

Bilal-G Family of Distributions with Applications to Biomedical and Reliability Engineering Data

Joy I. Udobi¹, Happiness O. Obiora-Ilouno², Okechukwu J. Obulezi^{2,*}, Doaa M.H. Ahmed³, Ehab M. Almetwally⁴ and Mohammed Elgarhy^{5,6}

¹Department of Statistics, School of Applied Sciences, Federal Polytechnic, Oke, Nigeria

²Department of Statistics, Faculty of Physical Sciences, Nnamdi Azikiwe University, P.O. Box 5025 Awka, Nigeria

³Department of Insurance and Risk Management, College of Business, Imam Mohammad Ibn Saud Islamic University (IMSIU), Riyadh, 11432, Saudi Arabia

⁴Department of Mathematics and Statistics, College of Science, Imam Mohammad Ibn Saud Islamic University, Riyadh, Saudi Arabia

⁵Department of Basic Sciences, Higher Institute of Administrative Sciences, Belbeis, AlSharkia, Egypt

⁶Department of Computer Engineering, Biruni University, 34010, Istanbul, Turkey

Abstract: This paper introduces the Bilal-G (B-G) family of distributions, a novel generator-based method for enhancing the flexibility of existing probability models to better accommodate complex data structures prevalent in biomedical and reliability engineering. Data from these fields frequently exhibit features like high skewness, significant outliers, and non-monotone hazard rates that challenge conventional distributions. Using the Bilal distribution as the generator, we construct the new family's general cumulative distribution function (CDF) and probability density function (PDF), from which a key, parsimonious sub-model, the two-parameter Bilal-Exponential (BE) distribution, is derived. We thoroughly analyze the BE distribution's properties, including its capability to model an increasing hazard rate, which is supported by Total Time on Test (TTT) plots of the application datasets. A comprehensive simulation study evaluates the performance of fifteen distinct non-Bayesian estimators, revealing that the Minimum Spacing Linex Distance (MSLNDE) method consistently provides the most accurate and precise parameter estimates across various sample sizes. Finally, the superiority of the BE distribution is demonstrated through its successful application to two real datasets: one on patient mortality rates and one on component failure times. For the mortality data (Data I), the BE model reduced the Akaike Information Criterion (AIC) by 1.99 units compared to the classical Weibull distribution. For the component failure data (Data II), the Bayesian Information Criterion (BIC) was reduced by 0.41 units compared to the best-fitting competing model (TIHTE), confirming the BE distribution's exceptional goodness-of-fit and reliability as a practical lifetime model.

Keywords: Bilal distribution, Bilal-G family of distributions, estimation, biomedical studies, engineering reliability data.

1. INTRODUCTION

In fields like reliability engineering and survival analysis, researchers frequently encounter complex empirical data that traditional statistical models (e.g., Normal, Exponential, Weibull) fail to adequately characterize. Data sets in these areas often exhibit crucial characteristics for accurate prediction, such as strong skewness, the presence of significant outliers, and critically, non-monotone hazard rates (e.g., bathtub or unimodal shapes) which capture various life cycles of a system or organism. This necessity—driven by the inability of traditional models to simultaneously capture features like skewness, kurtosis, multimodality, or varied tail behavior—is the primary factor spurring the creation of various methodologies for constructing more flexible models. Early conceptual work includes

[1]'s system of normalization transformations, which laid the groundwork for defining generalized distribution families. Other foundational techniques include using generalized distributions derived from concepts like the hazard rate [2], the Exponentiation Method [3], and the [4] technique, which adds a single parameter to an existing distribution's cumulative distribution function (CDF) to improve its goodness-of-fit. These foundational efforts established that the addition of parameters or the transformation of the baseline distributions could significantly enhance the versatility of the modeling.

Modern research has expanded these ideas through sophisticated generator-based families. A major innovation was the Beta-Generated (Beta-G) family [5], which uses the Beta distribution as a generator to transform a baseline distribution's CDF, leading to highly flexible models like the Beta-Normal distribution. Leveraging the advantages of simpler

*Address correspondence to this author at the Department of Statistics, Faculty of Physical Sciences, Nnamdi Azikiwe University, P.O. Box 5025 Awka, Nigeria; E-mail: oj.obulezi@unizik.edu.ng

forms, the Kumaraswamy-G (Kw-G) distribution [6] was subsequently introduced, offering similar shape properties to the Beta-G family with simpler explicit formulas and moment calculations. To address limitations in the Beta-G family's support range, the general Transformed-Transformer (T-X) method [7] was introduced, which allows any non-negative continuous random variable to serve as the generator. These generator methods—and their subsequent variants based on distributions like Weibull [8], Lomax [9], and Lindley [10]—consistently demonstrate greater flexibility, allowing researchers to model diverse data shapes and accommodate various hazard rates (e.g., constant, increasing, decreasing, bathtub, unimodal), often outperforming traditional models in statistical fit tests. Recently, families of distributions that capture these seemingly data behaviors have emerged in the literature. They include, but are not limited to, a new family of generalized distributions using transformed Lomax-x [11], Weibull sine generalized family [12], Kavya-Manoharan DUS family [13], Alpha sine power transformation [14], a new odd reparameterized exponential transformed-x family [15], type-II heavy-tailed family [16], a tangent DUS family [17], sine generalized family [18], exponential arctan family [19], arcsine Topp-Leone family [20], arctan-x family [21]. For more details see [22-27].

Due to the inherent characteristics of emerging datasets from various human endeavors, especially biomedical and engineering studies, a distribution that fits well with a particular dataset may not fit well with related datasets. On the above premise, the primary motivation and key contributions (novelty) of this study are encapsulated in the following points:

1. The proposed family of distributions exhibits desirable tractability, as all derived statistical properties and measures possess analytical, simple, and often linear solutions.
2. Adhering to the principle of parsimony in statistical inference, the Bilal-G family introduces only a single, non-redundant shape parameter, resulting in new distributions that maintain simplicity and avoid parameter superfluity.
3. The derived Bilal-Exponential submodel demonstrates effectiveness and robustness in modeling data characterized by significant skewness and the presence of outliers, suggesting that any distribution from the broader Bilal-G family may offer similarly robust distributional forms.
4. The Bilal-Exponential distribution consistently achieved superior goodness-of-fit, evidenced by its high probability values in statistical tests and minimum values across key model performance criteria when fitted to the two empirical datasets.

Traditional probability distributions, such as the Exponential, Normal, and Weibull models, often fall short when analyzing complex, emerging datasets in modern biomedical and reliability engineering. These models fail to adequately capture the inherent characteristics of real-world data, which frequently exhibit:

- **Extreme Skewness and Outliers:** Life-testing and survival data, such as patient survival times or component failure times, are commonly characterized by a high degree of asymmetry and the presence of significant, influential outliers that distort traditional parameter estimation.
- **Non-Monotone Failure Rates:** A critical requirement for accurate lifetime modeling is the ability to handle varying patterns of risk over time. Systems and organisms often pass through stages (infancy, useful life, wear-out), resulting in bathtub-shaped or unimodal (upside-down bathtub) hazard functions. Models that only allow for simple increasing or decreasing failure rates are practically inadequate for these scenarios.

To address these critical data challenges, this paper introduces the Bilal-G family of distributions, a novel methodology designed to enhance the flexibility and descriptive power of existing probability models. Our core contributions, specifically tailored to practical application in these fields, are as follows:

- **Modeling Complex Risk Patterns:** We demonstrate that the derived Bilal-Exponential (BE) submodel is capable of capturing an increasing failure rate, making it a versatile tool for modeling the wear-out phase in reliability and the later stages of mortality in biomedical studies.
- **Robustness to Real-World Noise:** The BE distribution demonstrates effectiveness and robustness in modeling data characterized by high skewness and the presence of outliers, providing a more reliable fit than competing models in non-ideal conditions.

- Demonstrated Superior Goodness-of-Fit: Through detailed application to two real-world datasets—one concerning patient mortality rates (biomedical) and another involving fixed component failure times (reliability)—the BE distribution is shown to consistently achieve a superior goodness-of-fit over several established lifetime models. This confirmed practical utility establishes the Bilal-G family as a powerful new tool for engineers and biostatisticians.

The study workflow is summarized in the following schematic diagram:

Introduction → Bilal – G Family →

BED istribution → Statistical Properties →

Estimation → Simulation → Applications → Concluding Remarks.

2. THE BILAL-G FAMILY

In this section, the Bilal-G family of distributions, with a special submodel, will be constructed and its properties derived and studied. Subsequently, the parameters of the submodel will be estimated under the complete sample using some non-Bayesian and Bayesian estimators.

Definition 2.1 Suppose a differentiable and monotone right-continuous function $W\{F(x)\}$ is defined as $W(F(x)) = -\log(1 - F(x))$. It is easy to see that $\frac{dW\{F(x)\}}{dx} = \frac{f(x)}{1 - F(x)}$.

Let T : the Bilal (θ) distribution, an innovative one-parameter longevity model introduced by [28] with probability density function (PDF) given as

$$r(t) = \frac{6}{\theta} e^{-2t/\theta} (1 - e^{-t/\theta}); \quad t > 0, \quad \theta > 0. \quad (2.1)$$

The associated cumulative distribution function (CDF) can be expressed as;

$$R(t) = 1 - (3 - 2e^{-t/\theta})e^{-2t/\theta}. \quad (2.2)$$

Then, the CDF of the Bilal-G family of distribution utilizing the $T - X$ generator proposed by [7] is given as

$$\begin{aligned} G(x) &= \int_0^{-\ln(1-F(x))} r(t) dt = R\{-\ln(1-F(x))\} \\ &= 1 - (3 - 2(1-F(x)))^{\frac{1}{\theta}} (1-F(x))^{\frac{2}{\theta}}; \quad x \in \mathbb{R}, \quad \theta > 0. \end{aligned} \quad (2.3)$$

The corresponding PDF to Eq. (2.3) is

$$g(x) = \frac{dW\{F(x)\}}{dx} \cdot r\{W[F(x)]\} = \frac{6f(x)}{\theta} (1-F(x))^{\frac{2}{\theta}-1} \left(1 - (1-F(x))^{\frac{1}{\theta}}\right). \quad (2.4)$$

3. BILAL - EXPONENTIAL (BE) DISTRIBUTION

In this section, the Bilal-G family of distributions is used to modify the exponential distribution If X has an exponential distribution, then the PDF and CDF are given by:

$$f(x; \beta) = \beta e^{-\beta x}, \quad F(x) = 1 - e^{-\beta x} \quad (3.1)$$

Now, substituting Eq. (3.1) into Eqs. (2.3) and (2.4), we obtain

$$G(x, \theta, \beta) = 1 - (3 - 2e^{-\beta x/\theta})e^{-2\beta x/\theta}, \quad x > 0, \quad \theta, \beta > 0. \quad (3.2)$$

The corresponding PDF is

$$g(x, \theta, \beta) = \frac{6\beta}{\theta} e^{-2\beta x/\theta} \left(1 - e^{-\frac{\beta}{\theta} x}\right). \quad (3.3)$$

The survival and hazard functions are given as

$$\begin{aligned} S(x) &= (3 - 2e^{-\beta x/\theta})e^{-2\beta x/\theta}; \quad \text{and} \\ h(x) &= \frac{g(x)}{S(x)} = \frac{6\beta}{\theta} \cdot \frac{(1 - e^{-(\beta/\theta)x})}{(3 - 2e^{-\beta x/\theta})}, \quad \text{respectively.} \end{aligned}$$

Based on Figure 1, the plots illustrate how the BE distribution's shape changes in response to its two parameters, β and θ . Figure 1a shows the probability density function (PDF), $g(x)$, which can take on various shapes, including right-skewed and nearly symmetric forms, depending on the values of β and θ . For instance, when $\beta = 0.75$ and $\theta = 2.5$ (the solid black line), the distribution has a gentle, long tail, while smaller values of θ and larger values of β , like $\beta = 1.5$ and $\theta = 1.5$ (the green dotted line), result in a sharper, more peaked density near $x = 0$. Figure 1b displays the hazard function, $h(x)$, where all curves consistently show an increasing failure rate over time, indicating that the probability of the event occurring increases as x increases. This characteristic of an increasing hazard rate is common to all parameter combinations shown, though the rate of increase is much steeper for combinations with smaller parameter values, such as $\beta = 0.005$ and $\theta = 0.05$ (the pink dashed line), compared to the shallower curve seen with $\beta = 0.75$ and $\theta = 2.5$ (the solid black line).

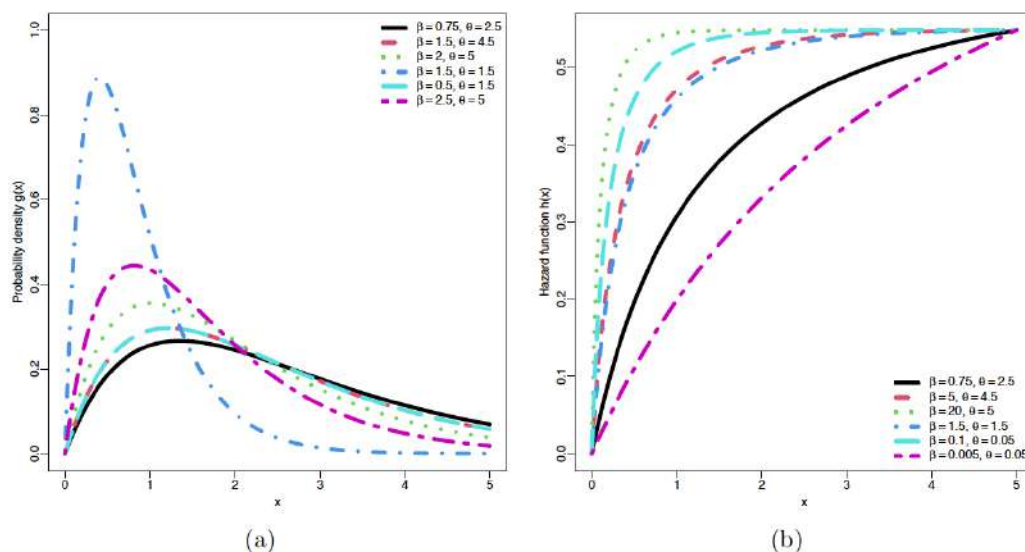


Figure 1: Plots of (a) PDF (b) hazard function for BE distribution.

4. STATISTICAL PROPERTIES

In this section, we study some of the essential properties of the BE model.

4.1. Moment and its Measures

The k th raw moment of a random variable X is given by:

$$E(X^k) = \int_0^\infty x^k f(x) dx = 6\Gamma(k+1) \left(\frac{\theta}{\beta}\right)^k \left[2^{-(k+1)} - 3^{-(k+1)}\right]. \quad (4.1)$$

From Eq. (4.1), the mean is when $k=1$, which is $E(X) = \frac{5\theta}{6\beta}$. The second crude moment is when $k=2$,

which is $E(X^2) = \frac{19}{18} \left(\frac{\theta}{\beta}\right)^2$. The variance of a random variable X is estimated as $Var(X) = E(X^2) - [E(X)]^2 = \frac{13}{36} \left(\frac{\theta}{\beta}\right)^2$. The third crude

moment of X is when $k=3$, which is $E(X^3) = \frac{65}{36} \left(\frac{\theta}{\beta}\right)^3$.

4.2. Moment Generating Function (MGF)

The moment generating function (MGF) of X is defined as:

$$\begin{aligned} M_X(t) &= E(e^{tX}) = \int_0^\infty e^{tx} f(x) dx \\ &= \frac{6\beta^2}{(2\beta - \theta t)(3\beta - \theta t)}, \quad t \in \mathfrak{R}. \end{aligned}$$

To compute the moments, we rewrite the MGF using partial fraction decomposition by setting $\lambda_1 = \frac{2\beta}{\theta}$

and $\lambda_2 = \frac{3\beta}{\theta}$:

$$M_X(t) = \frac{6\beta}{\theta} \left[\frac{1}{\frac{2\beta}{\theta} - t} - \frac{1}{\frac{3\beta}{\theta} - t} \right]$$

The first moment (μ'_1) is $\mu'_1 = \frac{d}{dt} M_X(t) \big|_{t=0}$:

$$\mu'_1 = \frac{6\beta}{\theta} \left[\left(\frac{2\beta}{\theta}\right)^{-2} - \left(\frac{3\beta}{\theta}\right)^{-2} \right]$$

$$\mu'_1 = E(X) = \frac{5\theta}{6\beta}$$

The second moment (μ'_2) is $\mu'_2 = \frac{d^2}{dt^2} M_X(t) \big|_{t=0}$:

$$\mu'_2 = \frac{12\beta}{\theta} \left[\left(\frac{2\beta}{\theta}\right)^{-3} - \left(\frac{3\beta}{\theta}\right)^{-3} \right]$$

$$\mu'_2 = E(X^2) = \frac{19\theta^2}{18\beta^2}$$

The mean of the BE distribution, illustrated in Figure 2a, generally decreases as the parameter α increases, particularly for smaller values of β . Conversely, as the parameter β increases, the mean exhibits a complex behavior but generally shows an increasing trend

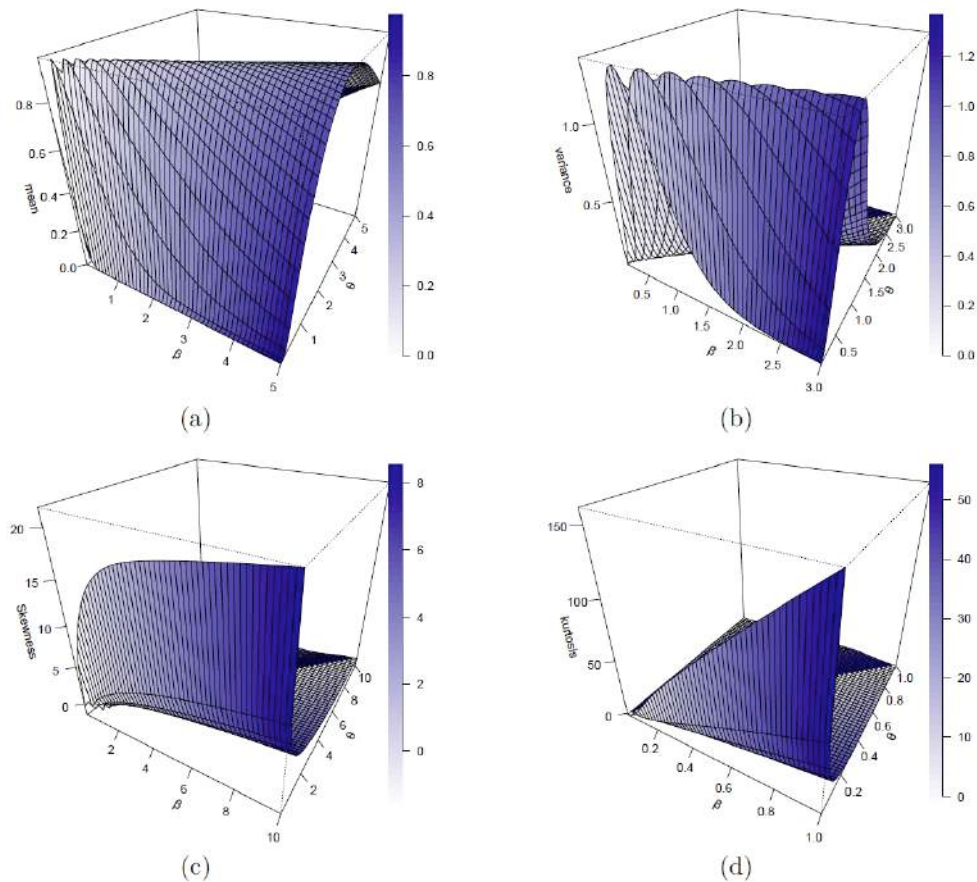


Figure 2: Plots (a) Mean (b) Variance (c) Skewness (d) Kurtosis of BE distribution.

before stabilizing, especially when α is small. The mean's value is observed to range from nearly 0 to approximately 0.8 within the plotted parameter space.

The variance, shown in Figure 2b, typically decreases with increasing values of both α and β . For instance, when α is close to 0, the variance is at its highest, approaching 1.0, but it drops rapidly as α increases towards 3.0. This indicates that a larger α leads to a more concentrated distribution. The variance surface demonstrates relatively sharp changes at the boundaries of the parameter space and remains generally low over a broad region, suggesting that the BE distribution is often less dispersed.

The skewness, displayed in Figure 2c, is strictly positive across the entire parameter space, which confirms the BE distribution is right-skewed. The maximum skewness, which is quite high, occurs when both α and β are small, indicating a pronounced long right tail under these conditions. As α and β increase, the skewness value rapidly decreases towards zero, implying the distribution becomes more symmetrical and approaches the shape of a normal distribution.

Finally, the kurtosis, presented in Figure 2d, also displays a similar pattern to the skewness, being

highest when the parameters α and β are small. This large kurtosis suggests that the distribution is leptokurtic (has heavy tails and a sharp peak) for small parameter values. The kurtosis decreases dramatically as α and β increase, approaching a value closer to 3 (the kurtosis of a normal distribution), indicating that the tail behavior becomes less extreme and the peak flatter as the parameters grow.

4.3. Order Statistics

The CDF and PDF of the k th order statistic, $X_{(k)}$, based on a random sample of size n , are derived from the general forms:

$$F_k(x) = \sum_{j=k}^n \binom{n}{j} (F_X(x))^j (1 - F_X(x))^{n-j}$$

$$f_k(x) = \frac{n!}{(k-1)!(n-k)!} (F_X(x))^{k-1} (1 - F_X(x))^{n-k} f_X(x)$$

Given the BE distribution's CDF, $F(x) = 1 - (3 - 2e^{-\frac{\beta x}{\theta}})e^{-\frac{2\beta x}{\theta}}$, and PDF, $f(x) = \frac{6\beta}{\theta} e^{-\frac{2\beta x}{\theta}} (1 - e^{-\frac{\beta x}{\theta}})$, we have:

The CDF of $X_{(k)}$ is:

$$F_k(x; \theta, \beta) = \sum_{j=k}^n \binom{n}{j} \left[1 - (3 - 2e^{-\frac{\beta x}{\theta}}) e^{-\frac{2\beta x}{\theta}} \right]^j \left[(3 - 2e^{-\frac{\beta x}{\theta}}) e^{-\frac{2\beta x}{\theta}} \right]^{n-j}$$

The PDF of $X_{(k)}$ is:

$$f_k(x; \theta, \beta) = n \binom{n-1}{k-1} \left[1 - (3 - 2e^{-\frac{\beta x}{\theta}}) e^{-\frac{2\beta x}{\theta}} \right]^{k-1} \left[(3 - 2e^{-\frac{\beta x}{\theta}}) e^{-\frac{2\beta x}{\theta}} \right]^{n-k} \cdot \frac{6\beta}{\theta} e^{-\frac{2\beta x}{\theta}} \left(1 - e^{-\frac{\beta x}{\theta}} \right)$$

4.4. Quantile Function

The quantile function, $Q(p)$, is the inverse of the CDF, $G(x; \theta, \beta) = 1 - 3e^{-\frac{2\beta x}{\theta}} + 2e^{-\frac{3\beta x}{\theta}}$, such that $G(Q(p)) = p$.

Solving the equation $1 - 3e^{-\frac{2\beta x}{\theta}} + 2e^{-\frac{3\beta x}{\theta}} = p$ for x requires numerical methods or approximation. Using the approximation $\ln\left(1 - 3e^{-\frac{2\beta x}{\theta}} + 2e^{-\frac{3\beta x}{\theta}}\right) \approx \ln(6) - \frac{5\beta x}{\theta}$, we solve:

$$\ln\left(1 - 3e^{-\frac{2\beta x}{\theta}} + 2e^{-\frac{3\beta x}{\theta}}\right) \approx \ln(p)$$

The approximate quantile function is:

$$Q(p) \approx \frac{\theta(\ln(p) - 1.7918)}{5\beta}$$

The proposed approximate quantile function, $Q_{approx}(p) \approx \frac{\theta(\ln(p) - 1.7918)}{5\beta}$, relies on a log-linear approximation that requires critical validation. Given the complexity of the exact CDF, $G(x; \theta, \beta) = 1 - 3e^{-\frac{2\beta x}{\theta}} + 2e^{-\frac{3\beta x}{\theta}}$, this approximation must be rigorously assessed to ensure the fidelity of subsequent simulation results. To quantify the error, a numerical comparison against the exact quantile function, $Q_{num}(p)$, must be performed.

- Numerical Inversion: Obtain $Q_{num}(p)$ by numerically solving the equation $G(x; \theta, \beta) - p = 0$ for x , using a stable root-finding algorithm (e.g., Bisection or Newton-Raphson) across a representative grid of probabilities $p \in (0, 1)$.

- Error Metric: The Relative Error (RE) is the most appropriate metric for evaluation, particularly over the tails of the distribution ($p \rightarrow 0$ and $p \rightarrow 1$).

$$RE(p) = \left| \frac{Q_{approx}(p) - Q_{num}(p)}{Q_{num}(p)} \right|$$

The approximation introduces a systematic bias into the simulated random variates.

- The magnitude of the RE directly correlates with the inaccuracy of the simulation results (e.g., estimated mean, variance, and, crucially, extreme percentiles).
- If the error is high in the tails, the simulation will fail to accurately represent extreme events, compromising any risk or reliability analyses performed using the simulated data.

If the quantified error is deemed unacceptable (e.g., $RE > 5\%$ for critical values of p), it is strongly recommended that numerical inversion be used in practice for generating random variates. Root-finding algorithms are highly efficient on modern computing platforms, making the increased accuracy achieved via numerical inversion a worthwhile trade-off for any minor increase in computational time. For very low quantiles ($p \rightarrow 0$), an approximation derived from a Taylor expansion ($G(x) \approx 3(\beta x / \theta)^2$) may yield better accuracy: $Q(p) \approx \frac{\theta}{\beta} \sqrt{\frac{p}{3}}$. This alternative should also be tested.

4.5. Entropy

The Rényi entropy of order α is defined as $H_\alpha(X) = \frac{1}{1-\alpha} \log\left(\int p^\alpha(x) dx\right)$. For the BE distribution:

$$H_\alpha(X) = \frac{1}{1-\alpha} \log\left(\int_0^\infty \left[\frac{6\beta}{\theta} e^{-\frac{2\beta}{\theta}x} \left(1 - e^{-\frac{\beta}{\theta}x}\right)\right]^\alpha dx\right)$$

Using the substitution $u = e^{-\frac{\beta}{\theta}x}$, the integral evaluates to a Beta function:

$$\int_0^\infty [f(x)]^\alpha dx = \left(\frac{6\beta}{\theta}\right)^\alpha \cdot \frac{\theta}{\beta} B(2\alpha, \alpha + 1)$$

The final expression for the Rényi entropy is:

$$H_\alpha(X) = \frac{1}{1-\alpha} \left[\alpha \log\left(\frac{6\beta}{\theta}\right) + \log\left(\frac{\theta}{\beta}\right) + \log \Gamma(\alpha + 1) + \log \Gamma(2\alpha) - \log \Gamma(3\alpha + 1) \right]$$

4.6. Mean Residual Life Function

The mean residual life function (MRL) is $mrl(x) = \frac{\int_x^\infty S(t)dt}{S(x)}$, where the survival function is $S(x) = (3 - 2e^{\frac{\beta}{\theta}x})e^{-\frac{2\beta}{\theta}x}$.

