (2013). A new method for generating families of continuous distributions. *Metron*, 71(1):63–79.
9. Alzaatreh, A., Lee, C., and Famoye, F. (2014). T-normal family of distributions: a new approach to generalize the normal distribution. *Journal of Statistical Distributions and Applications*, 1:1–18.

10. Anyiam, K. E., Alghamdi, F. M., Nwaigwe, C. C., Aljohani, H. M., and Obulezi, O. J. (2024). A new extension of burr-hatke exponential distribution with engineering and biomedical applications. *Heliyon*, 10(19).
11. Anzagra, L., Sarpong, S., and Nasiru, S. (2022). Odd chen-g family of distributions. *Annals of Data Science*, 9(2):369–391.
12. Bourguignon, M., Silva, R. B., Zea, L. M., and Cordeiro, G. M. (2013). The kumaraswamy pareto distribution. *Journal of statistical theory and applications*, 12(2):129–144.
13. Brooks, S. (1998). Markov chain monte carlo method and its application. *Journal of the royal statistical society: series D (the Statistician)*, 47(1):69–100.
14. Calabria, R. and Pulcini, G. (1996). Point estimation under asymmetric loss functions for left-truncated exponential samples. *Communications in Statistics-Theory and Methods*, 25(3):585–600.
15. Cheng, R. and Amin, N. (1979). Maximum product of spacings estimation with application to the lognormal distribution (mathematical report 79-1). *Cardiff: University of Wales IST*.
16. Cordeiro, G. M., Alizadeh, M., and Diniz Marinho, P. R. (2016). The type i half-logistic family of distributions. *Journal of Statistical Computation and Simulation*, 86(4):707–728.
17. Cordeiro, G. M. and De Castro, M. (2011). A new family of generalized distributions. *Journal of statistical computation and simulation*, 81(7):883–898.
18. Doostparast, M., Akbari, M. G., and Balakrishna, N. (2011). Bayesian analysis for the two-parameter pareto distribution based on record values and times. *Journal of Statistical Computation and Simulation*, 81(11):1393–1403.
19. Ekemezie, D.-F. N., Alghamdi, F. M., Aljohani, H. M., Riad, F. H., Abd El-Raouf, M., and Obulezi, O. J. (2024). A more flexible lomax distribution: Characterization, estimation, group acceptance sampling plan and applications. *Alexandria Engineering Journal*, 109:520–531.
20. Elbatal, I., Alotaibi, N., Almetwally, E. M., Alyami, S. A., and Elgarhy, M. (2022). On odd perks-g class of distributions: properties, regression model, discretization, bayesian and non-bayesian estimation, and applications. *Symmetry*, 14(5):883.
21. Elbatal, I. and Elgarhy, M. (2013). Statistical properties of kumaraswamy quasi lindley distribution. *International Journal of Mathematics Trends and Technology-IJMTT*, 4.
22. Gomes-Silva, F. S., Percontini, A., de Brito, E., Ramos, M. W., Venâncio, R., and Cordeiro, G. M. (2017). The odd lindley-g family of distributions. *Austrian journal of statistics*, 46(1):65–87.
23. Hassan, A. S. and Nassr, S. G. (2019). Power lindley-g family of distributions. *Annals of Data Science*, 6:189–210.
24. Jafari, A. A. and Tahmasebi, S. (2016). Gompertz-power series distributions. *Communications in Statistics-Theory and Methods*, 45(13):3761–3781.
25. Jamal, F., Chesneau, C., and Elgarhy, M. (2020). Type ii general inverse exponential family of distributions. *Journal of Statistics and Management Systems*, 23(3):617–641.
26. Lad, F., Sanfilippo, G., and Agro, G. (2015). Extropy: Complementary dual of entropy. *Statistical Science*, 30(1):40–58.

27. Lehmann, E. L. (2011). Ordered families of distributions. In *Selected Works of EL Lehmann*, pages 757–777. Springer.
28. Mansoor, M., Cordeiro, G. M., and Zubair, M. (2020). A study of the logistic exponentiated-exponential distribution and its applications. *Journal of Advances in Applied & Computational Mathematics*, 7:38–48.
29. Maurya, S. K. and Nadarajah, S. (2021). Poisson generated family of distributions: a review. *Sankhya B*, 83:484–540.
30. Oramulu, D. O., Alsadat, N., Kumar, A., Bahloul, M. M., and Obulezi, O. J. (2024). Sine generalized family of distributions: Properties, estimation, simulations and applications. *Alexandria Engineering Journal*, 109:532–552.
31. Rényi, A. (1961). On measures of entropy and information. In *Proceedings of the fourth Berkeley symposium on mathematical statistics and probability, volume 1: contributions to the theory of statistics*, volume 4, pages 547–562. University of California Press.
32. Swain, J. J., Venkatraman, S., and Wilson, J. R. (1988). Least-squares estimation of distribution functions in johnson's translation system. *Journal of Statistical Computation and Simulation*, 29(4):271–297.
33. Tahir, M., 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.
34. Tolba, A. H., Onyekwere, C. K., El-Saeed, A. R., Alsadat, N., Alohal, H., and Obulezi, O. J. (2023). A new distribution for modeling data with increasing hazard rate: A case of covid-19 pandemic and vinyl chloride data. *Sustainability*, 15(17):12782.
35. Van Ravenzwaaij, D., Cassey, P., and Brown, S. D. (2018). A simple introduction to markov chain monte-carlo sampling. *Psychonomic bulletin & review*, 25(1):143–154.
36. Varian, H. R. (1975). A bayesian approach to real estate assessment. *Studies in Bayesian econometrics and statistics in honor of Leonard J. Savage*.
37. Weibull, W. (1939). A statistical theory of strength of materials. *IVB-Handl*.
38. Yousof, H. M., Altun, E., Ramires, T. G., Alizadeh, M., and Rasekhi, M. (2018). A new family of distributions with properties, regression models and applications. *Journal of Statistics and Management Systems*, 21(1):163–188.



© 2025 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 Sphinx Scientific Press (SSP) or the editor(s). SSP 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.