Substituting $S(x)$ and letting $\eta = \frac{\beta}{\theta}$, we evaluate the integral:

$$mrl(x) = \frac{\int_x^\infty (3e^{-2\eta t} - 2e^{-3\eta t})dt}{(3 - 2e^{-\eta x})e^{-2\eta x}}$$

After evaluating the definite integrals and simplifying:

$$mrl(x) = \frac{\theta}{\beta} \cdot \frac{3 - 2e^{-\frac{\beta}{\theta}x}}{3 - 2e^{-\frac{\beta}{\theta}x}}$$

4.7. Stress-Strength Reliability for the BE Distribution

The stress-strength reliability is $R = P(X > Y) = \int_0^\infty [1 - F_X(y)]f_Y(y)dy$. Substituting the BE distributions for X and Y :

$$R = \frac{6\beta_Y}{\theta_Y} \int_0^\infty \left(3e^{-\frac{2\beta_X y}{\theta_X}} - 2e^{-\frac{3\beta_X y}{\theta_X}} \right) \cdot e^{-\frac{2\beta_Y y}{\theta_Y}} \left(1 - e^{-\frac{\beta_Y y}{\theta_Y}} \right) dy$$

Evaluating the integral, which involves terms of the form $\int_0^\infty e^{-ky} dy = 1/k$, gives the result:

$$R = \frac{6\beta_Y}{\theta_Y} \left[3 \left(\frac{1}{k_1} - \frac{1}{k_2} \right) - 2 \left(\frac{1}{k_3} - \frac{1}{k_4} \right) \right]$$

Where $k_1 = \frac{2\beta_X}{\theta_X} + \frac{2\beta_Y}{\theta_Y}$, $k_2 = \frac{2\beta_X}{\theta_X} + \frac{3\beta_Y}{\theta_Y}$,
 $k_3 = \frac{3\beta_X}{\theta_X} + \frac{2\beta_Y}{\theta_Y}$, and $k_4 = \frac{3\beta_X}{\theta_X} + \frac{3\beta_Y}{\theta_Y}$.

5. ESTIMATION

In this comprehensive analysis, we utilized fifteen distinct non-Bayesian estimation approaches to determine the unknown parameters (θ, β) of the BE distribution. These methods—spanning likelihood-based, distance-based, and spacing-based techniques—include Maximum Likelihood (MLE),

Anderson-Darling (ADE), Cramér-von Mises (CVME), Maximum Product of Spacings (MPSE), Ordinary Least Squares (OLSE), Right-Tail Anderson-Darling (RTADE), Weighted Least Squares (WLSE), Left-Tail Anderson-Darling (LTADE), Minimum Spacing Absolute Distance (MSADE), Minimum Spacing Absolute-Log Distance (MSALDE), Anderson-Darling Second-Order Left-Tail (ADSOE), Kolmogorov (KE), Minimum Spacing Square Distance (MSSDE), Minimum Spacing Squared-Log Distance (MSSLDE), and Minimum Spacing Linex Distance (MSLNDE). The specific objective functions for each method are detailed in the following subsections.

5.1. Maximum Likelihood Estimation (MLE)

The Method of Maximum Likelihood, introduced by [29] and [30], is widely preferred due to its favorable asymptotic properties, including consistency and asymptotic normality. Given a random sample $x_1, \dots, x_n$ of size n from the BE distribution's PDF $g(x; \theta, \beta)$, the likelihood function $L(\theta, \beta)$ is defined as the product of the PDF evaluated at each observation.

The log-likelihood function, $\ell(\theta, \beta) = \ln L(\theta, \beta)$, is derived by substituting the PDF given in Equation (3.3):

$$\ell = n \ln(6\beta) - n \ln \theta - \frac{2\beta}{\theta} \sum_{i=1}^n x_i + \sum_{i=1}^n \ln \left(1 - e^{-\frac{\beta}{\theta} x_i} \right).$$

The maximum likelihood estimators $\hat{\theta}, \hat{\beta}$ are obtained by solving the system of non-linear equations derived by setting the partial derivatives of ℓ with respect to each parameter to zero:

$$\frac{\partial \ell}{\partial \theta} = -\frac{n}{\theta} + \frac{2\beta}{\theta^2} \sum_{i=1}^n x_i - \frac{\beta}{\theta^2} \sum_{i=1}^n \frac{x_i}{e^{\frac{\beta}{\theta} x_i} - 1} = 0$$

$$\frac{\partial \ell}{\partial \beta} = \frac{n}{\beta} - \frac{2}{\theta} \sum_{i=1}^n x_i + \frac{1}{\theta} \sum_{i=1}^n \frac{x_i}{e^{\frac{\beta}{\theta} x_i} - 1} = 0$$

Since these equations are analytically intractable (non-linear), the estimates $\hat{\theta}, \hat{\beta}$ are solved numerically using iterative optimization algorithms.

The Hessian matrix $\mathbf{H}$ is defined as:

$$\mathbf{H} = \begin{pmatrix} \frac{\partial^2 \ell}{\partial \theta^2} & \frac{\partial^2 \ell}{\partial \theta \partial \beta} \\ \frac{\partial^2 \ell}{\partial \beta \partial \theta} & \frac{\partial^2 \ell}{\partial \beta^2} \end{pmatrix}.$$

$$\frac{\partial^2 \ell}{\partial \theta^2} = \frac{n}{\theta^2} - \frac{4\beta}{\theta^3} \sum_{i=1}^n x_i + \frac{2\beta}{\theta^3} \sum_{i=1}^n \frac{x_i}{e^{\frac{\beta}{\theta} x_i} - 1} - \frac{\beta^2}{\theta^4} \sum_{i=1}^n \frac{x_i^2 e^{\frac{\beta}{\theta} x_i}}{\left(e^{\frac{\beta}{\theta} x_i} - 1\right)^2}.$$

$$\frac{\partial^2 \ell}{\partial \theta \partial \beta} = \frac{\partial^2 \ell}{\partial \beta \partial \theta} = \frac{2}{\theta^2} \sum_{i=1}^n x_i - \frac{1}{\theta^2} \sum_{i=1}^n \frac{x_i}{e^{\frac{\beta}{\theta} x_i} - 1} + \frac{\beta}{\theta^3} \sum_{i=1}^n \frac{x_i^2 e^{\frac{\beta}{\theta} x_i}}{\left(e^{\frac{\beta}{\theta} x_i} - 1\right)^2}.$$

$$\frac{\partial^2 \ell}{\partial \beta^2} = -\frac{n}{\beta^2} - \frac{1}{\theta^2} \sum_{i=1}^n \frac{x_i^2 e^{\frac{\beta}{\theta} x_i}}{\left(e^{\frac{\beta}{\theta} x_i} - 1\right)^2}.$$

The Newton-Raphson Algorithm (Algorithm 5.1) was used to solve the system of non-linear score equations and obtain the MLEs.

Algorithm 1: Newton-Raphson for BE MLE

Require: Observed data $\mathbf{x} = \{x_1, \dots, x_n\}$, initial parameters $\theta^{(0)} = (\theta^{(0)}, \beta^{(0)})$, convergence tolerance ε , max iterations M .

Ensure: Maximum Likelihood Estimates $\hat{\theta} = (\hat{\theta}, \hat{\beta})$ and Asymptotic Variance-Covariance Matrix $\mathbf{V}$.

1: $k \leftarrow 0$

2: $\theta \leftarrow \theta^{(0)}$

3: **repeat**

4: Step 1: Calculate the Score Vector (First Derivatives) at $\theta^{(k)}$

$$5: \mathbf{S}^{(k)} \leftarrow \nabla \ell(\theta^{(k)}) = \begin{pmatrix} \frac{\partial \ell}{\partial \theta} \\ \frac{\partial \ell}{\partial \beta} \end{pmatrix}_{\theta^{(k)}}$$

6: Step 2: Calculate the Hessian Matrix (Second Derivatives) at $\theta^{(k)}$

$$7: \mathbf{H}^{(k)} \leftarrow \begin{pmatrix} \frac{\partial^2 \ell}{\partial \theta^2} & \frac{\partial^2 \ell}{\partial \theta \partial \beta} \\ \frac{\partial^2 \ell}{\partial \beta \partial \theta} & \frac{\partial^2 \ell}{\partial \beta^2} \end{pmatrix}_{\theta^{(k)}}$$

8: Step 3: Calculate the Step Direction $\delta^{(k)}$

9: Solve the linear system: $\mathbf{H}^{(k)} \delta^{(k)} = -\mathbf{S}^{(k)}$

10: Step 4: Update Parameters

11: $\theta^{(k+1)} \leftarrow \theta^{(k)} + \delta^{(k)}$

12: Step 5: Check Convergence

13: **If** $\|\delta^{(k)}\| < \varepsilon$ **or** $k \geq M$ **then**

14: **break**

15: **end if**

16: $k \leftarrow k + 1$

17: $\theta^{(k)} \leftarrow \theta^{(k+1)}$

18: **until** Convergence

19: $\hat{\theta} \leftarrow \theta^{(k+1)}$ {MLE Estimates}

20: Step 6: Calculate Asymptotic Variance-Covariance Matrix

21: $\mathbf{I}_o \leftarrow -\mathbf{H}(\hat{\theta})$ {Observed Fisher Information}

22: $\mathbf{V} \leftarrow \mathbf{I}_o^{-1}$ {Asymptotic Variance-Covariance Matrix}

5.2. Anderson-Darling Method (ADE)

The ADE approach, pioneered by [31] and [32], seeks to minimize a weighted goodness-of-fit statistic that emphasizes the tails of the distribution. Parameter estimates for the BE model are derived by minimizing the following Anderson-Darling statistic, A :

$$A(\theta, \beta) = -n - \frac{1}{n} \sum_{i=1}^n (2i-1) [\log G(x_{i:n}; \theta, \beta) + \log S(x_{n-i+1:n}; \theta, \beta)],$$

where $x_{i:n}$ are the order statistics, $G(x; \theta, \beta)$ is the CDF, and $S(x; \theta, \beta) = 1 - G(x; \theta, \beta)$ is the survival function of the BE distribution.

5.3. Method of Cramér-von Mises (CVME)

The CVME approach, often credited to [33], estimates parameters by finding values that minimize the Cramér-von Mises criterion. The parameters $\hat{\theta}, \hat{\beta}$ for the BE distribution are obtained by minimizing the following statistic, C :

$$C(\theta, \beta) = \frac{1}{12n} + \sum_{i=1}^n \left[G(x_{i:n}; \theta, \beta) - \frac{2i-1}{2n} \right]^2.$$

5.4. Ordinary Least Squares Estimation (OLSE)

The OLSE technique, utilized in distributional fitting by [34], aims to minimize the squared differences

between the empirical CDF and the theoretical CDF. The BE parameters are estimated by minimizing the Ordinary Least Squares statistic, V :

$$V(\theta, \beta) = \sum_{i=1}^n \left[G(x_{in}; \theta, \beta) - \frac{i}{n+1} \right]^2.$$

5.5. Right-Tail Anderson-Darling Estimation (RTADE)

The parameters $\hat{\theta}, \hat{\beta}$ are found by minimizing the RTADE statistic, R :

$$R(\theta, \beta) = \frac{n}{2} - 2 \sum_{i=1}^n G(x_{in}; \theta, \beta) - \frac{1}{n} \sum_{i=1}^n (2i-1) \log S(x_{in}; \theta, \beta).$$

5.6. Weighted Least Squares Estimation (WLSE)

As a robust variant of OLSE, the WLSE method [34] introduces weights to the squared differences. The parameter estimates for the BE distribution are derived by minimizing the Weighted Least Squares statistic, W :

$$W(\theta, \beta) = \sum_{i=1}^n \frac{(n+1)^2 (n+2)}{i(n-i+1)} \left[G(x_{in}; \theta, \beta) - \frac{i}{n+1} \right]^2.$$

5.7. Left-Tail Anderson-Darling Estimation (LTADE)

The LTADE method is based on optimizing the fit of the theoretical model to the observed data in the lower tail, following procedures detailed in studies such as [35]. The parameters $\hat{\theta}, \hat{\beta}$ for the BE model are estimated by minimizing the LTADE statistic, L :

$$L(\theta, \beta) = -\frac{3}{2}n + 2 \sum_{i=1}^n G(x_{in}; \theta, \beta) - \frac{1}{n} \sum_{i=1}^n (2i-1) \log G(x_{in}; \theta, \beta).$$

5.8. Anderson-Darling Second-Order Left-Tail (ADSOE)

This refined technique [36] investigates the second-order behavior of the distribution in the left tail. The BE parameters are estimated by minimizing the ADSOE statistic, LTS :

$$LTS(\theta, \beta) = 2 \sum_{i=1}^n \log G(x_i; \theta, \beta) + \frac{1}{n} \sum_{i=1}^n \frac{(2i-1)}{G(x_i; \theta, \beta)}.$$

5.9. Kolmogorov Estimation (KE)

The Kolmogorov method, also popularized by [36], focuses on minimizing the maximum distance between

the empirical and theoretical CDFs. The BE parameters are estimated by minimizing the Kolmogorov statistic, KM :

$$KM(\theta, \beta) = \max_{1 \leq i \leq n} \left[\frac{i}{n} - G(x_i; \theta, \beta), G(x_i; \theta, \beta) - \frac{i-1}{n} \right].$$

5.10. Estimation Methods Based on Product of Spacings (PS)

These methods rely on the spacings $I_i = G(x_{in}; \theta, \beta) - G(x_{i-1n}; \theta, \beta)$, where $G(x_{0n}; \theta, \beta) = 0$ and $G(x_{n+1n}; \theta, \beta) = 1$. The spacings I_i are now a function of only θ and β .

5.11. Maximum Product of Spacings Estimation (MPSE)

The MPSE method, proposed by [37], is a nonparametric alternative to MLE. The BE parameters are found by minimizing the negative log-product of spacings:

$$\delta(\theta, \beta) = -\frac{1}{n+1} \sum_{i=1}^{n+1} \log I_i(\theta, \beta).$$

5.12. Minimum Spacing Absolute Distance (MSADE)

The MSADE approach [38] minimizes the sum of absolute distances. The MSADE statistic, ζ , is:

$$\zeta(\theta, \beta) = \sum_{i=1}^{n+1} \left| I_i - \frac{1}{n+1} \right|.$$

5.13. Minimum Spacing Absolute-Log Distance (MSALDE)

The MSALDE method, also proposed by [38], uses a logarithmic transformation of the spacings. The BE parameters are estimated by minimizing the MSALDE statistic, Y :

$$Y(\theta, \beta) = \sum_{i=1}^{n+1} \left| \log I_i - \log \frac{1}{n+1} \right|.$$

5.14. Minimum Spacing Square Distance (MSSDE)

The BE parameters are estimated by minimizing the MSSDE statistic, φ :

$$\varphi(\theta, \beta) = \sum_{i=1}^{n+1} \left(I_i - \frac{1}{n+1} \right)^2.$$

5.15. Minimum Spacing Squared-Log Distance (MSSLDE)

The BE parameters are derived by minimizing the MSSLDE statistic, Ψ :

$$\Psi(\theta, \beta) = \sum_{i=1}^{n+1} \left(\log I_i - \log \frac{1}{n+1} \right)^2.$$

5.16. Minimum Spacing Linex Distance (MSLNDE)

The MSLNDE technique employs a Linex-type loss function. The BE parameters are estimated by minimizing the MSLNDE statistic, Δ :

$$\Delta(\theta, \beta) = \sum_{i=1}^{n+1} \left[e^{I_i - \frac{1}{n+1}} - \left(I_i - \frac{1}{n+1} \right) - 1 \right].$$

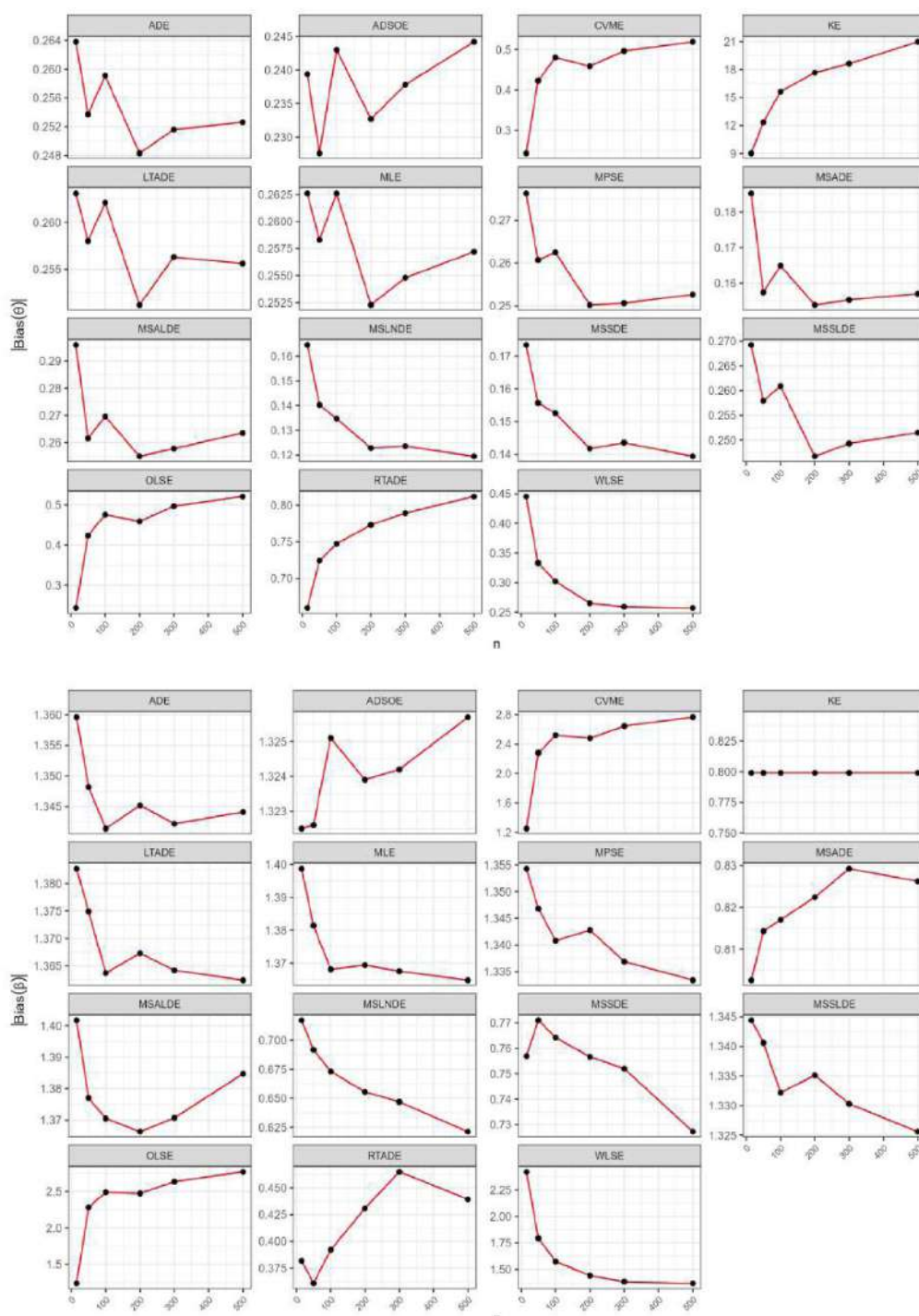


Figure 3: Graphical representations for the Bias values presented in Table 1.

6. SIMULATION

This section evaluates the efficacy of several estimation techniques in predicting the parameters of the BE distribution using a substantial amount of simulated data. Random data sets were created for various sample sizes ($n = 15, 50, 100, 200, 300$, and 500) using the quantile function of the suggested model in

the simulation run. This section will evaluate the efficacy and performance of our model estimators. Furthermore, we will assess the effectiveness of several estimation methods considering many aspects. The evaluation will focus on the accuracy, precision, and computing efficiency of each method. By conducting a comprehensive examination of data from diverse sample sizes, our objective was to provide

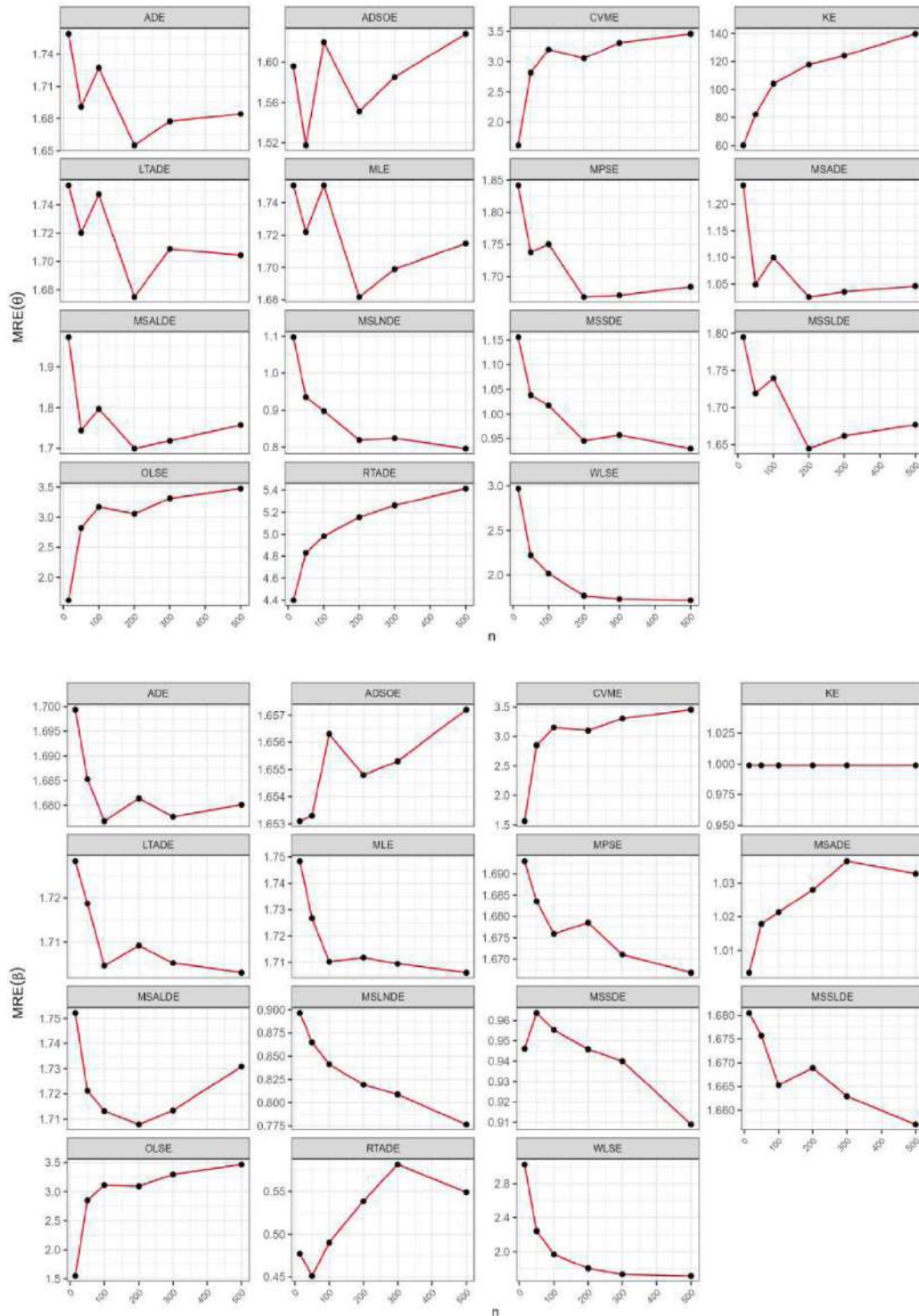


Figure 4: Graphical representations for the MRE values presented in Table 1.

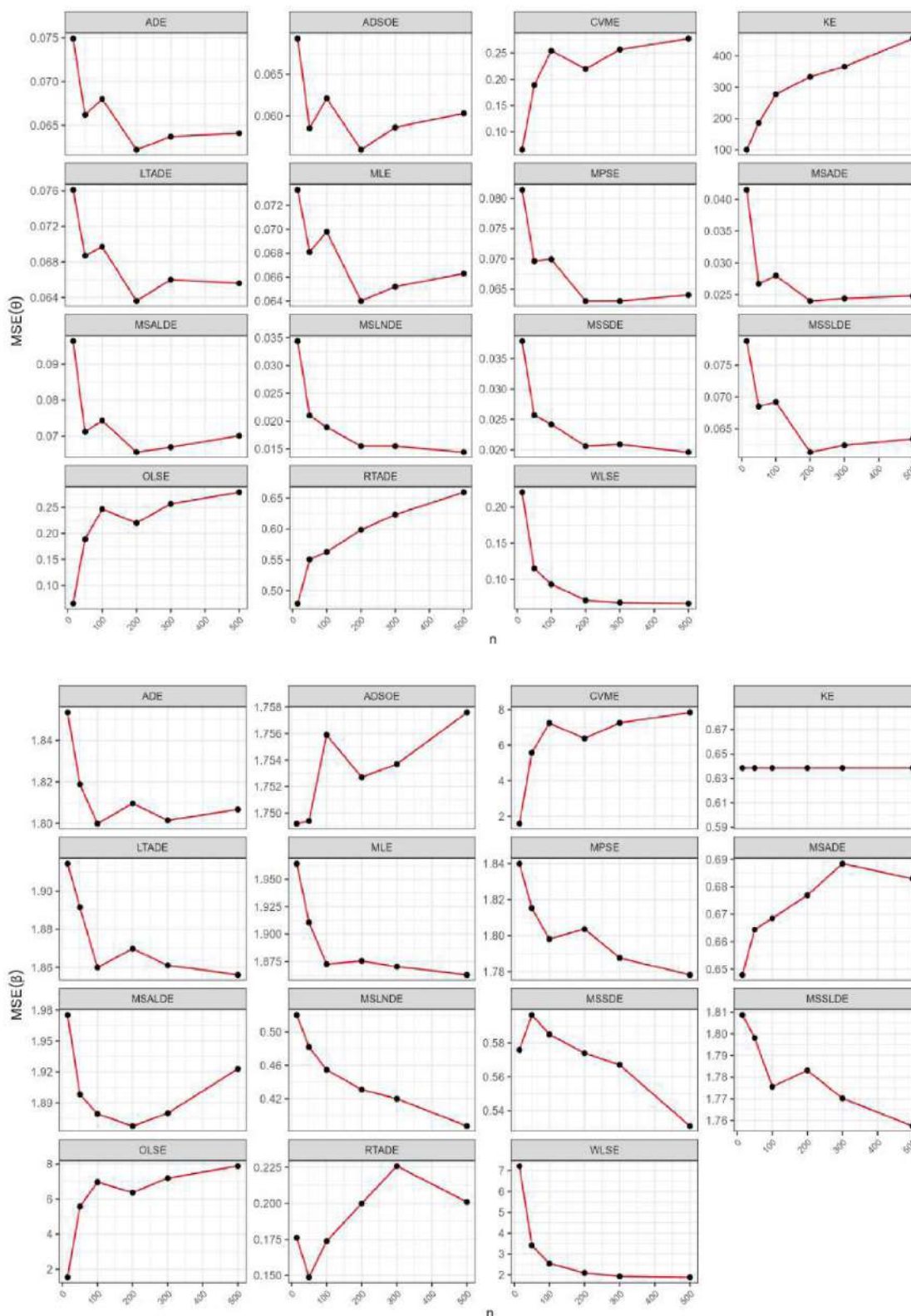


Figure 5: Graphical representations for the MSE values presented in Table 1.

significant insight into the most reliable estimation approach for the BE distribution. In addition, we will evaluate the efficacy of each estimating approach in various settings to determine its dependability and applicability in actual contexts. This research will help

researchers select the most suitable estimation approach for their specific needs. Furthermore, we will assess the effectiveness of various estimation methodologies, taking into account several elements such as

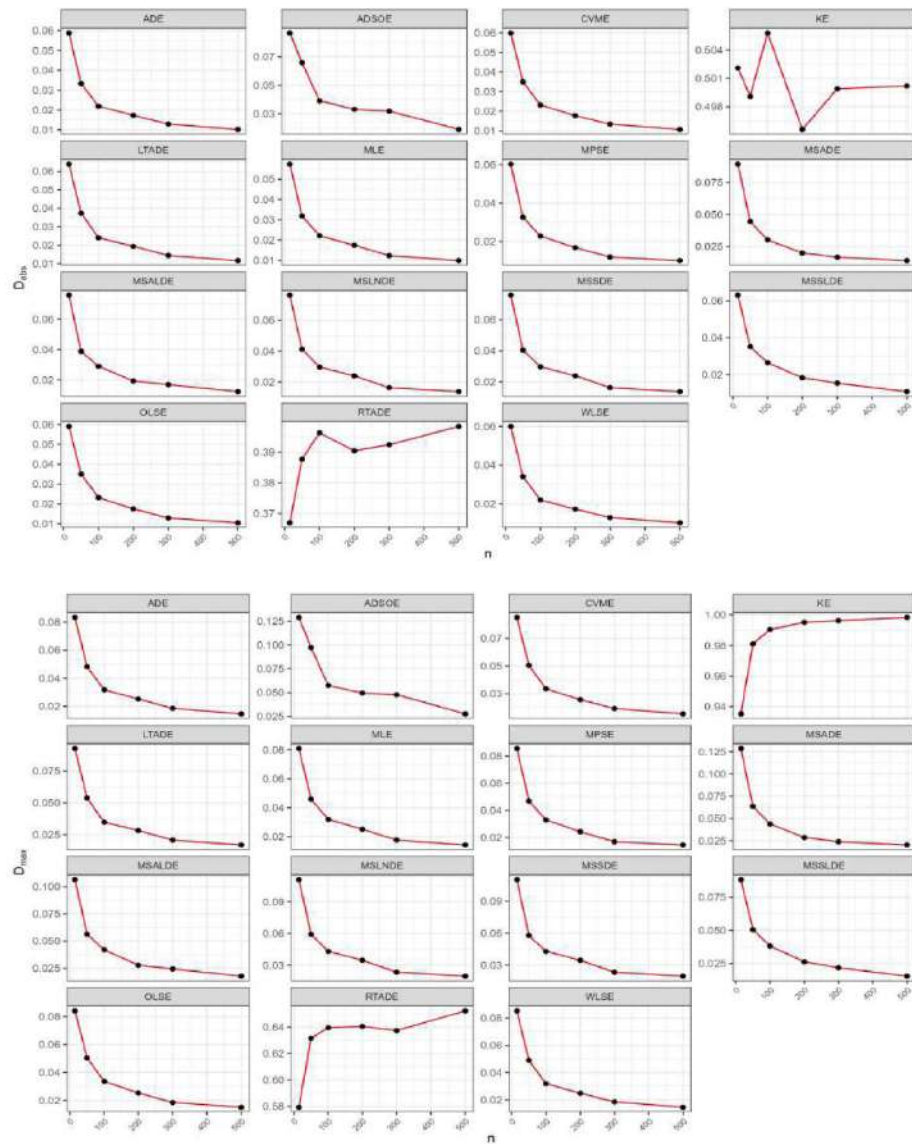


Figure 6: Graphical representations for the D_{abs} and D_{max} values presented in Table 1.

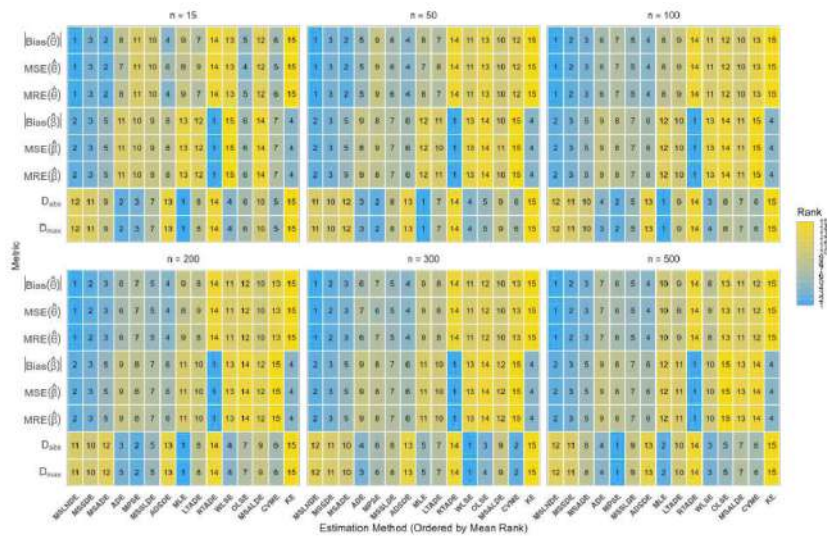


Figure 7: Partial Rank Heatmap for simulation in Table 1.

Table 1: Numerical Values of Simulation Measures for $\theta = 0.15, \beta = 0.8$.

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLSE	LTADE	MSADE	MSALDE	ADSOE	KE	MSSE	MSLDE	MSLNDE	
15	$ Bias(\hat{\theta}) $	0.2626 ⁽⁷⁾	0.2638 ⁽⁸⁾	0.2430 ⁽⁶⁾	0.2763 ⁽¹¹⁾	0.2431 ⁽⁶⁾	0.6597 ⁽¹⁴⁾	0.4454 ⁽¹³⁾	0.2631 ⁽⁸⁾	0.1852 ⁽⁶⁾	0.2569 ⁽¹²⁾	0.2394 ⁽⁴⁾	9.0219 ⁽¹²⁾	0.1734 ⁽²⁾	0.2692 ⁽¹⁰⁾	0.1646 ⁽³⁾	
	$ Bias(\hat{\beta}) $	1.3987 ⁽¹⁰⁾	1.3596 ⁽¹¹⁾	1.2508 ⁽⁷⁾	1.3543 ⁽¹⁰⁾	1.2370 ⁽⁶⁾	0.3818 ⁽¹⁾	2.4205 ⁽¹⁶⁾	1.3827 ⁽¹²⁾	0.8026 ⁽³⁾	1.4017 ⁽¹⁴⁾	1.3225 ⁽⁸⁾	0.7990 ⁽⁴⁾	0.7569 ⁽⁸⁾	1.3444 ⁽⁶⁾	0.7172 ⁽²⁾	
	MSE ($\hat{\theta}$)	0.0733 ⁽⁷⁾	0.0749 ⁽⁸⁾	0.0654 ⁽¹¹⁾	0.0814 ⁽¹¹⁾	0.0654 ⁽¹¹⁾	0.4794 ⁽¹⁴⁾	0.2204 ⁽¹³⁾	0.0761 ⁽¹²⁾	0.0415 ⁽³⁾	0.0693 ⁽⁹⁾	0.0693 ⁽⁹⁾	100.0431 ⁽¹⁵⁾	0.0379 ⁽²⁾	0.0788 ⁽¹⁰⁾	0.0344 ⁽¹⁾	
	MSE ($\hat{\beta}$)	1.9643 ⁽¹³⁾	1.8534 ⁽¹¹⁾	1.6731 ⁽⁷⁾	1.8399 ⁽¹⁰⁾	1.5369 ⁽⁶⁾	0.1762 ⁽³⁾	7.2129 ⁽¹⁶⁾	1.9143 ⁽¹²⁾	0.6479 ⁽³⁾	1.9753 ⁽¹⁴⁾	1.7492 ⁽⁶⁾	6.6384 ⁽⁴⁾	0.5760 ⁽³⁾	1.8087 ⁽⁹⁾	0.5205 ⁽²⁾	
	MRE ($\hat{\theta}$)	1.7507 ⁽⁷⁾	1.7590 ⁽⁸⁾	1.6203 ⁽³⁾	1.8419 ⁽¹¹⁾	1.6206 ⁽⁶⁾	4.3082 ⁽¹⁴⁾	2.9690 ⁽¹⁶⁾	1.7537 ⁽¹²⁾	1.2348 ⁽³⁾	1.9729 ⁽¹²⁾	1.5959 ⁽⁴⁾	60.1460 ⁽¹⁵⁾	1.1561 ⁽²⁾	1.7947 ⁽¹⁰⁾	1.0971 ⁽¹⁾	
	MRE ($\hat{\beta}$)	1.7484 ⁽¹³⁾	1.6929 ⁽¹¹⁾	1.5635 ⁽⁷⁾	1.6929 ⁽¹⁰⁾	1.5463 ⁽⁶⁾	4.4773 ⁽¹³⁾	3.0256 ⁽¹⁶⁾	1.7283 ⁽¹²⁾	1.0033 ⁽³⁾	1.7521 ⁽¹⁴⁾	1.6531 ⁽⁴⁾	50.9881 ⁽¹⁵⁾	0.9461 ⁽²⁾	1.6805 ⁽⁹⁾	0.8965 ⁽²⁾	
	D _{abs}	0.0574 ⁽¹⁾	0.0588 ⁽²⁾	0.0596 ⁽⁴⁾	0.0603 ⁽³⁾	0.0591 ⁽³⁾	0.3670 ⁽¹⁴⁾	0.0600 ⁽¹⁾	0.0639 ⁽³⁾	0.0892 ⁽¹³⁾	0.0759 ⁽¹⁰⁾	0.0866 ⁽¹²⁾	0.5021 ⁽¹⁶⁾	0.0758 ⁽⁹⁾	0.0630 ⁽⁷⁾	0.0762 ⁽¹¹⁾	
	D _{max}	0.0809 ⁽¹⁾	0.0833 ⁽²⁾	0.0832 ⁽⁴⁾	0.0855 ⁽³⁾	0.0841 ⁽³⁾	0.5794 ⁽¹⁴⁾	0.0854 ⁽⁵⁾	0.0926 ⁽³⁾	0.1285 ⁽¹²⁾	0.1065 ⁽⁹⁾	0.1290 ⁽¹³⁾	0.9353 ⁽¹⁵⁾	0.1104 ⁽¹⁰⁾	0.0882 ⁽⁷⁾	0.1109 ⁽¹¹⁾	
	$\sum$ Ranks	62 ⁽⁶⁾	63 ⁽⁷⁾	44 ⁽¹⁾	75 ⁽¹¹⁾	40 ⁽³⁾	73 ⁽¹⁰⁾	94 ⁽¹⁴⁾	77 ⁽¹²⁾	49 ⁽³⁾	63 ⁽⁷⁾	63 ⁽⁷⁾	87 ⁽¹²⁾	34 ⁽²⁾	71 ⁽⁹⁾	31 ⁽¹⁾	
	50	$ Bias(\hat{\theta}) $	0.2583 ⁽⁸⁾	0.2537 ⁽⁶⁾	0.4229 ⁽¹²⁾	0.2607 ⁽⁸⁾	0.4234 ⁽¹³⁾	0.7243 ⁽¹⁴⁾	0.3330 ⁽¹¹⁾	0.2580 ⁽⁷⁾	0.1574 ⁽⁸⁾	0.2615 ⁽¹⁰⁾	0.2276 ⁽⁴⁾	12.3294 ⁽¹⁵⁾	0.1557 ⁽²⁾	0.2579 ⁽⁸⁾	0.1403 ⁽¹⁾
		$ Bias(\hat{\beta}) $	1.3814 ⁽¹²⁾	1.3482 ⁽⁹⁾	2.2502 ⁽¹⁴⁾	1.3468 ⁽⁸⁾	2.2518 ⁽¹⁵⁾	0.3607 ⁽¹⁾	1.7924 ⁽¹⁶⁾	1.3749 ⁽¹⁰⁾	0.8143 ⁽³⁾	1.3770 ⁽¹¹⁾	1.3226 ⁽⁸⁾	0.7990 ⁽⁴⁾	0.7710 ⁽³⁾	1.3406 ⁽⁷⁾	0.6918 ⁽²⁾
		MSE ($\hat{\theta}$)	0.0681 ⁽¹⁾	0.0662 ⁽²⁾	0.1886 ⁽¹²⁾	0.0696 ⁽³⁾	0.1889 ⁽¹³⁾	0.5505 ⁽¹⁴⁾	0.1152 ⁽¹¹⁾	0.0687 ⁽³⁾	0.0267 ⁽³⁾	0.0712 ⁽¹⁰⁾	0.0585 ⁽⁴⁾	185.8662 ⁽¹⁵⁾	0.0257 ⁽²⁾	0.0685 ⁽⁷⁾	0.0210 ⁽¹⁾
		MSE ($\hat{\beta}$)	1.9106 ⁽¹²⁾	1.8188 ⁽⁹⁾	5.5754 ⁽¹⁶⁾	1.8152 ⁽⁸⁾	5.5666 ⁽¹⁴⁾	3.4187 ⁽¹¹⁾	3.2201 ⁽¹³⁾	1.8915 ⁽¹⁰⁾	0.6744 ⁽³⁾	1.8982 ⁽¹¹⁾	1.7494 ⁽⁶⁾	6.6384 ⁽⁴⁾	0.5067 ⁽³⁾	1.7981 ⁽⁷⁾	0.4820 ⁽²⁾
		MRE ($\hat{\theta}$)	1.7220 ⁽⁸⁾	1.6911 ⁽⁵⁾	2.8191 ⁽¹²⁾	1.7378 ⁽⁹⁾	2.8228 ⁽¹³⁾	4.8925 ⁽¹⁴⁾	2.4201 ⁽¹¹⁾	1.7203 ⁽⁷⁾	1.0493 ⁽³⁾	1.7433 ⁽¹⁰⁾	1.5175 ⁽⁴⁾	82.1950 ⁽¹⁵⁾	1.0382 ⁽²⁾	1.7190 ⁽⁶⁾	0.9352 ⁽¹⁾
		MRE ($\hat{\beta}$)	1.7267 ⁽¹²⁾	1.6853 ⁽⁹⁾	2.8503 ⁽¹⁴⁾	1.6835 ⁽⁸⁾	2.8523 ⁽¹³⁾	4.5091 ⁽¹³⁾	2.2405 ⁽¹³⁾	1.7187 ⁽¹⁰⁾	1.0178 ⁽³⁾	1.7212 ⁽¹¹⁾	1.6533 ⁽⁶⁾	0.9988 ⁽⁴⁾	0.9637 ⁽³⁾	1.6757 ⁽⁷⁾	0.8647 ⁽²⁾
D _{abs}		0.0319 ⁽¹⁾	0.0333 ⁽²⁾	0.0350 ⁽⁴⁾	0.0326 ⁽³⁾	0.0357 ⁽¹⁾	0.0377 ⁽¹⁴⁾	0.0340 ⁽¹⁾	0.0373 ⁽³⁾	0.0445 ⁽¹²⁾	0.0389 ⁽⁹⁾	0.0558 ⁽¹⁰⁾	0.4991 ⁽¹⁶⁾	0.0404 ⁽¹⁰⁾	0.0351 ⁽⁷⁾	0.0413 ⁽¹¹⁾	
D _{max}		0.0459 ⁽¹⁾	0.0481 ⁽²⁾	0.0505 ⁽⁴⁾	0.0469 ⁽³⁾	0.0505 ⁽⁵⁾	0.0631 ⁽¹⁴⁾	0.0490 ⁽¹⁾	0.0538 ⁽³⁾	0.0637 ⁽¹²⁾	0.0564 ⁽⁹⁾	0.0973 ⁽¹³⁾	0.9810 ⁽¹⁶⁾	0.0580 ⁽¹⁰⁾	0.0506 ⁽⁷⁾	0.0594 ⁽¹¹⁾	
$\sum$ Ranks		60 ⁽⁸⁾	48 ⁽³⁾	91 ⁽¹⁴⁾	55 ⁽⁶⁾	93 ⁽¹⁰⁾	73 ⁽¹⁰⁾	80 ⁽¹¹⁾	68 ⁽⁹⁾	48 ⁽³⁾	81 ⁽¹²⁾	56 ⁽⁷⁾	87 ⁽¹³⁾	35 ⁽²⁾	54 ⁽⁵⁾	31 ⁽¹⁾	
100		$ Bias(\hat{\theta}) $	0.2626 ⁽⁶⁾	0.2591 ⁽⁵⁾	0.4802 ⁽¹³⁾	0.2625 ⁽⁶⁾	0.4757 ⁽¹²⁾	0.7473 ⁽¹⁴⁾	0.3023 ⁽¹¹⁾	0.2621 ⁽⁷⁾	0.1649 ⁽³⁾	0.2696 ⁽¹⁰⁾	0.2430 ⁽⁴⁾	15.6263 ⁽¹⁵⁾	0.1526 ⁽²⁾	0.2609 ⁽⁶⁾	0.1347 ⁽¹⁾
		$ Bias(\hat{\beta}) $	1.3681 ⁽¹¹⁾	1.3414 ⁽⁹⁾	2.5198 ⁽¹⁶⁾	1.3408 ⁽⁸⁾	2.4912 ⁽¹⁴⁾	0.3920 ⁽¹⁾	1.5757 ⁽¹²⁾	1.3637 ⁽¹⁰⁾	0.8170 ⁽³⁾	1.3705 ⁽¹²⁾	1.3251 ⁽⁶⁾	0.7990 ⁽⁴⁾	0.7642 ⁽³⁾	1.3322 ⁽⁷⁾	0.6731 ⁽²⁾
		MSE ($\hat{\theta}$)	0.0698 ⁽¹⁾	0.0680 ⁽²⁾	0.2542 ⁽¹³⁾	0.0699 ⁽³⁾	0.2467 ⁽¹²⁾	0.5628 ⁽¹⁴⁾	0.0930 ⁽¹¹⁾	0.0697 ⁽³⁾	0.0280 ⁽³⁾	0.0743 ⁽¹⁰⁾	0.0621 ⁽⁴⁾	277.4996 ⁽¹⁵⁾	0.0242 ⁽²⁾	0.0692 ⁽⁶⁾	0.0189 ⁽¹⁾
		MSE ($\hat{\beta}$)	1.8726 ⁽¹¹⁾	1.7599 ⁽⁹⁾	7.2335 ⁽¹⁶⁾	1.7981 ⁽⁸⁾	6.9525 ⁽¹⁴⁾	1.7138 ⁽¹⁾	2.5431 ⁽¹³⁾	1.8600 ⁽¹⁰⁾	0.6685 ⁽³⁾	1.8792 ⁽¹²⁾	1.7559 ⁽⁶⁾	6.6384 ⁽⁴⁾	0.5852 ⁽³⁾	1.7756 ⁽⁷⁾	0.4547 ⁽²⁾
		MRE ($\hat{\theta}$)	1.7507 ⁽⁸⁾	1.7272 ⁽⁵⁾	3.2011 ⁽¹²⁾	1.7503 ⁽⁹⁾	3.1714 ⁽¹²⁾	4.9820 ⁽¹⁴⁾	2.0151 ⁽¹¹⁾	1.7474 ⁽⁷⁾	1.0994 ⁽³⁾	1.7971 ⁽¹⁰⁾	1.6201 ⁽⁴⁾	104.1753 ⁽¹⁵⁾	1.0176 ⁽²⁾	1.7393 ⁽⁶⁾	0.8978 ⁽¹⁾
		MRE ($\hat{\beta}$)	1.7102 ⁽¹¹⁾	1.6768 ⁽⁹⁾	3.1140 ⁽¹⁶⁾	1.6759 ⁽⁸⁾	3.1140 ⁽¹⁴⁾	4.9000 ⁽¹⁾	1.9697 ⁽¹³⁾	1.7046 ⁽¹⁰⁾	1.0213 ⁽³⁾	1.7132 ⁽¹¹⁾	1.6563 ⁽⁶⁾	0.9988 ⁽⁴⁾	0.9553 ⁽³⁾	1.6653 ⁽⁷⁾	0.8413 ⁽²⁾
	D _{abs}	0.0221 ⁽¹⁾	0.0219 ⁽¹⁾	0.0232 ⁽⁴⁾	0.0228 ⁽³⁾	0.0232 ⁽⁵⁾	0.0363 ⁽¹⁴⁾	0.0220 ⁽¹⁾	0.0240 ⁽³⁾	0.0302 ⁽¹²⁾	0.0290 ⁽⁹⁾	0.0391 ⁽¹⁰⁾	0.5058 ⁽¹⁶⁾	0.0297 ⁽¹⁰⁾	0.0264 ⁽⁶⁾	0.0297 ⁽¹¹⁾	
	D _{max}	0.0320 ⁽¹⁾	0.0318 ⁽¹⁾	0.0336 ⁽⁴⁾	0.0331 ⁽³⁾	0.0336 ⁽⁵⁾	0.0336 ⁽¹⁴⁾	0.0318 ⁽¹⁾	0.0347 ⁽³⁾	0.0438 ⁽¹²⁾	0.0419 ⁽⁹⁾	0.0575 ⁽¹³⁾	0.9904 ⁽¹⁶⁾	0.0430 ⁽¹⁰⁾	0.0381 ⁽⁸⁾	0.0431 ⁽¹¹⁾	
	$\sum$ Ranks	65 ⁽⁹⁾	44 ⁽³⁾	95 ⁽¹⁶⁾	57 ⁽⁶⁾	80 ⁽¹¹⁾	73 ⁽¹⁰⁾	76 ⁽¹¹⁾	65 ⁽⁹⁾	48 ⁽³⁾	81 ⁽¹²⁾	56 ⁽⁷⁾	87 ⁽¹³⁾	35 ⁽²⁾	55 ⁽⁵⁾	31 ⁽¹⁾	
	200	$ Bias(\hat{\theta}) $	0.2523 ⁽⁹⁾	0.2483 ⁽⁶⁾	0.4584 ⁽¹²⁾	0.2502 ⁽⁷⁾	0.4587 ⁽¹³⁾	0.7732 ⁽¹⁴⁾	0.2652 ⁽¹¹⁾	0.2512 ⁽⁸⁾	0.1539 ⁽³⁾	0.2549 ⁽¹⁰⁾	0.2327 ⁽⁴⁾	17.6594 ⁽¹⁵⁾	0.1418 ⁽²⁾	0.2467 ⁽⁶⁾	0.1228 ⁽¹⁾
		$ Bias(\hat{\beta}) $	1.3694 ⁽¹²⁾	1.3452 ⁽⁹⁾	2.4773 ⁽¹⁶⁾	1.3428 ⁽⁸⁾	2.4751 ⁽¹⁴⁾	0.4309 ⁽¹⁾	1.4409 ⁽¹³⁾	1.3673 ⁽¹¹⁾	0.8224 ⁽³⁾	1.3663 ⁽¹⁰⁾	1.3239 ⁽⁶⁾	0.7990 ⁽⁴⁾	0.7566 ⁽³⁾	1.3351 ⁽⁷⁾	0.6555 ⁽²⁾
		MSE ($\hat{\theta}$)	0.0640 ⁽⁶⁾	0.0622 ⁽²⁾	0.2196 ⁽¹²⁾	0.0630 ⁽³⁾	0.2200 ⁽¹³⁾	0.5985 ⁽¹⁴⁾	0.0709 ⁽¹¹⁾	0.0636 ⁽³⁾	0.0240 ⁽³⁾	0.0655 ⁽¹⁰⁾	0.0559 ⁽⁴⁾	333.2915 ⁽¹⁵⁾	0.0206 ⁽²⁾	0.0613 ⁽⁶⁾	0.0155 ⁽¹⁾
		MSE ($\hat{\beta}$)	1.8756 ⁽¹²⁾	1.8096 ⁽⁹⁾	6.3765 ⁽¹⁶⁾	1.8036 ⁽⁸⁾	6.3688 ⁽¹⁴⁾	2.0921 ⁽¹³⁾	2.0921 ⁽¹³⁾	1.8698 ⁽¹¹⁾	0.6769 ⁽³⁾	1.8674 ⁽¹⁰⁾	1.7527 ⁽⁶⁾	6.6384 ⁽⁴⁾	0.5740 ⁽³⁾	1.7831 ⁽⁷⁾	0.4309 ⁽²⁾
		MRE ($\hat{\theta}$)	1.6817 ⁽⁹⁾	1.6551 ⁽⁶⁾	3.0558 ⁽¹²⁾	1.6683 ⁽⁷⁾	3.0579 ⁽¹³⁾	5.1545 ⁽¹⁴⁾	1.7678 ⁽¹¹⁾	1.6749 ⁽⁸⁾	1.0257 ⁽³⁾	1.6991 ⁽¹⁰⁾	1.5513 ⁽⁴⁾	117.7293 ⁽¹⁵⁾	0.9454 ⁽²⁾	1.6446 ⁽⁶⁾	0.8189 ⁽¹⁾
		MRE ($\hat{\beta}$)	1.7118 ⁽¹²⁾	1.6814 ⁽⁹⁾	3.0966 ⁽¹⁶⁾	1.6785 ⁽⁸⁾	3.0938 ⁽¹⁴⁾	5.3877 ⁽¹⁴⁾	1.8011 ⁽¹³⁾	1.7092 ⁽¹¹⁾	1.0280 ⁽³⁾	1.7079 ⁽¹⁰⁾	1.6548 ⁽⁶⁾	0.9988 ⁽⁴⁾	0.9458 ⁽³⁾	1.6689 ⁽⁷⁾	0.8194 ⁽²⁾
D _{abs}		0.0174 ⁽¹⁾	0.0173 ⁽¹⁾	0.0177 ⁽⁴⁾	0.0168 ⁽³⁾	0.0175 ⁽⁵⁾	0.0304 ⁽¹⁴⁾	0.0172 ⁽¹⁾	0.0194 ⁽³⁾	0.0199 ⁽¹⁰⁾	0.0193 ⁽⁸⁾	0.0331 ⁽¹⁰⁾	0.4956 ⁽¹⁶⁾	0.0239 ⁽¹¹⁾	0.0182 ⁽⁷⁾	0.0240 ⁽¹²⁾	
D _{max}		0.0251 ⁽¹⁾	0.0251 ⁽¹⁾	0.0256 ⁽⁴⁾	0.0243 ⁽³⁾	0.0253 ⁽⁵⁾	0.0406 ⁽¹⁴⁾	0.0248 ⁽¹⁾	0.0248 ⁽²⁾	0.0287 ⁽¹⁰⁾	0.0278 ⁽⁹⁾	0.0495 ⁽¹³⁾	0.9951 ⁽¹⁶⁾	0.0346 ⁽¹¹⁾	0.0263 ⁽⁷⁾	0.0347 ⁽¹²⁾	
$\sum$ Ranks		71 ⁽¹⁵⁾	51 ⁽⁶⁾	93 ⁽¹⁶⁾	47 ⁽⁴⁾	91 ⁽¹⁴⁾	73 ⁽¹⁰⁾	76 ⁽¹¹⁾	75 ⁽¹⁰⁾	44 ⁽³⁾	76 ⁽¹¹⁾	56 ⁽⁷⁾	87 ⁽¹³⁾	37 ⁽²⁾	50 ⁽⁵⁾	33 ⁽¹⁾	
300		$ Bias(\hat{\theta}) $	0.2548 ⁽⁶⁾	0.2516 ⁽⁷⁾	0.4964 ⁽¹²⁾	0.2506 ⁽⁶⁾	0.4966 ⁽¹³⁾	0.7893 ⁽¹⁴⁾	0.2593 ⁽¹¹⁾	0.2563 ⁽⁹⁾	0.1553 ⁽³⁾	0.2578 ⁽¹⁰⁾	0.2378 ⁽⁴⁾	18.6238 ⁽¹⁵⁾	0.1436 ⁽²⁾	0.2493 ⁽⁶⁾	0.1236 ⁽¹⁾
		$ Bias(\hat{\beta}) $	1.3675 ⁽¹¹⁾	1.3422 ⁽⁹⁾	2.6440 ⁽¹⁶⁾	1.3369 ⁽⁸⁾	2.6353 ⁽¹⁴⁾	0.4651 ⁽¹⁾	1.3847 ⁽¹³⁾	1.3642 ⁽¹⁰⁾	0.8292 ⁽³⁾	1.3707 ⁽¹²⁾	1.3242 ⁽⁶⁾	0.7990 ⁽⁴⁾	0.7520 ⁽³⁾	1.3303 ⁽⁷⁾	0.6469 ⁽²⁾
		MSE ($\hat{\theta}$)	0.0652 ⁽⁸⁾	0.0637 ⁽⁷⁾	0.2565 ⁽¹²⁾	0.0630 ⁽³⁾	0.2567 ⁽¹³⁾	0									

Table 2: Partial and Overall Ranks for all Estimation Methods of the BE Distribution from 1

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLSE	LTADE	MSADE	MSALDE	ADSOE	KE	MSSDE	MSSLDE	MSLNDE
15	$ Bias (\hat{\theta})$	7.0	9.0	5.0	11.0	6.0	14.0	13.0	8.0	3.0	12.0	4.0	15.0	2.0	10.0	1.0
	$ Bias (\hat{\beta})$	13.0	11.0	7.0	10.0	6.0	1.0	15.0	12.0	5.0	14.0	8.0	4.0	3.0	9.0	2.0
	MSE ($\hat{\theta}$)	7.0	8.0	5.0	11.0	4.0	14.0	13.0	9.0	3.0	12.0	6.0	15.0	2.0	10.0	1.0
	MSE ($\hat{\beta}$)	13.0	11.0	7.0	10.0	6.0	1.0	15.0	12.0	5.0	14.0	8.0	4.0	3.0	9.0	2.0
	MRE ($\hat{\theta}$)	7.0	9.0	5.0	11.0	6.0	14.0	13.0	8.0	3.0	12.0	4.0	15.0	2.0	10.0	1.0
	MRE ($\hat{\beta}$)	13.0	11.0	7.0	10.0	6.0	1.0	15.0	12.0	5.0	14.0	8.0	4.0	3.0	9.0	2.0
	D_{abs}	1.0	2.0	4.0	6.0	3.0	14.0	5.0	8.0	13.0	10.0	12.0	15.0	9.0	7.0	11.0
	D_{max}	1.0	2.0	4.0	6.0	3.0	14.0	5.0	8.0	13.0	10.0	12.0	15.0	9.0	7.0	11.0
	$\sum Ranks$	62 ⁽⁶⁾	63 ⁽⁷⁾	44 ⁽⁴⁾	75 ⁽¹¹⁾	40 ⁽³⁾	73 ⁽¹⁰⁾	94 ⁽¹⁴⁾	77 ⁽¹²⁾	49 ⁽⁵⁾	97 ⁽¹⁶⁾	63 ⁽⁷⁾	87 ⁽¹³⁾	34 ⁽²⁾	71 ⁽⁹⁾	31 ⁽¹⁾
	$ Bias (\hat{\theta})$	8.0	5.0	12.0	9.0	13.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	6.0	1.0
	$ Bias (\hat{\beta})$	12.0	9.0	14.0	8.0	15.0	1.0	13.0	10.0	5.0	11.0	6.0	4.0	3.0	7.0	2.0
	MSE ($\hat{\theta}$)	6.0	5.0	12.0	9.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	7.0	1.0
	MSE ($\hat{\beta}$)	12.0	9.0	14.0	8.0	15.0	1.0	13.0	10.0	5.0	11.0	6.0	4.0	3.0	7.0	2.0
	MRE ($\hat{\theta}$)	8.0	5.0	12.0	9.0	13.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	6.0	1.0
	MRE ($\hat{\beta}$)	12.0	9.0	14.0	8.0	15.0	1.0	13.0	10.0	5.0	11.0	6.0	4.0	3.0	7.0	2.0
50	D_{abs}	1.0	3.0	6.0	2.0	5.0	14.0	4.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	D_{max}	1.0	3.0	6.0	2.0	5.0	14.0	4.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	$\sum Ranks$	60 ⁽⁸⁾	48 ⁽³⁾	91 ⁽¹⁴⁾	55 ⁽⁶⁾	93 ⁽¹⁵⁾	73 ⁽¹⁰⁾	80 ⁽¹¹⁾	68 ⁽⁹⁾	48 ⁽³⁾	81 ⁽¹²⁾	56 ⁽⁷⁾	87 ⁽¹³⁾	35 ⁽²⁾	54 ⁽⁵⁾	31 ⁽¹⁾
	$ Bias (\hat{\theta})$	9.0	5.0	13.0	8.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	6.0	1.0
	$ Bias (\hat{\beta})$	11.0	9.0	15.0	8.0	14.0	1.0	13.0	10.0	5.0	12.0	6.0	4.0	3.0	7.0	2.0
	MSE ($\hat{\theta}$)	8.0	5.0	13.0	9.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	6.0	1.0
	MSE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	1.0	13.0	10.0	5.0	12.0	6.0	4.0	3.0	7.0	2.0
	MRE ($\hat{\theta}$)	9.0	5.0	13.0	8.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	6.0	1.0
	MRE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	1.0	13.0	10.0	5.0	12.0	6.0	4.0	3.0	7.0	2.0
	D_{abs}	3.0	1.0	6.0	4.0	5.0	14.0	2.0	7.0	12.0	9.0	13.0	15.0	10.0	8.0	11.0
	D_{max}	3.0	1.0	5.0	4.0	6.0	14.0	2.0	7.0	12.0	9.0	13.0	15.0	10.0	8.0	11.0
	$\sum Ranks$	65 ⁽⁶⁾	44 ⁽³⁾	95 ⁽¹⁵⁾	57 ⁽⁷⁾	89 ⁽¹⁴⁾	73 ⁽¹⁰⁾	76 ⁽¹¹⁾	65 ⁽⁸⁾	48 ⁽⁴⁾	84 ⁽¹²⁾	56 ⁽⁶⁾	87 ⁽¹³⁾	35 ⁽²⁾	55 ⁽⁵⁾	31 ⁽¹⁾
	$ Bias (\hat{\theta})$	9.0	6.0	12.0	7.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias (\hat{\beta})$	12.0	9.0	15.0	8.0	14.0	1.0	13.0	11.0	5.0	10.0	6.0	4.0	3.0	7.0	2.0
	MSE ($\hat{\theta}$)	9.0	6.0	12.0	7.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
MSE ($\hat{\beta}$)	12.0	9.0	15.0	8.0	14.0	1.0	13.0	11.0	5.0	10.0	6.0	4.0	3.0	7.0	2.0	
200	MRE ($\hat{\theta}$)	9.0	6.0	12.0	7.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	12.0	9.0	15.0	8.0	14.0	1.0	13.0	11.0	5.0	10.0	6.0	4.0	3.0	7.0	2.0
	D_{abs}	4.0	3.0	6.0	1.0	5.0	14.0	2.0	9.0	10.0	8.0	13.0	15.0	11.0	7.0	12.0
	D_{max}	4.0	3.0	6.0	1.0	5.0	14.0	2.0	9.0	10.0	8.0	13.0	15.0	11.0	7.0	12.0
	$\sum Ranks$	71 ⁽⁶⁾	51 ⁽⁶⁾	93 ⁽¹⁵⁾	47 ⁽⁴⁾	91 ⁽¹⁴⁾	73 ⁽⁹⁾	76 ⁽¹¹⁾	75 ⁽¹⁰⁾	44 ⁽³⁾	76 ⁽¹¹⁾	56 ⁽⁷⁾	87 ⁽¹³⁾	37 ⁽²⁾	50 ⁽⁵⁾	33 ⁽¹⁾
	$ Bias (\hat{\theta})$	8.0	7.0	12.0	6.0	13.0	14.0	11.0	9.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias (\hat{\beta})$	11.0	9.0	15.0	8.0	14.0	1.0	13.0	10.0	5.0	12.0	6.0	4.0	3.0	7.0	2.0
	MSE ($\hat{\theta}$)	8.0	7.0	12.0	6.0	13.0	14.0	11.0	9.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MSE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	1.0	13.0	10.0	5.0	12.0	6.0	4.0	3.0	7.0	2.0
	MRE ($\hat{\theta}$)	8.0	7.0	12.0	6.0	13.0	14.0	11.0	9.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	1.0	13.0	11.0	5.0	10.0	6.0	4.0	3.0	7.0	2.0
	D_{abs}	2.0	3.0	6.0	1.0	5.0	14.0	4.0	7.0	11.0	12.0	13.0	15.0	9.0	8.0	10.0
	D_{max}	2.0	3.0	6.0	1.0	5.0	14.0	4.0	7.0	11.0	12.0	13.0	15.0	9.0	8.0	10.0
	$\sum Ranks$	61 ⁽⁶⁾	54 ⁽⁶⁾	93 ⁽¹⁵⁾	44 ⁽³⁾	91 ⁽¹⁴⁾	73 ⁽¹⁰⁾	80 ⁽¹¹⁾	71 ⁽⁹⁾	46 ⁽⁴⁾	90 ⁽¹³⁾	56 ⁽⁷⁾	87 ⁽¹²⁾	33 ⁽²⁾	52 ⁽⁵⁾	29 ⁽¹⁾
	$ Bias (\hat{\theta})$	10.0	6.0	12.0	7.0	13.0	14.0	11.0	9.0	3.0	11.0	4.0	15.0	2.0	5.0	1.0
$ Bias (\hat{\beta})$	11.0	9.0	14.0	8.0	15.0	1.0	12.0	8.0	3.0	13.0	7.0	4.0	3.0	6.0	2.0	
MSE ($\hat{\theta}$)	9.0	7.0	12.0	6.0	13.0	14.0	11.0	9.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0	
MSE ($\hat{\beta}$)	11.0	9.0	14.0	8.0	15.0	1.0	12.0	8.0	3.0	13.0	7.0	4.0	3.0	6.0	2.0	
500	MRE ($\hat{\theta}$)	10.0	6.0	12.0	7.0	13.0	14.0	11.0	9.0	3.0	11.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	11.0	9.0	14.0	8.0	15.0	1.0	12.0	8.0	3.0	13.0	7.0	4.0	3.0	6.0	2.0
	D_{abs}	1.0	3.0	6.0	2.0	5.0	14.0	4.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	D_{max}	1.0	3.0	6.0	2.0	5.0	14.0	4.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	$\sum Ranks$	64 ⁽⁶⁾	52 ⁽⁶⁾	90 ⁽¹³⁾	48 ⁽⁴⁾	94 ⁽¹⁵⁾	73 ⁽¹¹⁾	72 ⁽¹⁰⁾	70 ⁽⁹⁾	48 ⁽⁴⁾	90 ⁽¹³⁾	59 ⁽⁷⁾	87 ⁽¹²⁾	35 ⁽²⁾	47 ⁽³⁾	31 ⁽¹⁾

$$\left\{ \begin{array}{ll} \text{Bias :} & |Bias(\hat{\theta})| = \frac{1}{D} \sum_{i=1}^D |\hat{\theta} - \theta|, \\ \text{MeanSquaredErrors :} & MSE = \frac{1}{D} \sum_{i=1}^D (\hat{\theta} - \theta)^2, \\ \text{MeanRelativeErrors :} & MRE = \frac{1}{D} \sum_{i=1}^D |\hat{\theta} - \theta| / \theta, \\ \text{AverageAbsoluteDifference :} & D_{abs} = \frac{1}{nD} \sum_{i=1}^D \sum_{j=1}^n |G(x_{ij}|\theta) - G(x_{ij}|\hat{\theta})|, \\ \text{MaximumAbsoluteDifference :} & D_{max} = \frac{1}{D} \sum_{i=1}^D \max_j |G(x_{ij}|\theta) - G(x_{ij}|\hat{\theta})|, \end{array} \right.$$

where the parameter vector is $\Theta = (\theta, \alpha, \beta)$.

The partial and overall ranks for the fifteen estimation methods of the BE distribution, summarized in Table 2, reveal a clear hierarchy of performance across varying sample sizes n . The overall ranking, determined by the sum of ranks ($\sum Ranks$), consistently indicates that the MSLNDE method is the best performing estimator, achieving the first overall rank (rank 1) across all considered sample sizes, from $n=15$ to $n=500$. Following closely behind, the MSSDE method secures the second overall rank (rank 2) throughout the simulation study. This stability in the top two positions suggests that these methods are generally robust to changes in sample size. Further examination of the partial ranks shows that the MSLNDE method's superior overall performance is primarily driven by its effectiveness in estimating the parameter θ , where it achieves a rank of 1.0 for all bias, mean squared error, and mean relative error metrics ($|Bias|(\hat{\theta})$, $MSE(\hat{\theta})$, and $MRE(\hat{\theta})$) across all sample sizes. For the parameter β , MSLNDE consistently achieves a rank of 2.0. Similarly, the MSSDE method ranks 2.0 for all $\hat{\theta}$ metrics and 3.0 for all $\hat{\beta}$ metrics, reinforcing its second-place overall standing.

Conversely, methods like MSALDE and CVME generally exhibit poor performance, with MSALDE consistently placing near the bottom (rank 15), and CVME also ranking poorly for larger sample sizes ($n \geq 50$). The RTADE method is notably poor in estimating the parameter β , consistently ranking 1.0 for $\hat{\beta}$ metrics, but performs poorly for $\hat{\theta}$ metrics, leading to an overall mid-to-low rank. A different trend is observed for the goodness-of-fit metrics, D_{abs} and D_{max} , which measure distributional distance. Here, the best overall methods (MSLNDE and MSSDE) do not dominate. Instead, the MLE and the ADE frequently achieve the lowest ranks (1.0 to 4.0) for these metrics, especially as the sample size increases, suggesting they provide a better fit to the empirical distribution function despite having less favorable performance on the parameter-specific metrics. The overall ranking,

being a composite measure, thus prioritizes the accuracy of the parameter estimates (bias and error) over the fit statistics for this particular study. These inference from Table 2 are due to the simulation results contained in Table 1 and the plots which provide visual illustrations are in Figures 3, 4, 5 and 6. A visual summary of the performance of the estimation methods, ordered by rank, call heatmap for case 1 is in Figure 7.

The comprehensive ranking of estimation methods for the BE distribution, presented in Table 4, provides a clear assessment of their performance under different sample sizes, n . The overall performance, determined by the sum of ranks ($\sum Ranks$), indicates that the MSLNDE method is the best estimator, consistently achieving the first overall rank (rank 1) for all sample sizes from $n=50$ to $n=500$. For the smallest sample size, $n=15$, MSADDE performs the best overall (rank 1), though MSLNDE is a close second (rank 2). As n increases, MSLNDE, MSSDE, and MSADDE solidify their positions as the top three overall performers. Upon closer inspection of the partial ranks, MSLNDE's excellent performance stems from its effectiveness in estimating the θ parameter, for which it maintains a rank of 1.0 across all bias, mean squared error, and mean relative error metrics ($|Bias|(\hat{\theta})$, $MSE(\hat{\theta})$, and

$MRE(\hat{\theta})$) for all sample sizes. For the β parameter, the performance is less consistent, with the KE method dominating by achieving a rank of 1.0 for all $\hat{\beta}$ metrics for $n \geq 50$. The MSLNDE method's rank for $\hat{\beta}$ is 3.0 for all $n \geq 50$. The MSADDE method exhibits strong performance at the smallest sample size, $n=15$, where it achieves the best overall rank and ranks second or third for all $\hat{\theta}$ metrics. This suggests MSADDE is particularly well-suited for very small samples. However, as n increases, its overall rank generally falls, though it remains in the top five. The MSSDE method shows stable, strong performance, typically ranking second or fourth overall, and consistently ranks 2.0 or 4.0 for all $\hat{\theta}$ and $\hat{\beta}$ estimation metrics, respectively, for $n \geq 100$.

In contrast, several methods consistently perform poorly. The CVME and OLSE methods frequently fall into the lowest overall ranks, particularly as the sample size grows. Additionally, the WLSE, LTADDE, and KE methods generally perform poorly overall, with the exception that the KE method is highly accurate for the β parameter estimates, achieving rank 1.0 for all $\hat{\beta}$ metrics for $n \geq 50$. The RTADE method is consistently poor across all metrics and sample sizes. A noticeable pattern emerges for the goodness-of-fit metrics, D_{abs} and D_{max} , which measure distributional distance. Here,

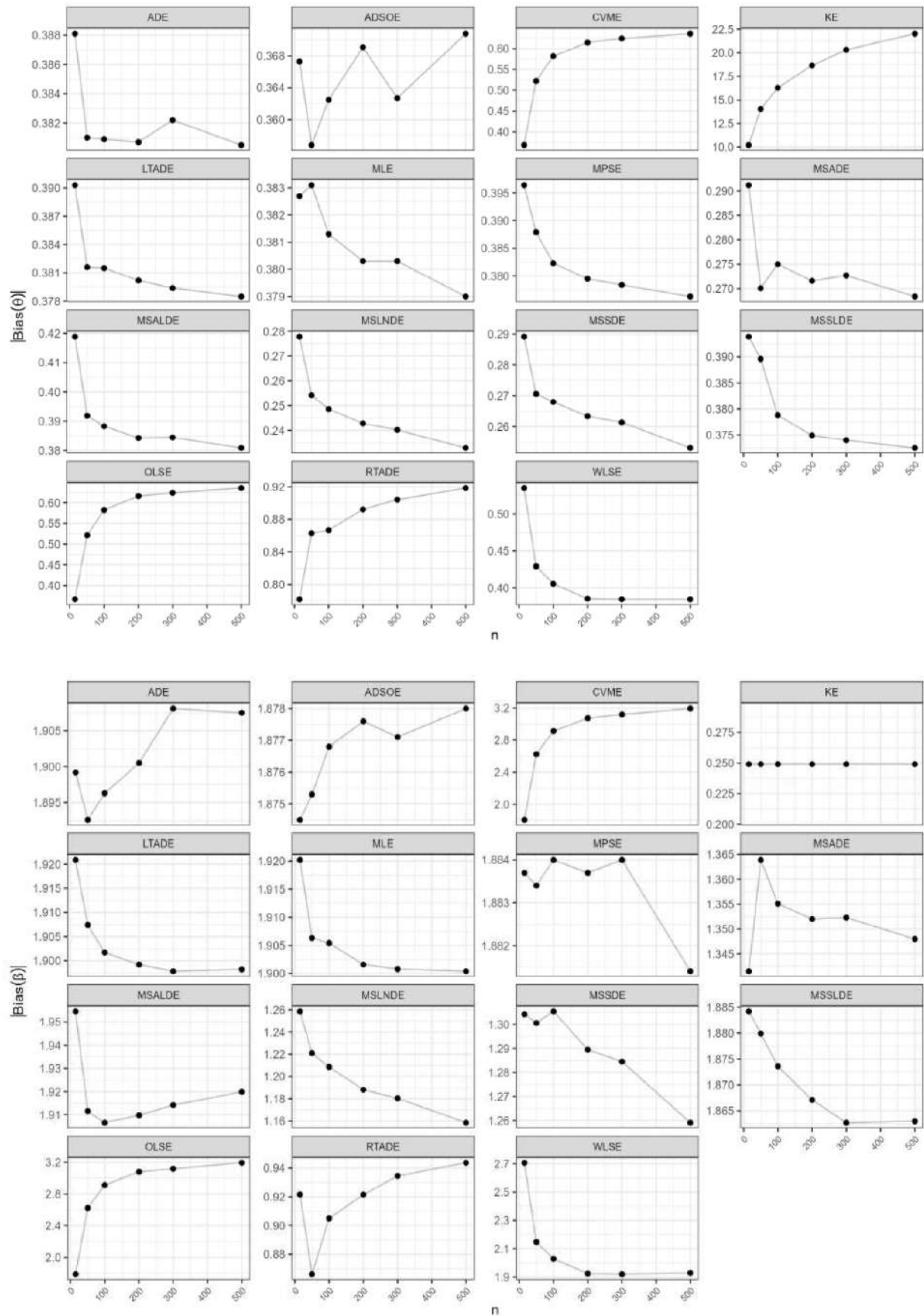


Figure 8: Graphical representations for the Bias values presented in Table 3.

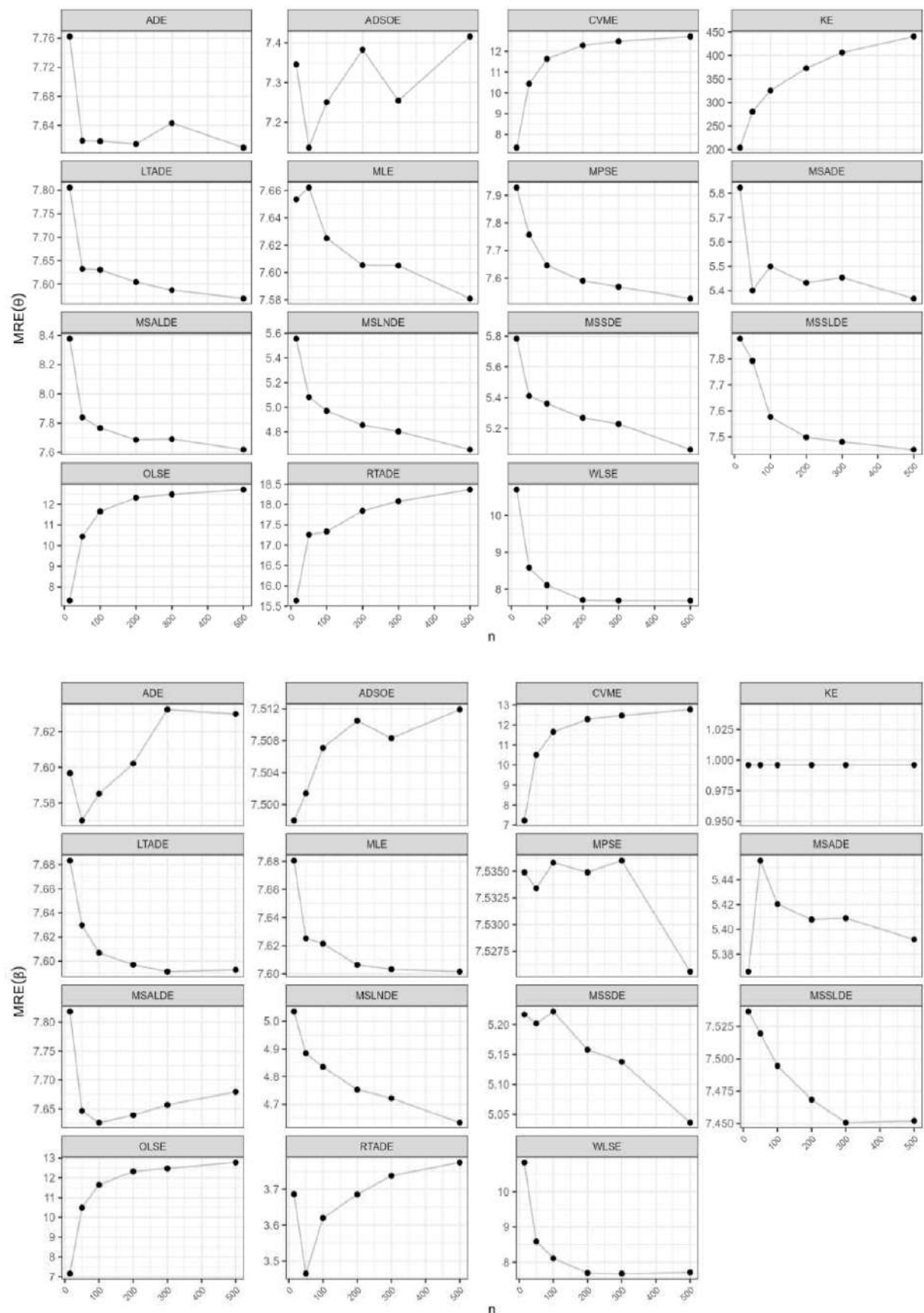


Figure 9: Graphical representations for the MRE values presented in Table 3.

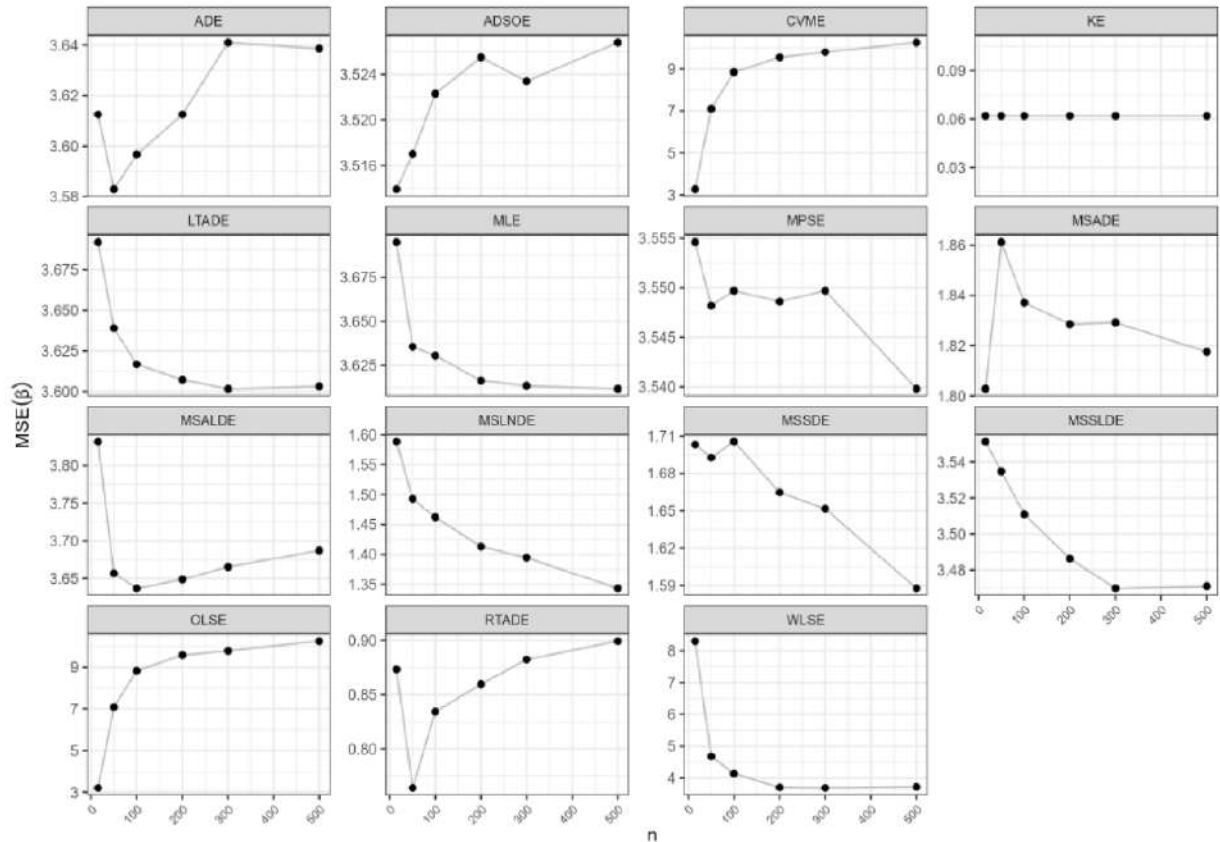
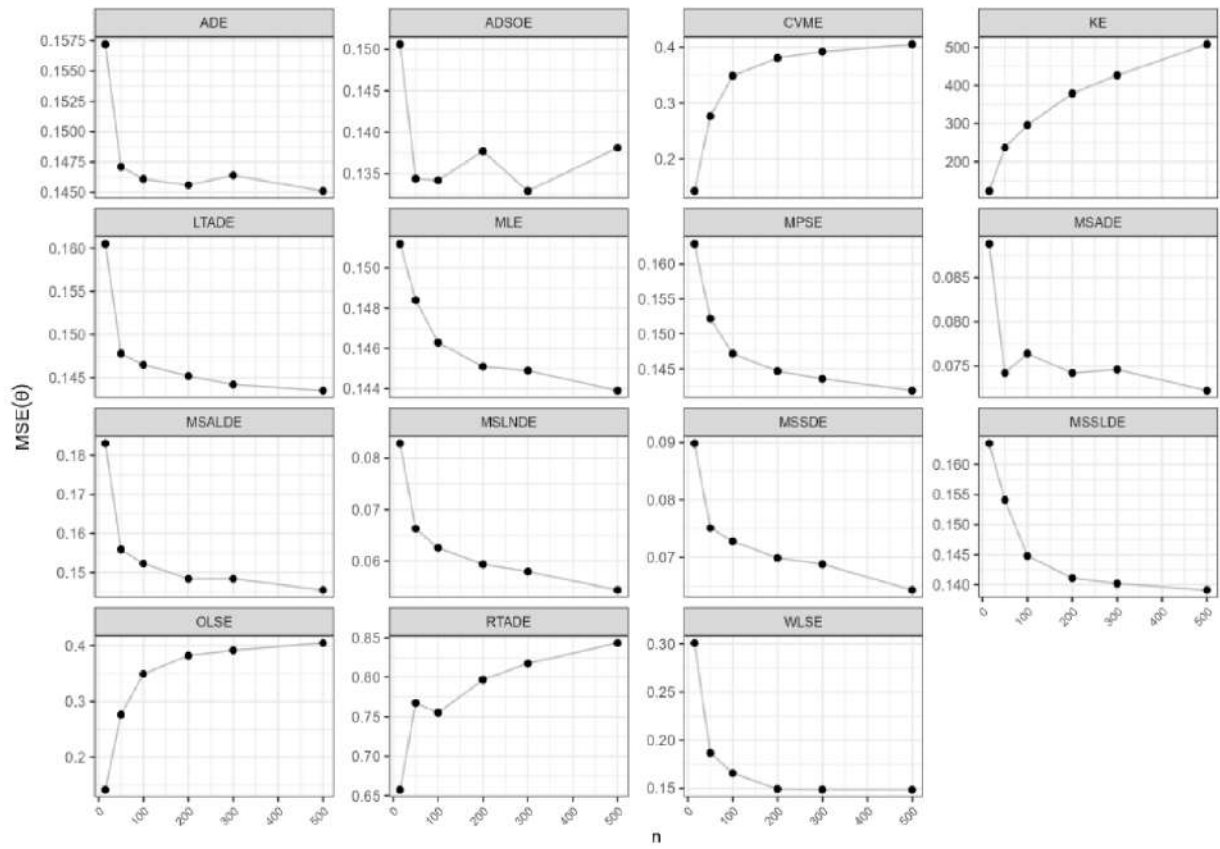


Figure 10: Graphical representations for the MSE values presented in Table 3.

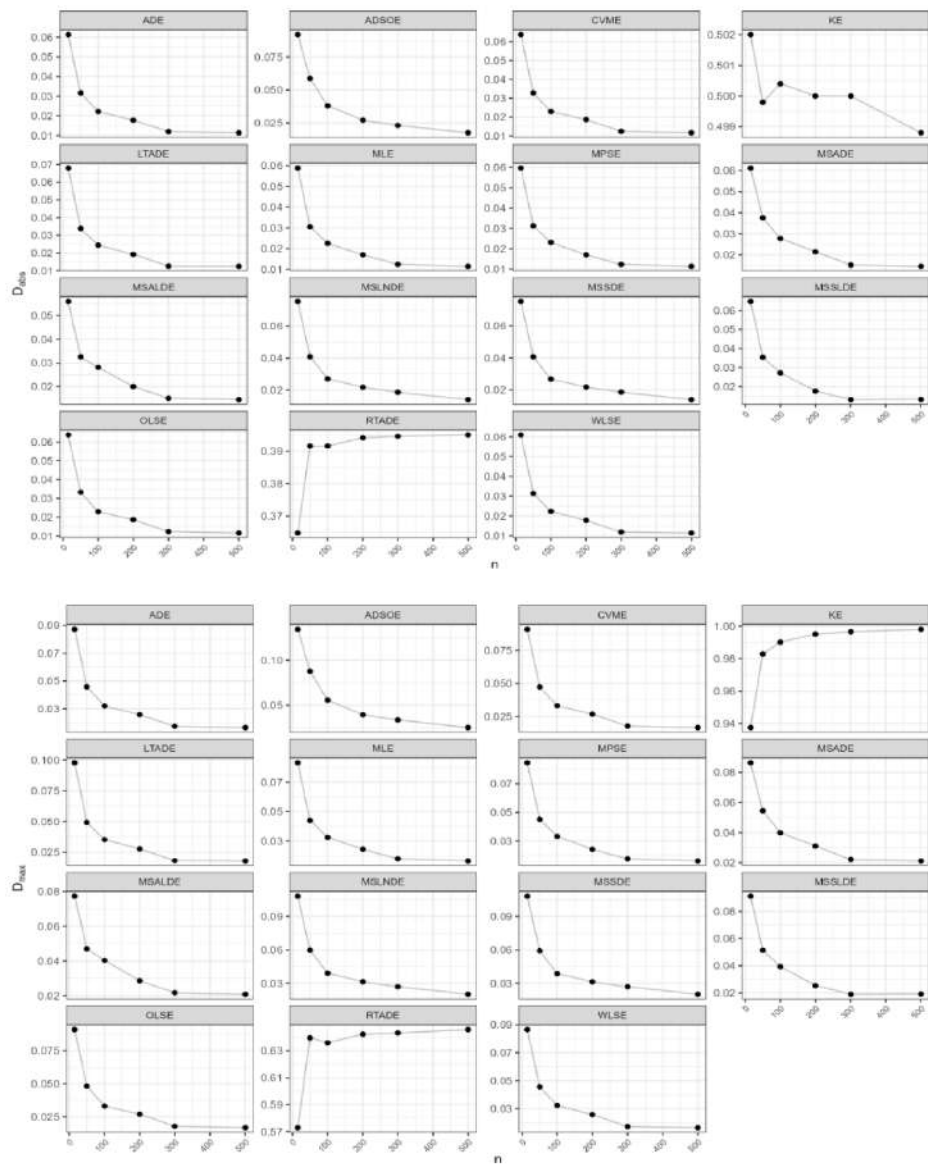


Figure 11: Graphical representations for the D_{abs} and D_{max} values presented in Table 3.

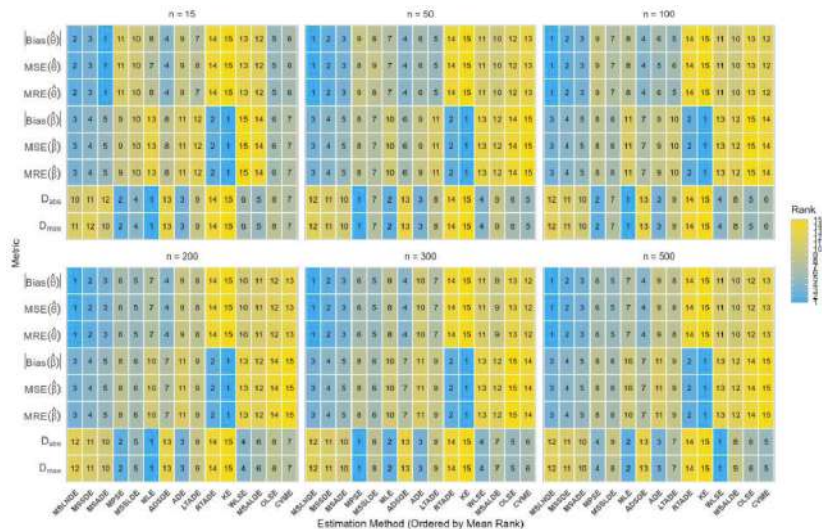


Figure 12: Partial Rank Heatmap for simulation in Table 3.

Table 3: Numerical Values of Simulation Measures for $\theta = 0.05, \beta = 0.25$.

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLSE	LTADE	MSADE	MSALDE	ADSOE	KE	MSSE	MSLSE	MSLNDE
15	$Bias(\hat{\theta})$	0.3827(7)	0.3881(6)	0.3687(6)	0.3667(9)	0.3667(9)	0.7818(14)	0.5352(13)	0.3903(9)	0.2912(3)	0.4189(12)	0.3673(5)	0.2032(15)	0.2892(2)	0.3939(10)	0.2779(1)
	$Bias(\hat{\beta})$	1.9202(12)	1.8902(11)	1.8060(7)	1.8837(9)	1.7889(9)	0.9216(2)	2.7069(16)	1.9209(12)	1.3415(6)	1.9546(1)	1.8745(6)	0.2490(1)	1.3042(4)	1.8842(10)	1.2586(3)
	MSE ($\hat{\theta}$)	0.1512(7)	0.1572(6)	0.1629(10)	0.1629(10)	0.1629(10)	0.6142(4)	0.3010(13)	0.1605(5)	0.0888(2)	0.1830(12)	0.1506(6)	0.1227(7)	0.0898(3)	0.1635(11)	0.0898(3)
	MSE ($\hat{\beta}$)	3.6951(13)	3.6125(11)	3.2693(9)	3.5546(10)	3.2069(9)	8.8734(24)	3.8303(16)	1.8023(6)	3.8318(14)	3.5139(8)	0.0620(1)	1.7033(4)	3.5514(9)	1.5886(3)	1.5886(3)
	MRE ($\hat{\theta}$)	7.6536(7)	7.7620(6)	7.3731(6)	7.9284(11)	7.3343(4)	15.6357(14)	10.7040(13)	5.8233(5)	8.3780(12)	7.3460(6)	204.0646(16)	5.7841(2)	7.8776(10)	5.5571(1)	5.5571(1)
	MRE ($\hat{\beta}$)	7.6906(12)	7.5968(11)	7.2241(7)	7.5958(9)	7.1558(6)	3.6848(2)	10.8275(13)	7.6835(13)	7.8182(14)	7.4980(6)	0.9960(1)	5.2169(4)	7.5365(10)	5.0346(3)	5.0346(3)
	D_{abs}	0.0588(2)	0.0613(6)	0.0637(7)	0.0597(8)	0.0638(9)	0.3648(14)	0.0610(9)	0.0680(10)	0.0559(1)	0.0919(12)	0.5020(15)	0.0752(11)	0.0647(9)	0.0752(11)	0.0752(11)
	D_{max}	0.0831(2)	0.0872(6)	0.0908(7)	0.0844(8)	0.0910(8)	0.5730(14)	0.0865(9)	0.0978(10)	0.0773(1)	0.1342(3)	0.9375(15)	0.1082(12)	0.0911(9)	0.1081(11)	0.1081(11)
	$\sum Ranks$	62(92)	69(92)	66(92)	66(92)	46(92)	76(70)	93(92)	32(92)	80(92)	66(92)	78(78)	42(92)	78(78)	35(92)	35(92)
	$Bias(\hat{\theta})$	0.3831(7)	0.3810(6)	0.3879(8)	0.3879(8)	0.5217(12)	0.8630(14)	0.4293(11)	0.2701(2)	0.3919(10)	0.3568(4)	0.2706(3)	0.3896(9)	0.2706(3)	0.2542(1)	0.2542(1)
	$Bias(\hat{\beta})$	1.8926(6)	1.8926(6)	2.6246(13)	1.8834(8)	2.6231(14)	0.8662(2)	2.1474(18)	1.9074(11)	1.3639(6)	1.9116(12)	1.8733(6)	1.3055(4)	1.8799(7)	1.2210(3)	1.2210(3)
	MSE ($\hat{\theta}$)	0.1484(7)	0.1471(6)	0.2766(13)	0.1522(8)	0.2762(12)	0.7676(14)	0.1866(10)	0.1478(6)	0.0742(2)	0.1559(10)	0.1344(4)	0.0751(3)	0.1541(9)	0.0653(3)	0.0653(3)
	MSE ($\hat{\beta}$)	3.6356(10)	3.2829(6)	7.0960(13)	3.5482(8)	7.0853(14)	17.2604(14)	4.6826(12)	1.8612(6)	3.6568(12)	3.5170(6)	0.0620(1)	1.6927(4)	3.5348(7)	1.4929(3)	1.4929(3)
	MRE ($\hat{\theta}$)	7.6622(7)	7.6191(6)	10.4374(13)	7.7576(8)	10.4334(14)	17.2604(14)	8.5855(12)	5.4013(6)	7.6323(6)	7.1364(4)	280.6091(16)	5.4114(3)	7.7924(9)	5.0834(1)	5.0834(1)
	MRE ($\hat{\beta}$)	7.6252(10)	7.5702(9)	10.4985(13)	7.5334(8)	10.4925(14)	3.4648(2)	8.5897(12)	7.6297(11)	7.6465(12)	7.5014(6)	0.9960(1)	5.2019(4)	7.5197(7)	4.8838(3)	4.8838(3)
D_{abs}	0.0305(1)	0.0317(6)	0.0328(6)	0.0313(2)	0.0333(7)	0.3916(14)	0.0314(9)	0.0339(8)	0.0326(5)	0.0586(12)	0.0407(11)	0.0355(9)	0.0409(12)	0.0409(12)	0.0409(12)	
D_{max}	0.0441(1)	0.0458(4)	0.0474(6)	0.0453(2)	0.0481(7)	0.6396(14)	0.0455(9)	0.0492(8)	0.0473(1)	0.0877(12)	0.9830(15)	0.0592(11)	0.0513(9)	0.0596(12)	0.0596(12)	
$\sum Ranks$	53(92)	50(92)	52(92)	52(92)	92(92)	76(76)	78(78)	41(92)	76(76)	56(92)	78(78)	43(92)	66(92)	36(92)	36(92)	
50	$Bias(\hat{\theta})$	0.3813(7)	0.3809(6)	0.5816(12)	0.3823(9)	0.5824(12)	0.8667(14)	0.4057(11)	0.3815(8)	0.2750(3)	0.3883(10)	0.3625(4)	0.2680(2)	0.3788(5)	0.2486(1)	0.2486(1)
	$Bias(\hat{\beta})$	1.9054(11)	1.8903(6)	2.9121(13)	1.8840(8)	2.9118(14)	0.9050(2)	2.0285(12)	1.9017(10)	1.3551(6)	1.9066(12)	1.8768(7)	1.3054(4)	1.8736(6)	1.2086(3)	1.2086(3)
	MSE ($\hat{\theta}$)	0.1463(7)	0.1461(6)	0.3487(12)	0.1472(9)	0.3492(12)	0.7552(14)	0.1655(11)	0.1465(8)	0.0764(6)	0.1523(10)	0.1342(4)	0.0728(2)	0.1445(9)	0.0626(1)	0.0626(1)
	MSE ($\hat{\beta}$)	3.5967(9)	3.8438(13)	3.5497(8)	8.8326(14)	8.8326(14)	0.8342(2)	4.1409(12)	3.6168(10)	1.8371(5)	3.6367(12)	3.5223(7)	1.7056(4)	3.5109(6)	1.4622(3)	1.4622(3)
	MRE ($\hat{\theta}$)	7.6251(7)	7.6185(6)	11.6317(12)	7.6461(9)	11.6484(13)	17.3348(14)	8.1135(11)	7.6305(8)	5.4997(5)	7.7657(12)	7.2504(7)	325.4170(16)	5.3699(2)	7.5769(4)	4.9716(1)
	MRE ($\hat{\beta}$)	7.5852(6)	11.6486(13)	7.5358(8)	11.6474(14)	11.6474(14)	3.6200(2)	8.1138(11)	7.6088(10)	5.4204(6)	7.6263(12)	7.5014(7)	5.2218(4)	7.4945(6)	4.8346(3)	4.8346(3)
	D_{abs}	0.0225(3)	0.0222(6)	0.0229(6)	0.0231(6)	0.0229(6)	0.3916(14)	0.0223(9)	0.0244(7)	0.0278(11)	0.0281(12)	0.0381(13)	0.0267(8)	0.0272(10)	0.0270(9)	0.0270(9)
	D_{max}	0.0325(3)	0.0320(6)	0.0332(6)	0.0334(6)	0.0331(6)	0.6360(14)	0.0323(9)	0.0352(7)	0.0399(11)	0.0404(12)	0.0557(12)	0.0386(6)	0.0393(10)	0.0390(9)	0.0390(9)
	$\sum Ranks$	60(92)	47(92)	91(92)	63(92)	89(92)	76(76)	76(76)	68(92)	46(92)	90(92)	59(92)	34(92)	53(92)	30(92)	30(92)
	$Bias(\hat{\theta})$	0.3803(8)	0.3807(6)	0.6142(12)	0.3705(6)	0.6160(12)	0.8922(14)	0.3853(11)	0.3802(7)	0.2716(6)	0.3843(10)	0.3691(4)	0.2634(2)	0.3740(5)	0.2428(1)	0.2428(1)
	$Bias(\hat{\beta})$	1.9016(11)	1.9005(10)	3.0720(13)	1.8837(8)	3.0797(13)	0.9215(2)	1.9241(12)	1.8992(9)	1.3520(5)	1.9098(12)	1.8776(7)	1.2895(4)	1.8671(6)	1.1882(3)	1.1882(3)
	MSE ($\hat{\theta}$)	0.1451(7)	0.1456(6)	0.3806(12)	0.1447(6)	0.3825(12)	0.7957(14)	0.1491(11)	0.1452(8)	0.0742(6)	0.1474(10)	0.1377(4)	0.0699(2)	0.1411(9)	0.0694(3)	0.0694(3)
	MSE ($\hat{\beta}$)	3.6163(11)	3.6125(10)	9.5434(14)	3.5486(8)	9.5842(15)	0.8506(2)	3.7064(12)	3.6072(9)	1.8285(5)	3.6486(12)	3.5255(7)	1.6648(4)	3.4863(6)	1.4131(3)	1.4131(3)
	MRE ($\hat{\theta}$)	7.6054(9)	7.6146(8)	12.2837(12)	7.5900(6)	12.3199(13)	17.8443(14)	7.7607(11)	7.6044(7)	5.4320(6)	7.6854(12)	7.3829(4)	373.1595(16)	5.2684(2)	4.8567(3)	4.8567(3)
	MRE ($\hat{\beta}$)	7.6064(11)	7.6021(10)	12.2880(13)	7.5349(8)	12.3188(13)	3.6858(2)	7.6965(12)	7.5969(9)	5.4079(6)	7.6391(12)	7.5105(7)	5.1580(4)	7.4685(6)	4.7528(3)	4.7528(3)
D_{abs}	0.0170(2)	0.0178(6)	0.0186(6)	0.0170(1)	0.0187(7)	0.3941(14)	0.0178(9)	0.0192(8)	0.0215(10)	0.0200(9)	0.0271(13)	0.0217(11)	0.0177(9)	0.0217(12)	0.0217(12)	
D_{max}	0.0245(3)	0.0257(6)	0.0268(6)	0.0244(1)	0.0270(7)	0.6423(14)	0.0257(9)	0.0277(8)	0.0310(10)	0.0287(9)	0.0393(13)	0.0254(13)	0.0313(11)	0.0314(12)	0.0314(12)	
$\sum Ranks$	60(92)	66(92)	90(92)	44(92)	98(92)	76(76)	81(92)	65(92)	44(92)	84(92)	59(92)	40(92)	39(92)	36(92)	36(92)	
200	$Bias(\hat{\theta})$	0.3822(6)	0.3822(6)	0.6240(13)	0.3784(6)	0.6239(12)	0.9041(14)	0.3846(11)	0.3794(7)	0.2727(6)	0.3845(10)	0.3627(4)	0.2614(2)	0.3740(5)	0.2402(1)	0.2402(1)
	$Bias(\hat{\beta})$	1.9008(10)	1.9081(11)	3.1181(13)	1.8840(8)	3.1166(14)	0.9345(2)	1.9197(12)	1.8978(9)	1.3523(5)	1.9142(12)	1.8771(7)	1.2845(4)	1.8627(6)	1.1805(3)	1.1805(3)
	MSE ($\hat{\theta}$)	0.1449(9)	0.1464(6)	0.3920(13)	0.1436(6)	0.3919(12)	0.8175(14)	0.1485(11)	0.1442(7)	0.0746(6)	0.1484(10)	0.1329(4)	0.0688(2)	0.1402(9)	0.0680(3)	0.0680(3)
	MSE ($\hat{\beta}$)	3.6134(10)	3.6410(11)	9.7968(13)	3.5497(8)	9.7874(14)	0.8824(2)	3.6889(12)	3.6017(9)	1.8292(6)	3.6654(12)	3.5234(7)	1.6517(4)	3.4697(6)	1.3946(3)	1.3946(3)
	MRE ($\hat{\theta}$)	7.6051(8)	7.6433(9)	12.4800(13)	7.5687(6)	12.4779(12)	18.0819(14)	7.6927(11)	7.5874(7)	5.4541(6)	7.6902(12)	7.2547(4)	5.2290(2)	7.4508(6)	4.8043(3)	4.8043(3)
	MRE ($\hat{\beta}$)	7.6033(10)	7.6324(11)	12.4723(13)	7.5360(8)	12.4662(14)	3.7379(2)	7.6790(12)	7.5912(9)	5.4090(6)	7.6569(12)	7.3083(7)	5.1378(4)	7.4507(6)	4.7219(3)	4.7219(3)
	D_{abs}	0.0124(3)	0.0124(6)	0.0124(6)	0.0123(3)	0.0124(6)	0.3946(14)	0.0119(9)	0.0126(7)	0.0153(10)	0.0151(9)	0.0231(13)	0.0132(4)	0.0156(12)	0.0156(11)	0.0156(11)
	D_{max}	0.0180(3)	0.0174(6)	0.0179(6)	0.0178(3)	0.0179(6)	0.6432(14)	0.0172(9)	0.0183(7)	0.0221(10)	0.0218(9)	0.0337(13)	0.0190(8)	0.0269(12)	0.0269(11)	0.0269(11)
	$\sum Ranks$	66(92)	64(92)	92(92)	48(92)	88(92)	76(76)	74(92)	62(92)	44(92)	84(92)	59(92)	42(92)	49(92)	34(92)	34(92)
	$Bias(\hat{\theta})$	0.3790(9)	0.3805(6)	0.6355(13)	0.3763(6)	0.6354(12)	0.9184(14)	0.3845(11)	0.3785(7)	0.2684(6)	0.3809(10)	0.3708(4)	0.2531(2)	0.3725(5)	0.2329(1)	0.2329(1)
	$Bias(\hat{\beta})$	1.9004(10)	1.9075(11)	3.1928(13)	1.8814(8)	3.1920(14)	0.9436(2)	1.9285(12)	1.8982(9)	1.3480(5)	1.9199(12)	1.8780(7)	1.2591(4)	1.8630(6)	1.1585(3)	1.1585(3)
	MSE ($\hat{\theta}$)	0.1439(9)	0.1451(6)	0.4051(13)	0.1419(6)	0.4051(12)	0.8435(14)	0.1483(11)	0.1435(7)	0.0722(6)	0.1455(10)	0.1381(4)	0.0643(2)	0.1391(9)	0.0544(3)	0.0544(3)
	MSE ($\hat{\beta}$)	3.6117(10)	3.6386(11)	10.2521(13)	3.5398(8)	10.2466(14)	0.8992(2)	3.7214(12)	3.6032(9)	1.8173(6)	3.6870(12)	7.5248(7)	5.0820(10)	3.4709(6)	1.3432(3)	1.3432(3)
	MRE ($\hat{\theta}$)	7.5808(8)	7.6097(9)	12.7091(13)	7.5257(6)	12.7081(12)	18.3675(14)	7.6903(11)	7.5693(7)	5.3679(6)	7.6183(12)	7.4159(4)	5.0616(2)	7.4510(6)	4.6573(3)	4.6573(3)
	MRE ($\hat{\beta}$)	7.6016(10)	7.6299(11)	12.7714(13)	7.5256(8)	12.7679(14)	3.7745(2)	7.7140(12)	7.5927(9)	5.3919(6)	7.6795(12)	7.5119(7)	5.0362(4)	7.4521(6)	4.6339(3)	4.6339(3)
D_{abs}	0.0166(4)	0.0161(6)	0.0166(6)	0.0114(3)	0.0116(6)	0.3950(14)	0.0114(9)	0.0116(7)	0.0147(10)	0.0145(9)	0.0175(13)	0.0139(4)	0.0156(12)	0.0140(10)	0.0140(10)	
D_{max}	0.0166(4)	0.0165(6)	0.0168(6)	0.0164(2)	0.0168(6)	0.6457(14)										

Table 4: Partial and Overall Ranks for all Estimation Methods of the BE Distribution from 3

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLSE	LTADE	MSADE	MSALDE	ADSOE	KE	MSSDE	MSSLDE	MSLNDE
15	$ Bias(\hat{\theta}) $	7.0	8.0	6.0	11.0	4.0	14.0	13.0	9.0	3.0	12.0	5.0	15.0	2.0	10.0	1.0
	$ Bias(\hat{\beta}) $	12.0	11.0	7.0	9.0	6.0	2.0	15.0	13.0	5.0	14.0	8.0	1.0	4.0	10.0	3.0
	MSE ($\hat{\theta}$)	7.0	8.0	5.0	10.0	4.0	14.0	13.0	9.0	2.0	12.0	6.0	15.0	3.0	11.0	1.0
	MSE ($\hat{\beta}$)	13.0	11.0	7.0	10.0	6.0	2.0	15.0	12.0	5.0	14.0	8.0	1.0	4.0	9.0	3.0
	MRE ($\hat{\theta}$)	7.0	8.0	6.0	11.0	4.0	14.0	13.0	9.0	3.0	12.0	5.0	15.0	2.0	10.0	1.0
	MRE ($\hat{\beta}$)	12.0	11.0	7.0	9.0	6.0	2.0	15.0	13.0	5.0	14.0	8.0	1.0	4.0	10.0	3.0
	D _{abs}	2.0	6.0	7.0	3.0	8.0	14.0	4.0	10.0	5.0	1.0	13.0	15.0	11.0	9.0	12.0
	D _{max}	2.0	6.0	7.0	3.0	8.0	14.0	4.0	10.0	5.0	1.0	13.0	15.0	11.0	9.0	11.0
	$\sum$ Ranks	62(6)	69(9)	52(5)	66(7)	46(4)	76(10)	93(15)	85(14)	32(1)	80(13)	66(7)	78(11)	42(3)	78(11)	35(2)
	$ Bias(\hat{\theta}) $	7.0	5.0	13.0	8.0	12.0	14.0	11.0	6.0	2.0	10.0	4.0	15.0	3.0	9.0	1.0
50	$ Bias(\hat{\beta}) $	10.0	9.0	15.0	8.0	14.0	2.0	13.0	11.0	5.0	12.0	6.0	1.0	4.0	7.0	3.0
	MSE ($\hat{\theta}$)	7.0	5.0	13.0	8.0	12.0	14.0	11.0	6.0	2.0	10.0	4.0	15.0	3.0	9.0	1.0
	MRE ($\hat{\theta}$)	7.0	5.0	13.0	8.0	12.0	14.0	11.0	6.0	2.0	10.0	4.0	15.0	3.0	9.0	1.0
	MRE ($\hat{\beta}$)	10.0	9.0	15.0	8.0	14.0	2.0	13.0	11.0	5.0	12.0	6.0	1.0	4.0	7.0	3.0
	D _{abs}	1.0	4.0	6.0	2.0	7.0	14.0	3.0	8.0	10.0	5.0	13.0	15.0	11.0	9.0	12.0
	D _{max}	1.0	4.0	6.0	2.0	7.0	14.0	3.0	8.0	10.0	5.0	13.0	15.0	11.0	9.0	12.0
	$\sum$ Ranks	53(6)	50(4)	96(13)	52(5)	92(14)	76(10)	78(12)	67(9)	41(2)	76(10)	56(7)	78(12)	43(3)	66(8)	36(1)
	$ Bias(\hat{\theta}) $	7.0	6.0	12.0	9.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\hat{\beta}) $	11.0	9.0	15.0	8.0	14.0	2.0	13.0	10.0	5.0	12.0	6.0	1.0	4.0	7.0	3.0
	MSE ($\hat{\theta}$)	7.0	6.0	12.0	9.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
100	MSE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	2.0	13.0	10.0	5.0	12.0	6.0	1.0	4.0	7.0	3.0
	MRE ($\hat{\theta}$)	7.0	6.0	12.0	9.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	11.0	9.0	15.0	8.0	14.0	2.0	13.0	10.0	5.0	12.0	6.0	1.0	4.0	7.0	3.0
	D _{abs}	3.0	1.0	5.0	6.0	4.0	14.0	2.0	7.0	11.0	12.0	13.0	15.0	8.0	10.0	9.0
	D _{max}	3.0	1.0	5.0	6.0	4.0	14.0	2.0	7.0	11.0	12.0	13.0	15.0	8.0	10.0	9.0
	$\sum$ Ranks	60(7)	47(4)	91(12)	63(8)	89(13)	76(10)	76(10)	68(9)	46(3)	90(14)	59(6)	78(12)	34(2)	53(5)	30(1)
	$ Bias(\hat{\theta}) $	8.0	9.0	12.0	6.0	13.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\hat{\beta}) $	11.0	10.0	14.0	8.0	15.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	MSE ($\hat{\theta}$)	7.0	9.0	12.0	6.0	13.0	14.0	11.0	8.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MSE ($\hat{\beta}$)	11.0	10.0	14.0	8.0	15.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
200	MRE ($\hat{\theta}$)	8.0	9.0	12.0	6.0	13.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	11.0	10.0	14.0	8.0	15.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	D _{abs}	2.0	5.0	6.0	1.0	7.0	14.0	4.0	8.0	10.0	9.0	13.0	15.0	11.0	9.0	12.0
	D _{max}	2.0	5.0	6.0	1.0	7.0	14.0	4.0	8.0	10.0	9.0	13.0	15.0	11.0	9.0	12.0
	$\sum$ Ranks	60(7)	66(9)	90(14)	44(4)	98(15)	76(10)	81(12)	65(8)	44(4)	84(13)	59(6)	78(11)	40(3)	39(2)	36(1)
	$ Bias(\hat{\theta}) $	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\hat{\beta}) $	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	MSE ($\hat{\theta}$)	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MSE ($\hat{\beta}$)	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	MRE ($\hat{\theta}$)	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
300	MRE ($\hat{\beta}$)	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	D _{abs}	6.0	2.0	4.0	3.0	5.0	14.0	1.0	7.0	10.0	9.0	13.0	15.0	12.0	8.0	11.0
	D _{max}	6.0	2.0	4.0	3.0	5.0	14.0	1.0	7.0	10.0	9.0	13.0	15.0	12.0	8.0	11.0
	$\sum$ Ranks	66(9)	64(8)	92(13)	48(4)	88(14)	76(11)	74(10)	62(7)	44(3)	84(13)	59(6)	78(12)	42(2)	49(3)	34(1)
	$ Bias(\hat{\theta}) $	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\hat{\beta}) $	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	MSE ($\hat{\theta}$)	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MSE ($\hat{\beta}$)	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
	MRE ($\hat{\theta}$)	8.0	9.0	13.0	6.0	12.0	14.0	11.0	7.0	3.0	10.0	4.0	15.0	2.0	5.0	1.0
	MRE ($\hat{\beta}$)	10.0	11.0	15.0	8.0	14.0	2.0	13.0	9.0	5.0	12.0	7.0	1.0	4.0	6.0	3.0
500	D _{abs}	4.0	3.0	6.0	2.0	5.0	14.0	1.0	7.0	12.0	11.0	13.0	15.0	9.0	8.0	10.0
	D _{max}	4.0	3.0	6.0	2.0	5.0	14.0	1.0	7.0	12.0	11.0	13.0	15.0	9.0	8.0	10.0
	$\sum$ Ranks	62(7)	66(9)	96(13)	46(3)	88(13)	76(11)	74(10)	62(7)	45(4)	85(13)	59(6)	78(12)	36(2)	49(3)	32(1)
	$ Bias(\hat{\theta}) $	7.0	8.0	11.0	9.0	10.0	14.0	13.0	9.0	3.0	12.0	5.0	15.0	2.0	10.0	1.0
	$ Bias(\hat{\beta}) $	12.0	11.0	14.0	10.0	11.0	2.0	14.0	10.0	5.0	13.0	8.0	1.0	4.0	9.0	3.0
	MSE ($\hat{\theta}$)	7.0	8.0	11.0	9.0	10.0	14.0	13.0	9.0	3.0	12.0	5.0	15.0	2.0	10.0	1.0
	MSE ($\hat{\beta}$)	12.0	11.0	14.0	10.0	11.0	2.0	14.0	10.0	5.0	13.0	8.0	1.0	4.0	9.0	3.0
	MRE ($\hat{\theta}$)	7.0	8.0	11.0	9.0	10.0	14.0	13.0	9.0	3.0	12.0	5.0	15.0	2.0	10.0	1.0
	MRE ($\hat{\beta}$)	12.0	11.0	14.0	10.0	11.0	2.0	14.0	10.0	5.0	13.0	8.0	1.0	4.0	9.0	3.0
	$\sum$ Ranks	62(7)	66(9)	96(13)	46(3)	88(13)	76(11)	74(10)	62(7)	45(4)	85(13)	59(6)	78(12)	36(2)	49(3)	32(1)

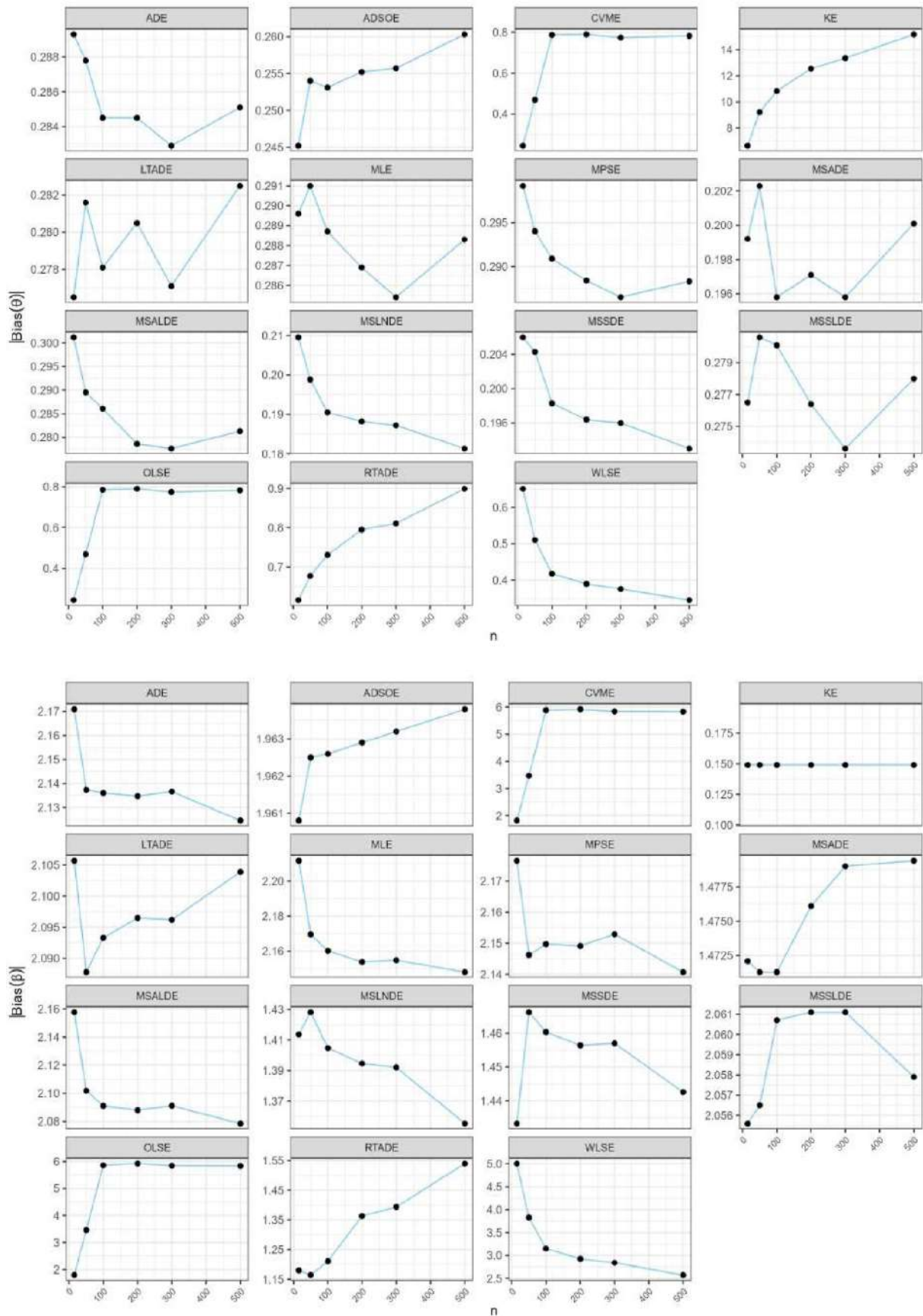


Figure 13: Graphical representations for the Bias values presented in Table 5.

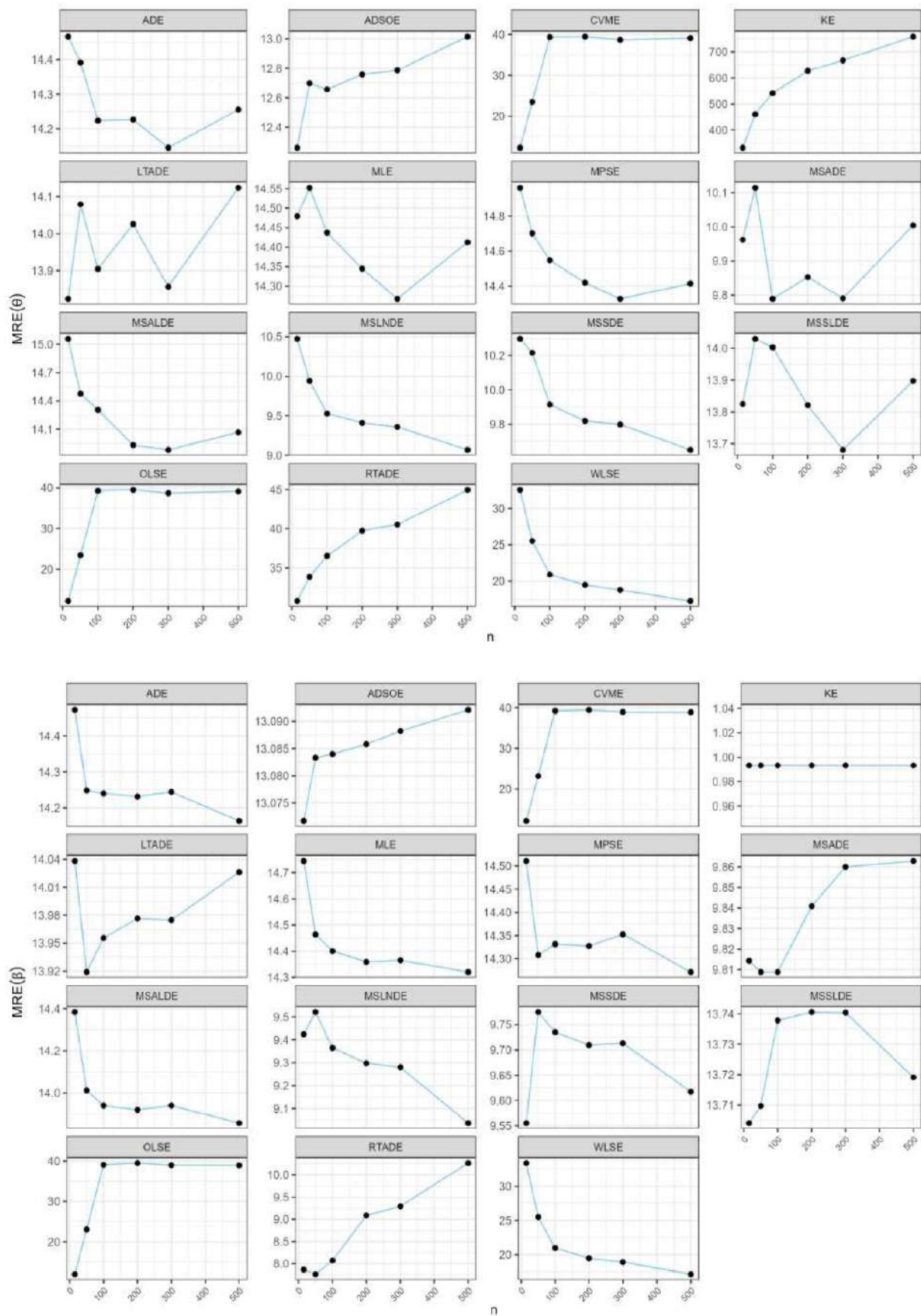


Figure 14: Graphical representations for the MRE values presented in Table 5.

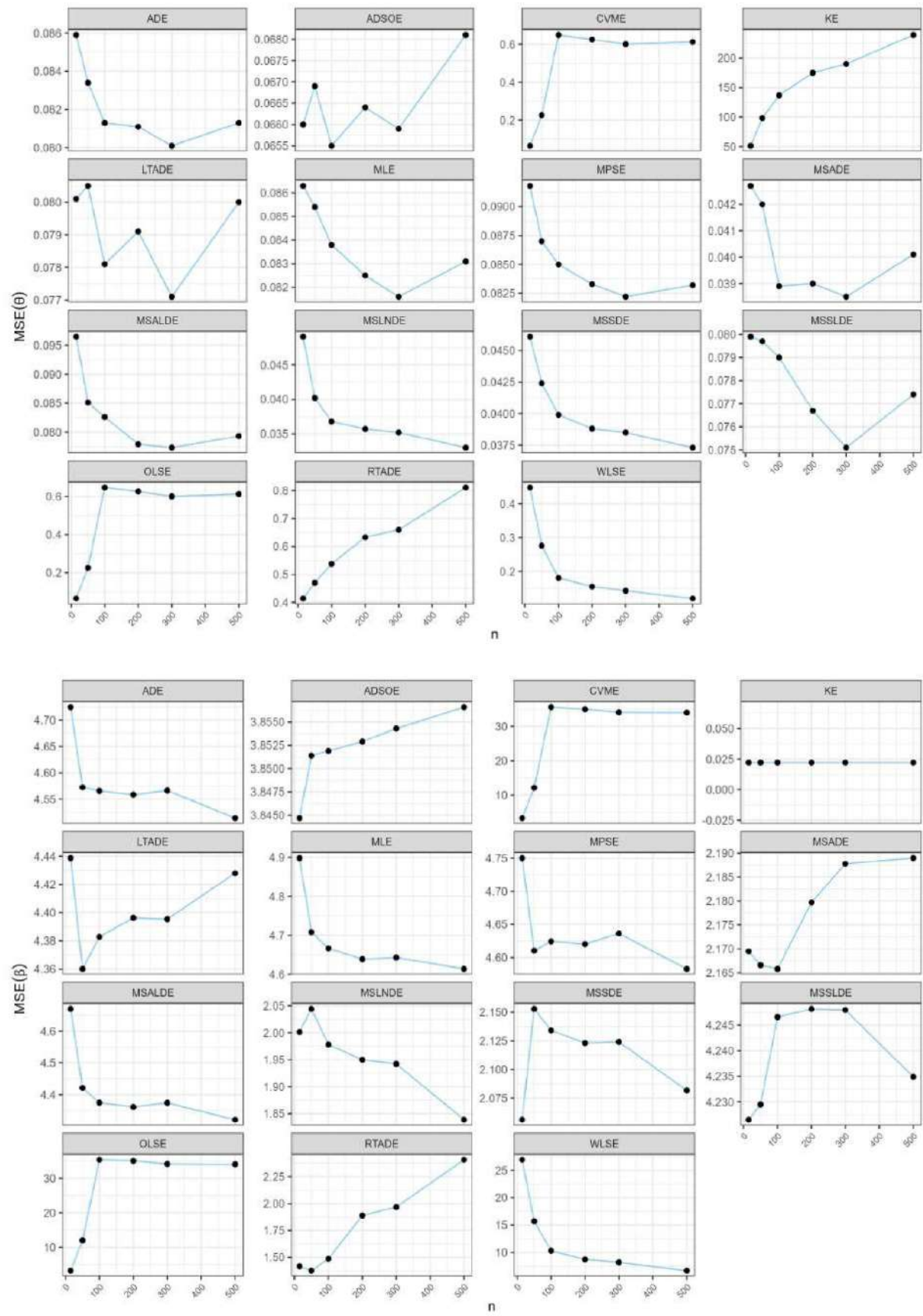


Figure 15: Graphical representations for the MSE values presented in Table 5.

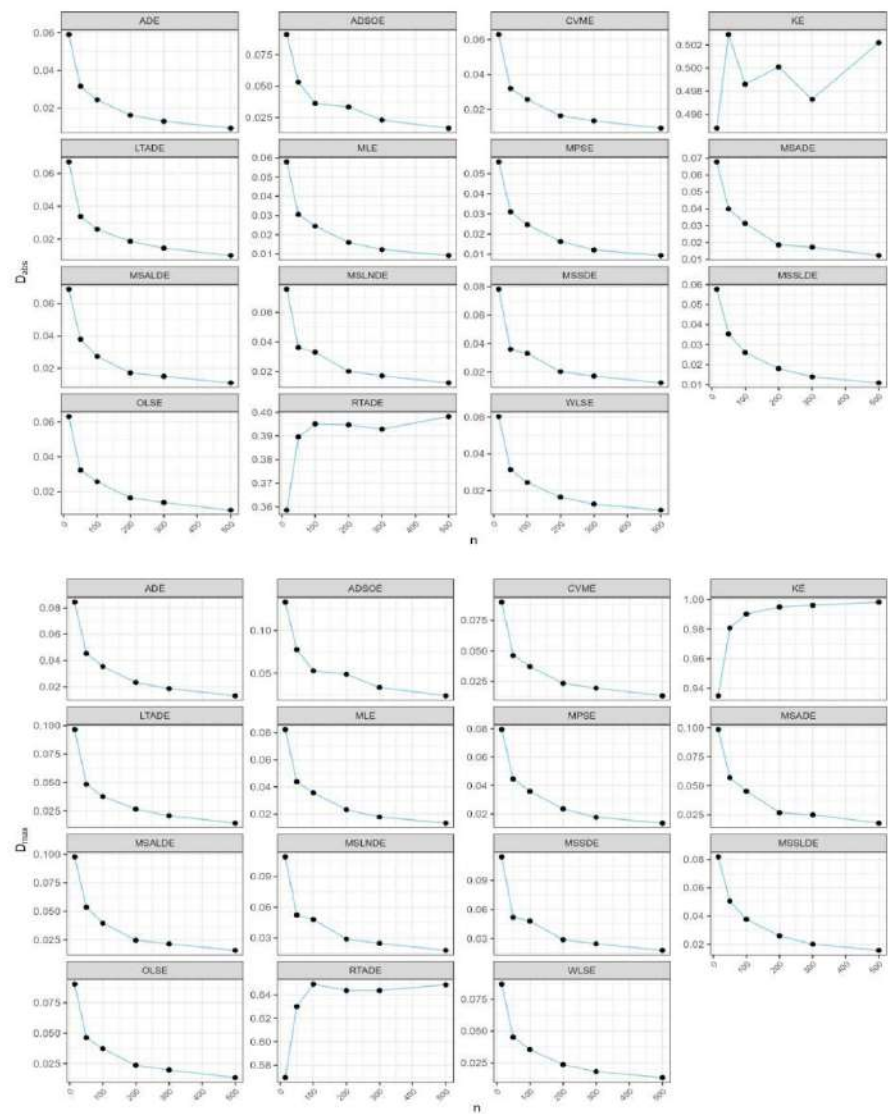


Figure 16: Graphical representations for the D_{abs} and D_{max} values presented in Table 5.

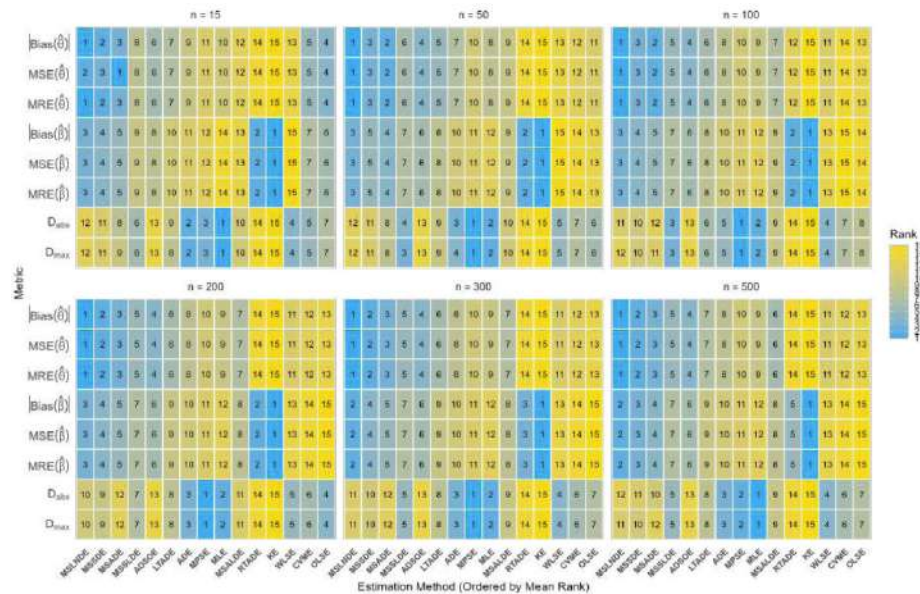


Figure 17: Partial Rank Heatmap for simulation in Table 5.

Table 5: Numerical Values of Simulation Measures for $\theta = 0.02, \beta = 0.15$

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLS	LTAD	MSADE	MSALDE	ADSOE	KE	MSSE	MISSDE	MSLNDE
15	$ Bias(\hat{\theta}) $	0.2896 ⁽¹⁰⁾	0.2893 ⁽⁹⁾	0.2454 ⁽⁶⁾	0.2992 ⁽¹¹⁾	0.2447 ⁽⁴⁾	0.6159 ⁽¹³⁾	0.6512 ⁽¹³⁾	0.2765 ⁽⁷⁾	0.1992 ⁽²⁾	0.3012 ⁽¹²⁾	0.2452 ⁽³⁾	6.6432 ⁽¹³⁾	0.2060 ⁽²⁾	0.2765 ⁽⁹⁾	0.2095 ⁽³⁾
	$ Bias(\hat{\beta}) $	2.2117 ⁽¹⁴⁾	2.1709 ⁽¹²⁾	1.8169 ⁽⁷⁾	2.1766 ⁽¹³⁾	1.7985 ⁽⁶⁾	1.1804 ⁽²⁾	5.0038 ⁽¹⁰⁾	2.1057 ⁽¹⁰⁾	1.4721 ⁽⁹⁾	2.1577 ⁽¹¹⁾	1.9608 ⁽⁸⁾	1.4190 ⁽¹⁾	1.4333 ⁽⁴⁾	2.0556 ⁽⁹⁾	1.4136 ⁽⁸⁾
	MSE ($\hat{\theta}$)	0.0863 ⁽¹⁰⁾	0.0859 ⁽⁹⁾	0.0643 ⁽⁶⁾	0.0918 ⁽¹¹⁾	0.0638 ⁽⁴⁾	0.4140 ⁽¹³⁾	0.4483 ⁽¹⁴⁾	0.0801 ⁽⁷⁾	0.0427 ⁽²⁾	0.0965 ⁽¹²⁾	0.0660 ⁽⁶⁾	51.2973 ⁽¹³⁾	0.0461 ⁽²⁾	0.0799 ⁽⁷⁾	0.0491 ⁽⁸⁾
	MSE ($\hat{\beta}$)	4.7245 ⁽¹²⁾	4.7245 ⁽¹²⁾	3.3194 ⁽⁷⁾	4.7502 ⁽¹³⁾	3.2510 ⁽⁶⁾	1.4163 ⁽²⁾	26.9135 ⁽¹⁰⁾	4.4387 ⁽¹⁰⁾	2.1695 ⁽⁹⁾	4.6701 ⁽¹¹⁾	4.8447 ⁽⁸⁾	0.0222 ⁽¹⁾	2.0561 ⁽²⁾	4.2265 ⁽⁹⁾	2.0014 ⁽⁸⁾
	MRE ($\hat{\theta}$)	14.4793 ⁽¹⁴⁾	14.4665 ⁽¹²⁾	12.2719 ⁽⁷⁾	14.9618 ⁽¹³⁾	12.2369 ⁽⁶⁾	30.7928 ⁽¹³⁾	32.5623 ⁽¹⁴⁾	13.8237 ⁽¹⁰⁾	9.9622 ⁽⁹⁾	15.0680 ⁽¹²⁾	12.2606 ⁽⁸⁾	332.1619 ⁽¹³⁾	10.2977 ⁽²⁾	13.8251 ⁽⁸⁾	10.4735 ⁽⁹⁾
	MRE ($\hat{\beta}$)	14.7445 ⁽¹⁴⁾	14.4724 ⁽¹²⁾	12.1129 ⁽⁷⁾	14.5109 ⁽¹³⁾	11.9898 ⁽⁶⁾	7.8696 ⁽²⁾	33.3584 ⁽¹⁰⁾	14.0382 ⁽¹⁰⁾	9.8143 ⁽⁹⁾	14.3845 ⁽¹¹⁾	13.0717 ⁽⁸⁾	0.9933 ⁽¹⁾	9.5551 ⁽⁴⁾	13.7040 ⁽⁸⁾	9.4240 ⁽⁸⁾
	D_{abs}	0.0579 ⁽³⁾	0.0591 ⁽⁴⁾	0.0629 ⁽⁶⁾	0.0538 ⁽⁵⁾	0.0631 ⁽⁴⁾	0.3586 ⁽¹⁴⁾	0.0604 ⁽⁴⁾	0.0670 ⁽⁸⁾	0.0673 ⁽⁹⁾	0.0911 ⁽¹⁰⁾	0.0577 ⁽³⁾	0.4948 ⁽¹³⁾	0.0783 ⁽¹¹⁾	0.0757 ⁽⁹⁾	0.0757 ⁽¹¹⁾
	D_{max}	0.0824 ⁽³⁾	0.0843 ⁽⁴⁾	0.0899 ⁽⁶⁾	0.0796 ⁽⁵⁾	0.0904 ⁽⁷⁾	0.5695 ⁽¹⁴⁾	0.0869 ⁽⁵⁾	0.0967 ⁽⁸⁾	0.0987 ⁽¹⁰⁾	0.0978 ⁽⁹⁾	0.1330 ⁽¹²⁾	0.9350 ⁽¹³⁾	0.1145 ⁽¹²⁾	0.0820 ⁽⁹⁾	0.1002 ⁽¹¹⁾
	$\sum Ranks$	78 ⁽¹²⁾	71 ⁽⁹⁾	50 ⁽⁶⁾	74 ⁽¹¹⁾	44 ⁽⁴⁾	75 ⁽¹³⁾	97 ⁽¹³⁾	68 ⁽⁶⁾	37 ⁽²⁾	88 ⁽¹²⁾	68 ⁽¹²⁾	78 ⁽¹²⁾	42 ⁽³⁾	54 ⁽⁹⁾	40 ⁽⁸⁾
	50	$ Bias(\hat{\theta}) $	0.2910 ⁽⁹⁾	0.2878 ⁽⁷⁾	0.4703 ⁽¹²⁾	0.2940 ⁽¹⁰⁾	0.4692 ⁽¹¹⁾	0.6775 ⁽¹⁴⁾	0.5106 ⁽¹³⁾	0.2816 ⁽⁶⁾	0.2023 ⁽²⁾	0.2895 ⁽⁸⁾	0.2540 ⁽⁴⁾	9.2116 ⁽¹³⁾	0.2043 ⁽³⁾	0.2806 ⁽⁵⁾
$ Bias(\hat{\beta}) $		2.1696 ⁽¹²⁾	2.1373 ⁽¹⁰⁾	3.4751 ⁽¹⁴⁾	2.1462 ⁽¹¹⁾	3.4597 ⁽¹³⁾	1.1646 ⁽²⁾	3.8297 ⁽¹⁰⁾	2.0878 ⁽⁸⁾	1.4713 ⁽⁹⁾	2.1018 ⁽⁶⁾	1.9625 ⁽⁶⁾	1.4190 ⁽¹⁾	1.4662 ⁽⁴⁾	2.0565 ⁽⁷⁾	1.4281 ⁽⁸⁾
MSE ($\hat{\theta}$)		0.0854 ⁽⁹⁾	0.0834 ⁽⁷⁾	0.2259 ⁽¹²⁾	0.0870 ⁽¹⁰⁾	0.2249 ⁽¹¹⁾	0.4702 ⁽¹⁴⁾	0.2761 ⁽¹³⁾	0.0805 ⁽⁵⁾	0.0426 ⁽²⁾	0.0851 ⁽⁸⁾	0.0669 ⁽⁴⁾	98.1426 ⁽¹³⁾	0.0424 ⁽³⁾	0.0797 ⁽⁷⁾	0.0402 ⁽⁸⁾
MSE ($\hat{\beta}$)		4.7083 ⁽¹²⁾	4.5722 ⁽¹⁰⁾	12.1351 ⁽¹⁴⁾	4.6104 ⁽¹¹⁾	12.0330 ⁽¹³⁾	1.3752 ⁽²⁾	15.6380 ⁽¹⁰⁾	4.3601 ⁽⁸⁾	2.1666 ⁽⁹⁾	4.4211 ⁽⁶⁾	3.8514 ⁽⁴⁾	0.0222 ⁽¹⁾	2.1529 ⁽⁴⁾	4.2295 ⁽⁷⁾	2.0442 ⁽⁸⁾
MRE ($\hat{\theta}$)		14.5121 ⁽¹⁴⁾	14.3919 ⁽¹²⁾	23.5163 ⁽¹³⁾	14.7010 ⁽¹⁰⁾	23.4611 ⁽¹¹⁾	33.8753 ⁽¹⁴⁾	25.5315 ⁽¹³⁾	14.0796 ⁽⁹⁾	11.1442 ⁽²⁾	14.4751 ⁽¹¹⁾	12.6981 ⁽¹⁰⁾	460.5796 ⁽¹³⁾	10.2161 ⁽⁹⁾	14.0301 ⁽⁸⁾	9.9416 ⁽¹⁾
MRE ($\hat{\beta}$)		14.4640 ⁽¹²⁾	14.2484 ⁽¹⁰⁾	23.1630 ⁽¹³⁾	14.3076 ⁽¹¹⁾	23.0647 ⁽¹⁰⁾	7.7637 ⁽²⁾	25.5047 ⁽¹³⁾	13.9190 ⁽⁸⁾	9.8088 ⁽⁹⁾	14.0121 ⁽¹¹⁾	13.0833 ⁽¹⁰⁾	0.9933 ⁽¹⁾	9.7749 ⁽⁴⁾	13.7097 ⁽⁷⁾	9.5208 ⁽⁶⁾
D_{abs}		0.0305 ⁽¹⁾	0.0315 ⁽⁴⁾	0.0321 ⁽⁶⁾	0.0311 ⁽²⁾	0.0322 ⁽⁵⁾	0.3895 ⁽¹⁴⁾	0.0314 ⁽⁴⁾	0.0338 ⁽⁷⁾	0.0399 ⁽¹⁰⁾	0.0379 ⁽¹¹⁾	0.0532 ⁽¹⁰⁾	0.5029 ⁽¹³⁾	0.0360 ⁽⁶⁾	0.0354 ⁽⁸⁾	0.0362 ⁽¹⁰⁾
D_{max}		0.0438 ⁽¹⁾	0.0453 ⁽⁴⁾	0.0462 ⁽⁶⁾	0.0447 ⁽²⁾	0.0463 ⁽⁵⁾	0.6299 ⁽¹⁴⁾	0.0452 ⁽³⁾	0.0485 ⁽⁷⁾	0.0571 ⁽¹²⁾	0.0537 ⁽¹¹⁾	0.0776 ⁽¹⁰⁾	0.9807 ⁽¹³⁾	0.0519 ⁽⁹⁾	0.0507 ⁽⁸⁾	0.0522 ⁽¹⁰⁾
$\sum Ranks$		65 ⁽⁸⁾	59 ⁽⁷⁾	88 ⁽¹⁴⁾	67 ⁽⁹⁾	84 ⁽¹⁰⁾	76 ⁽¹¹⁾	90 ⁽¹³⁾	56 ⁽⁶⁾	45 ⁽³⁾	73 ⁽¹⁰⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	39 ⁽²⁾	52 ⁽⁴⁾	32 ⁽¹⁾
100		$ Bias(\hat{\theta}) $	0.2887 ⁽⁹⁾	0.2845 ⁽⁷⁾	0.7874 ⁽¹⁴⁾	0.2909 ⁽¹⁰⁾	0.7555 ⁽¹³⁾	0.7306 ⁽¹²⁾	0.4177 ⁽¹¹⁾	0.2781 ⁽⁵⁾	0.1958 ⁽²⁾	0.2861 ⁽⁸⁾	0.2531 ⁽⁴⁾	10.8277 ⁽¹³⁾	0.1983 ⁽³⁾	0.2801 ⁽⁶⁾
	$ Bias(\hat{\beta}) $	2.1360 ⁽¹²⁾	2.1360 ⁽¹⁰⁾	5.8824 ⁽¹⁴⁾	2.1497 ⁽¹¹⁾	5.8569 ⁽¹³⁾	1.2105 ⁽²⁾	3.1477 ⁽¹⁰⁾	2.0933 ⁽⁸⁾	1.4713 ⁽⁹⁾	2.0910 ⁽⁶⁾	1.9626 ⁽⁶⁾	1.4190 ⁽¹⁾	1.4673 ⁽⁴⁾	2.0607 ⁽⁷⁾	1.4047 ⁽⁸⁾
	MSE ($\hat{\theta}$)	0.0838 ⁽⁹⁾	0.0813 ⁽⁷⁾	0.6492 ⁽¹⁴⁾	0.0830 ⁽¹⁰⁾	0.6471 ⁽¹³⁾	0.5377 ⁽¹²⁾	0.1867 ⁽¹¹⁾	0.0781 ⁽⁵⁾	0.0389 ⁽²⁾	0.0826 ⁽⁸⁾	0.0655 ⁽⁴⁾	136.7376 ⁽¹³⁾	0.0399 ⁽³⁾	0.0790 ⁽⁶⁾	0.0368 ⁽⁸⁾
	MSE ($\hat{\beta}$)	4.6667 ⁽¹²⁾	4.5653 ⁽¹⁰⁾	35.6407 ⁽¹⁴⁾	4.6243 ⁽¹¹⁾	35.3613 ⁽¹³⁾	1.4862 ⁽²⁾	10.3157 ⁽¹⁰⁾	4.3828 ⁽⁹⁾	2.1658 ⁽⁹⁾	4.3742 ⁽⁸⁾	3.8519 ⁽⁶⁾	0.0222 ⁽¹⁾	2.1342 ⁽⁴⁾	4.2466 ⁽⁷⁾	1.9782 ⁽⁸⁾
	MRE ($\hat{\theta}$)	14.4363 ⁽¹⁴⁾	14.2234 ⁽¹²⁾	39.2682 ⁽¹³⁾	14.5474 ⁽¹⁰⁾	39.2772 ⁽¹¹⁾	36.5321 ⁽¹²⁾	20.8846 ⁽¹¹⁾	13.9045 ⁽⁸⁾	9.7592 ⁽²⁾	14.3053 ⁽¹⁰⁾	12.6560 ⁽⁹⁾	541.3852 ⁽¹³⁾	9.9153 ⁽³⁾	14.0032 ⁽⁶⁾	9.5272 ⁽⁵⁾
	MRE ($\hat{\beta}$)	14.4004 ⁽¹²⁾	14.2399 ⁽¹⁰⁾	39.2225 ⁽¹³⁾	14.3314 ⁽¹¹⁾	39.0475 ⁽¹⁰⁾	8.0715 ⁽²⁾	20.9845 ⁽¹⁰⁾	13.9401 ⁽⁸⁾	9.8088 ⁽⁹⁾	13.9401 ⁽⁸⁾	13.7379 ⁽⁷⁾	0.9933 ⁽¹⁾	9.7352 ⁽⁴⁾	13.7379 ⁽⁷⁾	9.3646 ⁽⁶⁾
	D_{abs}	0.0245 ⁽³⁾	0.0244 ⁽⁴⁾	0.0258 ⁽⁶⁾	0.0247 ⁽²⁾	0.0256 ⁽⁵⁾	0.3951 ⁽¹⁴⁾	0.0244 ⁽⁴⁾	0.0260 ⁽⁸⁾	0.0314 ⁽¹⁰⁾	0.0273 ⁽⁹⁾	0.0362 ⁽¹⁰⁾	0.4986 ⁽¹³⁾	0.0330 ⁽¹¹⁾	0.0260 ⁽⁷⁾	0.0330 ⁽¹²⁾
	D_{max}	0.0356 ⁽³⁾	0.0355 ⁽⁴⁾	0.0374 ⁽⁶⁾	0.0339 ⁽²⁾	0.0372 ⁽⁵⁾	0.6491 ⁽¹⁴⁾	0.0354 ⁽¹⁾	0.0377 ⁽⁸⁾	0.0453 ⁽¹⁰⁾	0.0396 ⁽⁹⁾	0.0530 ⁽¹⁰⁾	0.9901 ⁽¹³⁾	0.0479 ⁽¹²⁾	0.0377 ⁽⁷⁾	0.0479 ⁽¹¹⁾
	$\sum Ranks$	69 ⁽⁹⁾	55 ⁽⁷⁾	99 ⁽¹⁴⁾	71 ⁽¹⁰⁾	91 ⁽¹¹⁾	70 ⁽¹¹⁾	74 ⁽¹⁰⁾	58 ⁽⁶⁾	41 ⁽³⁾	63 ⁽⁸⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	44 ⁽³⁾	53 ⁽⁴⁾	35 ⁽¹⁾
	200	$ Bias(\hat{\theta}) $	0.2869 ⁽⁹⁾	0.2845 ⁽⁷⁾	0.7895 ⁽¹⁴⁾	0.2884 ⁽¹⁰⁾	0.7900 ⁽¹³⁾	0.7949 ⁽¹⁴⁾	0.3867 ⁽¹¹⁾	0.2805 ⁽⁵⁾	0.1971 ⁽²⁾	0.2786 ⁽⁸⁾	0.2552 ⁽⁴⁾	12.5378 ⁽¹³⁾	0.1964 ⁽³⁾	0.2764 ⁽⁶⁾
$ Bias(\hat{\beta}) $		2.1347 ⁽¹⁰⁾	2.1347 ⁽¹⁰⁾	5.9180 ⁽¹⁴⁾	2.1491 ⁽¹¹⁾	5.9100 ⁽¹³⁾	1.3628 ⁽²⁾	2.9252 ⁽¹⁰⁾	2.0965 ⁽⁸⁾	1.4761 ⁽⁹⁾	2.0880 ⁽⁶⁾	1.9629 ⁽⁶⁾	1.4190 ⁽¹⁾	1.4564 ⁽⁴⁾	2.0611 ⁽⁷⁾	1.3946 ⁽⁸⁾
MSE ($\hat{\theta}$)		0.0825 ⁽⁹⁾	0.0811 ⁽⁷⁾	0.6257 ⁽¹⁴⁾	0.0833 ⁽¹⁰⁾	0.6265 ⁽¹³⁾	0.5329 ⁽¹²⁾	0.1553 ⁽¹¹⁾	0.0791 ⁽⁵⁾	0.0390 ⁽²⁾	0.0779 ⁽⁸⁾	0.0664 ⁽⁴⁾	175.0082 ⁽¹³⁾	0.0388 ⁽³⁾	0.0767 ⁽⁶⁾	0.0367 ⁽⁸⁾
MSE ($\hat{\beta}$)		4.6393 ⁽¹²⁾	4.5585 ⁽¹⁰⁾	35.0240 ⁽¹⁴⁾	4.6202 ⁽¹¹⁾	35.0357 ⁽¹³⁾	1.8889 ⁽²⁾	8.7379 ⁽¹⁰⁾	4.3963 ⁽⁹⁾	2.1797 ⁽⁹⁾	4.3610 ⁽⁸⁾	3.8529 ⁽⁶⁾	0.0222 ⁽¹⁾	2.1231 ⁽⁴⁾	4.2482 ⁽⁷⁾	1.9498 ⁽⁸⁾
MRE ($\hat{\theta}$)		14.3449 ⁽¹⁴⁾	14.2266 ⁽¹²⁾	39.4765 ⁽¹³⁾	14.4189 ⁽¹⁰⁾	39.5002 ⁽¹¹⁾	30.7442 ⁽¹²⁾	19.4801 ⁽¹¹⁾	14.0259 ⁽⁸⁾	9.8527 ⁽²⁾	13.9313 ⁽¹⁰⁾	12.7585 ⁽⁹⁾	626.8922 ⁽¹³⁾	9.8101 ⁽³⁾	13.8212 ⁽⁶⁾	9.4093 ⁽⁵⁾
MRE ($\hat{\beta}$)		14.3588 ⁽¹²⁾	14.2313 ⁽¹⁰⁾	39.4534 ⁽¹³⁾	14.3275 ⁽¹¹⁾	39.4600 ⁽¹⁰⁾	10.5015 ⁽²⁾	10.5015 ⁽²⁾	13.0765 ⁽⁹⁾	9.8409 ⁽⁸⁾	13.9201 ⁽⁸⁾	13.0858 ⁽⁶⁾	0.9933 ⁽¹⁾	9.7005 ⁽⁴⁾	13.7406 ⁽⁷⁾	9.2971 ⁽⁸⁾
D_{abs}		0.0160 ⁽¹⁾	0.0161 ⁽⁴⁾	0.0164 ⁽⁶⁾	0.0163 ⁽²⁾	0.0164 ⁽⁵⁾	0.3947 ⁽¹⁴⁾	0.0164 ⁽⁴⁾	0.0186 ⁽⁸⁾	0.0171 ⁽⁷⁾	0.0334 ⁽¹⁰⁾	0.0334 ⁽¹⁰⁾	0.5001 ⁽¹³⁾	0.0202 ⁽¹²⁾	0.0180 ⁽⁶⁾	0.0201 ⁽¹¹⁾
D_{max}		0.0231 ⁽¹⁾	0.0232 ⁽⁴⁾	0.0236 ⁽⁶⁾	0.0235 ⁽²⁾	0.0236 ⁽⁵⁾	0.6438 ⁽¹⁴⁾	0.0237 ⁽⁶⁾	0.0266 ⁽⁹⁾	0.0268 ⁽¹⁰⁾	0.0247 ⁽⁷⁾	0.0488 ⁽¹⁰⁾	0.9505 ⁽¹³⁾	0.0250 ⁽¹²⁾	0.0260 ⁽⁸⁾	0.0289 ⁽¹¹⁾
$\sum Ranks$		65 ⁽⁸⁾	58 ⁽⁷⁾	86 ⁽¹⁴⁾	69 ⁽⁹⁾	94 ⁽¹⁰⁾	76 ⁽¹¹⁾	84 ⁽¹³⁾	66 ⁽⁶⁾	44 ⁽³⁾	56 ⁽⁶⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	42 ⁽³⁾	52 ⁽⁴⁾	34 ⁽¹⁾
300		$ Bias(\hat{\theta}) $	0.2854 ⁽⁹⁾	0.2829 ⁽⁷⁾	0.7740 ⁽¹⁴⁾	0.2865 ⁽¹⁰⁾	0.7739 ⁽¹³⁾	0.8106 ⁽¹⁴⁾	0.3756 ⁽¹¹⁾	0.2771 ⁽⁵⁾	0.1958 ⁽²⁾	0.2776 ⁽⁸⁾	0.2557 ⁽⁴⁾	13.3428 ⁽¹³⁾	0.1960 ⁽³⁾	0.2736 ⁽⁶⁾
	$ Bias(\hat{\beta}) $	2.1548 ⁽¹²⁾	2.1366 ⁽¹⁰⁾	5.8409 ⁽¹⁴⁾	2.1520 ⁽¹¹⁾	5.8413 ⁽¹³⁾	1.3933 ⁽²⁾	2.8406 ⁽¹⁰⁾	2.0692 ⁽⁸⁾	1.4796 ⁽⁹⁾	2.0911 ⁽⁶⁾	1.9632 ⁽⁶⁾	1.4190 ⁽¹⁾	1.4570 ⁽⁴⁾	2.0611 ⁽⁷⁾	1.

Table 6: Partial and Overall Ranks for all Estimation Methods of the BE Model from 5

n	Metric	MLE	ADE	CVME	MPSE	OLSE	RTADE	WLSE	LTADE	MSADE	MSALDE	ADSOE	KE	MSSDE	MSSLDE	MSLNDE
15	$ Bias(\theta) $	10.0	9.0	6.0	11.0	4.0	13.0	14.0	7.0	1.0	12.0	5.0	15.0	2.0	8.0	3.0
	$ Bias(\beta) $	14.0	12.0	7.0	13.0	6.0	2.0	15.0	10.0	5.0	11.0	8.0	1.0	4.0	9.0	3.0
	MSE (θ)	10.0	9.0	5.0	11.0	4.0	13.0	14.0	8.0	1.0	12.0	6.0	15.0	2.0	7.0	3.0
	MSE (β)	14.0	12.0	7.0	13.0	6.0	2.0	15.0	10.0	5.0	11.0	8.0	1.0	4.0	9.0	3.0
	MRE (θ)	10.0	9.0	6.0	11.0	4.0	13.0	14.0	7.0	1.0	12.0	5.0	15.0	2.0	8.0	3.0
	MRE (β)	14.0	12.0	7.0	13.0	6.0	2.0	15.0	10.0	5.0	11.0	8.0	1.0	4.0	9.0	3.0
	D_{abs}	3.0	4.0	6.0	1.0	7.0	14.0	5.0	8.0	9.0	10.0	13.0	15.0	12.0	2.0	11.0
	D_{max}	3.0	4.0	6.0	1.0	7.0	14.0	5.0	8.0	9.0	10.0	13.0	15.0	12.0	2.0	11.0
	$\sum Ranks$	78 ⁽¹²⁾	71 ⁽⁹⁾	50 ⁽⁵⁾	74 ⁽¹¹⁾	44 ⁽⁴⁾	73 ⁽¹⁰⁾	97 ⁽¹⁵⁾	68 ⁽⁸⁾	37 ⁽¹⁾	88 ⁽¹⁴⁾	66 ⁽⁷⁾	78 ⁽¹²⁾	42 ⁽³⁾	54 ⁽⁶⁾	40 ⁽²⁾
	$ Bias(\theta) $	9.0	7.0	12.0	10.0	11.0	14.0	13.0	6.0	2.0	8.0	4.0	15.0	3.0	5.0	1.0
50	$ Bias(\beta) $	12.0	10.0	14.0	11.0	13.0	2.0	15.0	8.0	5.0	9.0	6.0	1.0	4.0	7.0	3.0
	MSE (θ)	9.0	7.0	12.0	10.0	11.0	14.0	13.0	6.0	2.0	8.0	4.0	15.0	3.0	5.0	1.0
	MSE (β)	12.0	10.0	14.0	11.0	13.0	2.0	15.0	8.0	5.0	9.0	6.0	1.0	4.0	7.0	3.0
	MRE (θ)	9.0	7.0	12.0	10.0	11.0	14.0	13.0	6.0	2.0	8.0	4.0	15.0	3.0	5.0	1.0
	MRE (β)	12.0	10.0	14.0	11.0	13.0	2.0	15.0	8.0	5.0	9.0	6.0	1.0	4.0	7.0	3.0
	D_{abs}	1.0	4.0	5.0	2.0	6.0	14.0	3.0	7.0	12.0	11.0	13.0	15.0	9.0	8.0	10.0
	D_{max}	1.0	4.0	5.0	2.0	6.0	14.0	3.0	7.0	12.0	11.0	13.0	15.0	9.0	8.0	10.0
	$\sum Ranks$	65 ⁽⁸⁾	59 ⁽⁷⁾	88 ⁽¹⁴⁾	67 ⁽⁹⁾	84 ⁽¹³⁾	76 ⁽¹¹⁾	90 ⁽¹⁵⁾	56 ⁽⁶⁾	45 ⁽³⁾	73 ⁽¹⁰⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	39 ⁽²⁾	52 ⁽⁴⁾	32 ⁽¹⁾
	$ Bias(\theta) $	9.0	7.0	14.0	10.0	13.0	12.0	11.0	5.0	2.0	8.0	4.0	15.0	3.0	6.0	1.0
	$ Bias(\beta) $	12.0	10.0	15.0	11.0	14.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
	MSE (θ)	9.0	7.0	14.0	10.0	13.0	12.0	11.0	5.0	2.0	8.0	4.0	15.0	3.0	6.0	1.0
100	MSE (β)	12.0	10.0	15.0	11.0	14.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
	MRE (θ)	9.0	7.0	14.0	10.0	13.0	12.0	11.0	5.0	2.0	8.0	4.0	15.0	3.0	6.0	1.0
	MRE (β)	12.0	10.0	15.0	11.0	14.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
	D_{abs}	3.0	2.0	6.0	4.0	5.0	14.0	1.0	8.0	10.0	9.0	13.0	15.0	11.0	7.0	12.0
	D_{max}	3.0	2.0	6.0	4.0	5.0	14.0	1.0	8.0	10.0	9.0	13.0	15.0	11.0	7.0	12.0
	$\sum Ranks$	69 ⁽⁹⁾	55 ⁽⁵⁾	99 ⁽¹⁵⁾	71 ⁽¹¹⁾	91 ⁽¹⁴⁾	70 ⁽¹⁰⁾	74 ⁽¹²⁾	58 ⁽⁷⁾	41 ⁽²⁾	66 ⁽⁸⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	44 ⁽³⁾	53 ⁽⁴⁾	35 ⁽¹⁾
	$ Bias(\theta) $	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\beta) $	12.0	10.0	14.0	11.0	15.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
	MSE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MSE (β)	12.0	10.0	14.0	11.0	15.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
200	MRE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MRE (β)	12.0	10.0	14.0	11.0	15.0	2.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	3.0
	D_{abs}	1.0	2.0	4.0	3.0	5.0	14.0	6.0	9.0	10.0	7.0	13.0	15.0	12.0	8.0	11.0
	D_{max}	1.0	2.0	4.0	3.0	5.0	14.0	6.0	9.0	10.0	7.0	13.0	15.0	12.0	8.0	11.0
	$\sum Ranks$	65 ⁽⁸⁾	58 ⁽⁷⁾	86 ⁽¹⁴⁾	69 ⁽¹⁰⁾	94 ⁽¹⁵⁾	76 ⁽¹¹⁾	84 ⁽¹³⁾	66 ⁽⁶⁾	44 ⁽³⁾	56 ⁽⁶⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	42 ⁽²⁾	52 ⁽⁴⁾	34 ⁽¹⁾
	$ Bias(\theta) $	9.0	8.0	13.0	10.0	12.0	14.0	11.0	6.0	2.0	7.0	4.0	15.0	3.0	5.0	1.0
	$ Bias(\beta) $	12.0	10.0	14.0	11.0	15.0	3.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	2.0
	MSE (θ)	9.0	8.0	13.0	10.0	12.0	14.0	11.0	6.0	2.0	7.0	4.0	15.0	3.0	5.0	1.0
	MSE (β)	12.0	10.0	14.0	11.0	15.0	3.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	2.0
	MRE (θ)	9.0	8.0	13.0	10.0	12.0	14.0	11.0	6.0	2.0	7.0	4.0	15.0	3.0	5.0	1.0
300	MRE (β)	12.0	10.0	14.0	11.0	15.0	3.0	13.0	9.0	5.0	8.0	6.0	1.0	4.0	7.0	2.0
	D_{abs}	2.0	4.0	5.0	1.0	6.0	14.0	3.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	D_{max}	2.0	4.0	5.0	1.0	6.0	14.0	3.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	$\sum Ranks$	67 ⁽¹⁰⁾	62 ⁽⁷⁾	91 ⁽¹⁴⁾	65 ⁽⁹⁾	93 ⁽¹⁵⁾	79 ⁽¹³⁾	78 ⁽¹¹⁾	61 ⁽⁶⁾	45 ⁽³⁾	63 ⁽⁶⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	41 ⁽²⁾	50 ⁽⁴⁾	31 ⁽¹⁾
	$ Bias(\theta) $	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\beta) $	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
	MSE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MSE (β)	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
	MRE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MRE (β)	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
500	D_{abs}	2.0	4.0	5.0	1.0	6.0	14.0	3.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	D_{max}	2.0	4.0	5.0	1.0	6.0	14.0	3.0	8.0	12.0	9.0	13.0	15.0	10.0	7.0	11.0
	$\sum Ranks$	67 ⁽¹⁰⁾	62 ⁽⁷⁾	91 ⁽¹⁴⁾	65 ⁽⁹⁾	93 ⁽¹⁵⁾	79 ⁽¹³⁾	78 ⁽¹¹⁾	61 ⁽⁶⁾	45 ⁽³⁾	63 ⁽⁶⁾	56 ⁽⁶⁾	78 ⁽¹²⁾	41 ⁽²⁾	50 ⁽⁴⁾	31 ⁽¹⁾
	$ Bias(\theta) $	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	$ Bias(\beta) $	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
	MSE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MSE (β)	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
	MRE (θ)	9.0	8.0	12.0	10.0	13.0	14.0	11.0	7.0	3.0	6.0	4.0	15.0	2.0	5.0	1.0
	MRE (β)	12.0	10.0	15.0	11.0	14.0	5.0	13.0	9.0	4.0	8.0	6.0	1.0	3.0	7.0	2.0
	D_{abs}	2.0	1.0	3.0	6.0	4.0	14.0	5.0	7.0	10.0	9.0	13.0	15.0	11.0	8.0	12.0
	D_{max}	2.0	1.0	3.0	6.0	4.0	14.0	5.0	7.0	10.0	9.0	13.0	15.0	11.0	8.0	12.0
	$\sum Ranks$	67 ⁽⁹⁾	56 ⁽⁵⁾	87 ⁽¹⁴⁾	75 ⁽¹⁰⁾	89 ⁽¹⁵⁾	85 ⁽¹³⁾	82 ⁽¹²⁾	62 ⁽⁸⁾	41 ⁽³⁾	60 ⁽⁷⁾	56 ⁽⁵⁾	78 ⁽¹¹⁾	37 ⁽²⁾	52 ⁽⁴⁾	33 ⁽¹⁾

methods like MLE, ADE, and WLSE achieve the best ranks (1.0 to 6.0), especially at larger sample sizes, despite their generally lower overall ranking based on the sum of all metrics. This indicates a trade-off where some estimators may be better at fitting the overall distribution shape but are less accurate in their specific parameter estimates (bias and error). The overall rankings, therefore, prioritize the accuracy of the parameter estimates. These inference from Table 4 are due to the simulation results contained in Table 3 and the plots which provide visual illustrations are in Figures 8, 9, 10 and 11. A visual summary of the performance of the estimation methods, ordered by rank, call heatmap for case 3 is in Figure 12.

The summary of partial and overall ranks for the estimation methods of the BE model, as detailed in Table 6, demonstrates clear patterns in estimator performance across different sample sizes, n . Based on the comprehensive $\sum Ranks$, the MSLNDE method is consistently the best overall estimator for larger sample sizes, securing the first rank (rank 1) for $n \geq 50$ and the second rank (rank 2) for $n = 15$. The MSSDE method provides the closest competition, consistently achieving the second rank (rank 2) for $n \geq 200$ and the third rank (rank 3) for $n < 200$. Considering parameter-specific performance reveals that MSLNDE's overall dominance for $n \geq 50$ is primarily attributable to its exceptional accuracy in estimating the θ parameter. It maintains a rank of 1.0 for all $\hat{\theta}$ metrics ($|Bias|(\hat{\theta})$, $MSE(\hat{\theta})$, and $MRE(\hat{\theta})$) when $n \geq 50$. For the β parameter, MSLNDE also performs very well, ranking 2.0 or 3.0 for all $\hat{\beta}$ metrics for $n \geq 50$. The MSSDE method follows a similar pattern for the θ parameter, consistently ranking 2.0 or 3.0 for all $\hat{\theta}$ metrics for

$n \geq 50$. A notable exception to the general trend occurs at the smallest sample size, $n = 15$, where the MSADE method takes the first overall rank (rank 1). At this sample size, MSADE achieves a rank of 1.0 for all $\hat{\theta}$ metrics, outperforming MSLNDE, which ranks 3.0 for these metrics. This suggests that MSADE is particularly effective in small sample situations. The KE method, while consistently ranking poorly overall (typically rank 11 to 13), exhibits superior performance in estimating the β parameter, achieving a rank of 1.0 for all $\hat{\beta}$ metrics for $n \geq 50$.

Conversely, several methods consistently perform poorly across the board, particularly for larger sample sizes. The CVME, OLSE, and WLSE methods, along with RTADE, are frequently found in the lowest overall rank positions (e.g., CVME with rank 15 for $n \geq 100$). When considering the goodness-of-fit metrics, D_{abs} and D_{max} , there is again a divergence from the overall ranking. Here, the MLE, ADE, and MPSE methods often achieve the best ranks (1.0 to 4.0), especially at larger sample sizes, suggesting these methods are better at modeling the distribution's shape even if their specific parameter estimates are less accurate than MSLNDE or MSSDE. The overall ranking, however, suggests that minimizing the error and bias in parameter estimation is given greater weight in the $\sum Ranks$ metric. These inference from Table 6 are due to the simulation results contained in Table 5 and the plots which provide visual illustrations are in Figures 13, 14, 15 and 16. A visual summary of the performance of the estimation methods, ordered by rank, call heatmap for case 5 is in Figure 17.

Table 7: Data I

0.9636	2.7852	3.8628	2.6436	3.0120	2.1780	1.7952	1.9236	1.0176	1.3272	2.9796	2.3520	2.8644	1.0488	1.1244	2.0904
.9852	3.0468	2.4324	2.0088	2.1444	1.9680	0.6228	1.1328	0.8964	1.0008	2.0436	2.4972	2.3556	2.5644	0.9684	2.2452
.9872	1.8420	1.4724	1.3980	1.6176	3.6120	2.6088	0.5436	0.9972	1.6212	1.8540	0.3120	0.5400	1.4844	1.2264	1.0068
.6204	0.9888	1.5948	1.6320	1.3668	1.2876	0.7500	1.9596	1.3944	1.4088	1.6368	1.2360	1.1760	0.9648	0.4200	0.7308
.9768	1.0896	0.9696	0.9072	0.7056	0.3612	0.9648	0.8772	0.7800	0.6192	0.9084	0.6168	0.6972	0.7512	0.5760	5.2956
.6624	5.6340	8.9772	6.2292	4.3596	7.9320	9.8988	6.9984	3.6084	6.5124	3.6732	3.9936	2.0640	3.5124	6.4104	6.0204
.1452	2.6064	3.0852	4.5780	8.7624	4.7412	3.8220	2.1216	3.7956	2.8380	1.9284	5.5704	7.7268	5.2872	6.0252	4.3560
.5904	3.8472	2.0364	2.6544	5.9604	4.7040	5.7300	2.0988	2.2500	4.1808	1.9716	6.0948	4.8648	4.0176	5.1300	1.9368
.4916	3.9744	3.6840	2.9448	2.7960	3.2352	7.2252	5.2176	1.0884	1.9956	3.2436	3.7092	0.6240	1.0800	2.9688	2.4528
.0148	1.2420	1.9788	3.1896	3.2652	2.7336	2.5752	1.5000	3.9204	2.7888	2.8176	3.2748	2.4720	2.3532	1.9308	0.8412
.4628	1.9536	2.1792	1.9392	1.8156	1.4112	0.7224	1.9308	1.7556	1.5600	1.2384	1.8168	1.2348	1.3956	1.5036	1.1532
.2360	2.9304	4.5072	7.1808	4.9836	3.4872	6.3456	7.9188	5.5992	2.8872	5.2104	2.9376	3.1944	1.6512	2.8092	5.1288
.8168	2.5152	2.0844	2.4684	3.6624	7.0092	3.7932	3.0576	1.6968	3.0360	2.2704	1.5432	4.4556	6.1812	4.6764	1.3188
.7068	6.6516	3.8244	3.1848	3.7476	4.5180	5.4912	7.3872	3.4908	3.0804	3.3684	4.1184	3.0912	1.3176	3.4884	4.9176

Table 8: Data II

1.43	0.11	0.71	0.77	2.63	1.49	3.46	2.46	0.59	0.74	1.23	0.94	4.36	0.4	1.74
.73	2.23	2.23	0.7	1.06	1.46	0.3	1.82	2.37	0.63	1.23	1.24	1.24	1.186	1.17

7. APPLICATIONS

Data I consists of 224 observations and represents the mortality rate of patients during the COVID-19 pandemic in Canada. The data is contained in [39] and reported in Table 7.

Data II is the failure interval time for 30 fixed components studied by [40-43] presented in Table 8.

Table 9: Summary of Basic Statistics

Measures	Data I	Data II
n	224	30
Q_1	1.4061	0.7475
Q_3	3.8226	2.1275
IQR	2.4165	1.3800
Outlier	.9772, 7.932,	
	.8988, 8.7624,	4.36, 4.73
	.7268, 7.9188	
Mean	2.9013	1.5552
Median	2.4702	1.2350
Standard Deviation	1.9053	1.1155
Variance	3.6301	1.2444
Range	0.5868	4.6300
Skewness	1.0794	1.3405
Kurtosis	3.8597	4.4203

Table 9 summarizes basic statistics for two datasets, Data I and Data II, which have vastly different sample sizes ($n = 224$ and $n = 30$ respectively). The measures of central tendency, Mean and Median, are higher for Data I (2.9013 and 2.4702) than for Data II (1.5552 and 1.2350). Both datasets show positive Skewness (1.0794 for Data I and 1.3405 for Data II), indicating a distribution with a longer tail on the right side, but this is more pronounced in Data II. The Median is noticeably lower than the Mean in both cases, which is consistent with the positive Skewness. Measures of dispersion, such as Standard Deviation (1.9053 vs 1.1155), Variance (3.6301 vs 1.2444), and Interquartile Range (2.4165 vs 1.3800), are all higher for Data I, suggesting that Data I is more spread out than Data II. Both datasets contain outliers, with Data I having six high values and Data II having two high values. The Kurtosis values (3.8597 for Data I and 4.4203 for Data II) are greater than the standard value

of 3 for a normal distribution, suggesting that both distributions are leptokurtic, meaning they have heavier tails and a sharper peak than a normal distribution, with Data II exhibiting a slightly higher degree of peakedness. Interestingly, the Range for Data I (0.5868) is much smaller than Data II (4.6300), despite Data I having a much larger sample size and greater measures of variability, which is unusual and suggests a potential miscalculation or special context for the Range in Data I, perhaps indicating it only represents the range before accounting for the listed outliers, or it is a typo since the minimum to accommodate the listed outliers must be much larger than 0.5868.

The competing models are the Weibull distribution by [44], the Gumbel distribution by [45], the new generalized logistic-x transformed exponential (NGLXTE) distribution by [42], and the type-I heavy-tailed exponential (TIHTE) distribution by [46].

A comparison of model fitness for Data I and Data II, presented in Table 10, suggests that the BE distribution provides the superior fit for both datasets. For Data I, the BE model yields the highest log-likelihood ($LL = -426.76$) and the lowest values for the information criteria: AIC (857.5228), CAIC (857.5771), BIC (864.3461), and HQIC (860.2770). Generally, lower values for these criteria indicate a better model fit relative to its complexity. Furthermore, the BE model has the smallest Anderson-Darling ($A = 0.4316$) and Cramér-von Mises ($W = 0.0600$) statistics, measures of overall goodness-of-fit, and a small Kolmogorov-Smirnov ($KS = 0.0440$) statistic with a large p-value (0.7795), confirming that the null hypothesis of the data following the BE distribution is not rejected. Among the competing models, the Weibull and KMDUSW distributions offer the next best fit, with similar statistics.

For Data II, the BE model again exhibits the highest log-likelihood ($LL = -39.44$) and the lowest values across all information criteria, such as $AIC = 82.8735$. It also maintains the smallest W (0.0437) and A (0.2616) statistics. The BE model's KS statistic (0.0922) and its associated p-value (0.9606) are strong indicators of its good fit to Data II. The TIHTE model is a strong contender for Data II, showing a very close log-likelihood and information criteria values, and the smallest W and A statistics, though the BE remains the clear overall best performer based on the combined criteria. The NGLXTE model consistently provides the poorest fit for both datasets, as indicated by its lowest log-likelihood and highest information criteria and goodness-of-fit statistics.

Table 10: Measures of Fitness and Model Performance

Data	Model	LL	AIC	CAIC	BIC	HQIC	W	A	KS	P-value
6* I	BE	-426.76	857.5228	857.5771	864.3461	860.2770	0.0600	0.4316	0.0440	0.7795
	Weibull	-427.75	859.5098	859.5641	866.3331	862.2640	0.1313	0.9181	0.0532	0.5503
	Gumbel	-433.14	870.2768	870.3311	877.1001	873.0310	0.1709	1.2285	0.0627	0.3418
	KMDUSW	-427.93	859.852	859.9063	866.6753	862.6062	0.1324	0.9236	0.0527	0.5619
	NGLXTE	-437.21	878.4287	878.483	885.2520	881.1829	0.3433	2.2470	0.0803	0.1115
	TIHTE	-431.04	866.0753	866.1296	872.8986	868.8295	0.0625	0.4154	0.0749	0.1615
6* II	BE	-39.44	82.8735	83.3179	85.6759	83.7700	0.0437	0.2616	0.0922	0.9606
	Weibull	-39.76	83.5171	83.9616	86.3195	84.4136	0.0606	0.3613	0.1160	0.8140
	Gumbel	-40.25	84.4974	84.9419	87.2998	85.3939	0.0577	0.3456	0.1235	0.7498
	KMDUSW	-39.67	83.3498	83.7943	86.1522	84.2463	0.0578	0.3429	0.1118	0.8475
	NGLXTE	-41.32	86.6406	87.0851	89.4430	87.5371	0.1136	0.6924	0.1622	0.4092
	TIHTE	-39.64	83.2799	83.7243	86.0827	84.1764	0.0357	0.2105	0.0357	0.2105

Table 11: Maximum Likelihood Estimates and Standard Errors

Data	Model	Scale	Shape
I	BE	.4548	.5865
		(0.4413)	(0.4578)
	Weibull	.2551	.6150
		(0.1424)	(0.0823)
	Gumbel	.0562	.3771
		(0.0965)	(0.0755)
	KMDUSW	.5751	.1566
		(0.0813)	(0.0210)
	NGLXTE	.8620	.1967
		(0.0467)	(0.0081)
	TIHTE	.1159	.5842
		(0.0353)	(0.9727)
II	BE	.7145	.3340
		(0.0588)	(0.1732)
	Weibull	.7293	.4987
		(0.2226)	(0.2063)
	Gumbel	.0847	.7541
		(0.1441)	(0.1128)
	KMDUSW	.4657	.4515
		(0.2036)	(0.1103)
	NGLXTE	.9363	.3633
		(0.1387)	(0.0439)
	TIHTE	.1804	.4521
		(0.1179)	(2.6624)

Table 11 presents the maximum likelihood estimates (MLEs) and their standard errors for the scale and shape parameters across all six models for Data I and Data II. A comparison of the point estimates reveals significant variability in the parameter values between the different models and datasets, reflecting distinct distributional characteristics. For Data I, the Weibull model has the largest scale parameter estimate (3.2551), while the TIHTE model exhibits the largest shape parameter estimate (3.5842). When evaluating the precision of these estimates via their standard errors, the NGLXTE model consistently shows the most precise estimation of the shape parameter across both datasets, with notably small standard errors (0.0081 for Data I and 0.0439 for Data II). Conversely, the TIHTE model for Data II shows a considerable lack of precision, particularly for its shape parameter, which has an exceptionally large standard error of 2.6624 relative to its estimate. Overall, the standard errors are generally smaller for Data II, indicating more precise parameter estimation for that dataset, with the BE model having the most precisely estimated scale parameter (0.7145, standard error 0.0588).

The BE distribution's shape parameter is a primary indicator of its failure rate (hazard function) behavior over time. The estimated shape parameters $\hat{\beta}$ for the BE model in Table 1 are both greater than 1, leading to specific conclusions regarding risk. The general trend is Increasing Failure Rate (IFR) For both Data I ($\hat{\beta}=1.5865$) and Data II ($\hat{\beta}=1.3340$), the shape estimates are >1 . This strongly suggests that the underlying failure process is governed by an Increasing Failure Rate (IFR) trend, where the risk of failure rises as the component or subject ages, characteristic of a wear-out phenomenon. The hazard shape is an upside-down bathtub. A shape parameter $\hat{\beta} > 1$ in the BE

model typically corresponds to an upside-down bathtub (unimodal) hazard function. This implies that the risk of failure:

- Initially increases rapidly from time $t = 0$.
- Reaches a peak level.
- Gradually decreases at very late ages.

This pattern often models initial high-stress or random defect failures that peak during the early to mid-life span of the system.

The estimates for Data II are more precise than those for Data I, providing greater confidence in the inferred risk behavior for the second dataset.

Figure 18a for Data I shows a distribution that is unimodal and slightly skewed, centered around a median of approximately 3.5, with the bulk of the data ranging from about 0 to 8 and a maximum around 12. Figure 18b for Data II displays a broader, more platykurtic distribution, also unimodal, centered near a median of 1.5, suggesting greater variability compared to Data I. Data I appears to have a higher median and less spread than Data II, whose data points are more dispersed across the range from roughly -1 to 6.

These plots compare the histograms of two datasets, Data I and Data II, against the probability density functions of six different statistical distributions. In Figure 19a, the histogram for Data I is roughly unimodal and slightly skewed right, with the BE, Weibull, Gumbel, and NGLXTE distributions appearing to provide a visually closer fit to the data's density than the others. Figure 19b shows that the histogram for Data II is also unimodal but less skewed than Data I, and the BE, KMDUSW and Gumbel distributions seem to align reasonably well with the shape of the observed data. Both plots illustrate how various theoretical

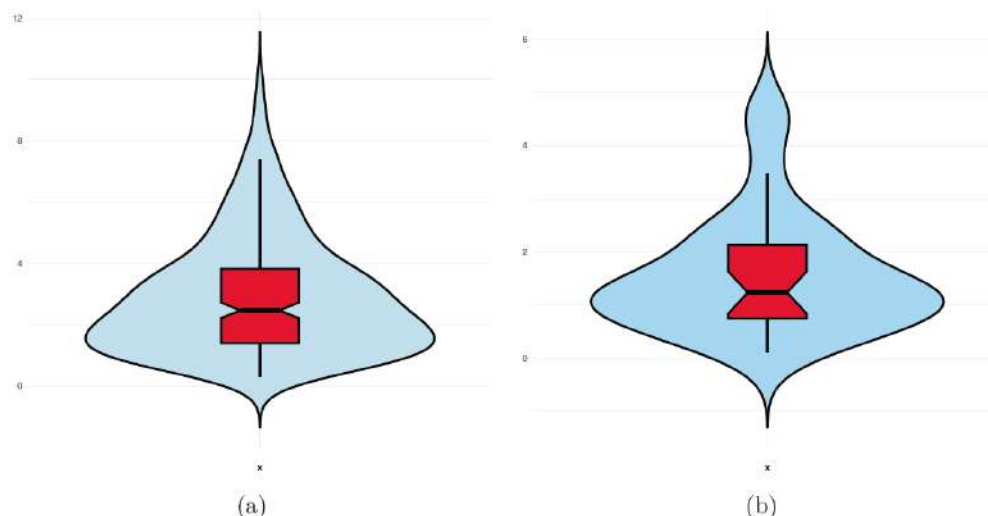


Figure 18: Boxplot in Violin of (a) Data I (b) Data II.

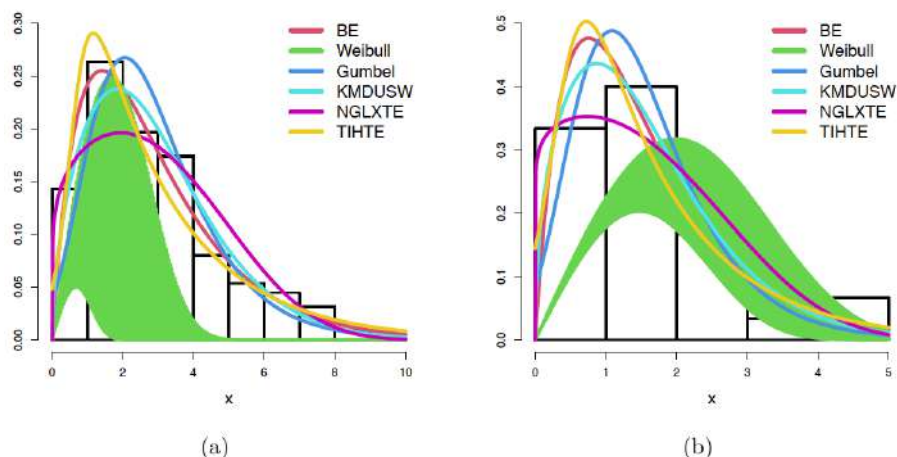


Figure 19: Histogram vs Density of (a) Data I (b) Data II.

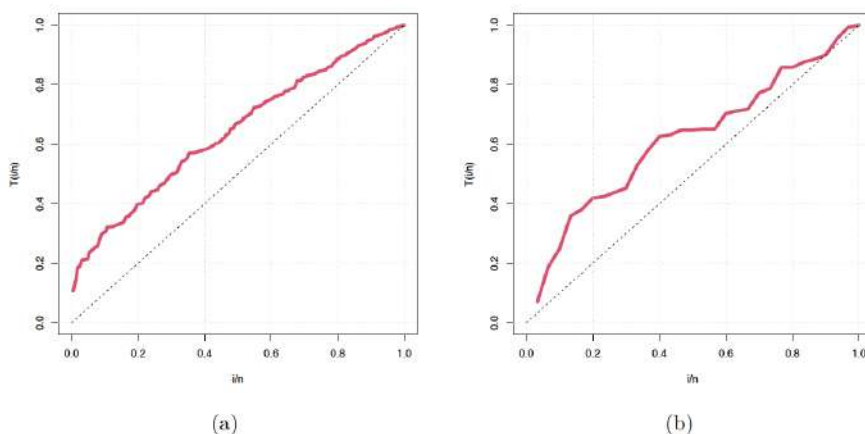


Figure 20: Total time on test (TTT) plots for (a) Data I (b) Data II.

distributions attempt to model the underlying frequency or density of the respective datasets.

Figure 20a for Data I shows the Total Time on Test (TTT) plot is mostly concave, remaining entirely above the diagonal line, which suggests a distribution with an increasing failure rate or Increasing Failure Rate (IFR) property. Figure 20b for Data II exhibits an S-shaped curve that is initially convex and then becomes concave, crossing the diagonal, indicating a distribution with a bathtub-shaped failure rate or a Decreasing Failure Rate (DFR) followed by an IFR. Since both plots lie above the diagonal for a significant portion, both distributions generally suggest an increasing failure rate over time, especially Data I, which shows a consistently increasing risk. These support the hazard profile in Figure 20b, hence validating the use of BE distribution in fitting the two datasets.

Figure 21 presents the profile log-likelihood plots for the parameters θ and β based on Data I. For both parameters, the profile log-likelihood curve is nearly flat across the displayed ranges, meaning the value of the negative log-likelihood changes very little for a wide

range of parameter values. The maximum likelihood estimates (MLEs), marked by the red dots, are approximately $\hat{\theta} = 2.1$ and $\hat{\beta} = 0.6$, and due to the flatness, the resulting 95% confidence intervals are extremely wide or essentially unbounded within the plotted regions, as the flat line never reaches the critical dashed threshold (which is not visible for θ or β).

The two plots in Figure 22 show the profile negative log-likelihood for parameters θ and β , respectively. The maximum likelihood estimate for θ is approximately 2.0 and for β is approximately 1.0, as indicated by the lowest point on each flat profile. Since the profile negative log-likelihood is constant and below the dashed line for the 95% confidence level (95% CL) across the plotted range, the interval is not closed and the data provides very little information to constrain these parameters.

The Probability-Probability (P-P) plots in Figure 23a show that the empirical data points for Data I are generally close to the 45° dashed reference line for

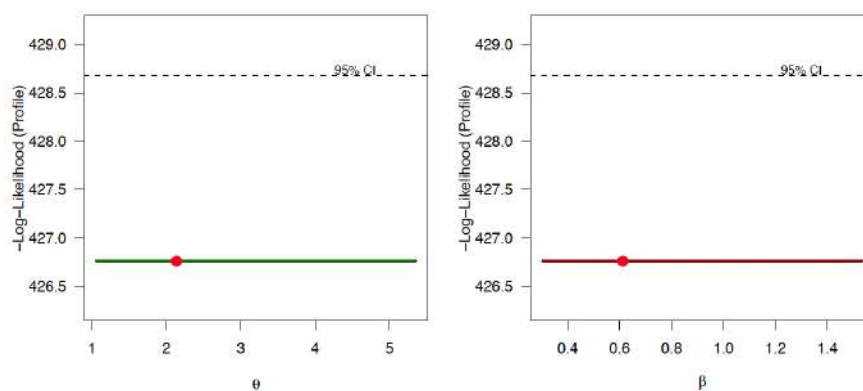


Figure 21: Log-likelihood Profile for Data I.

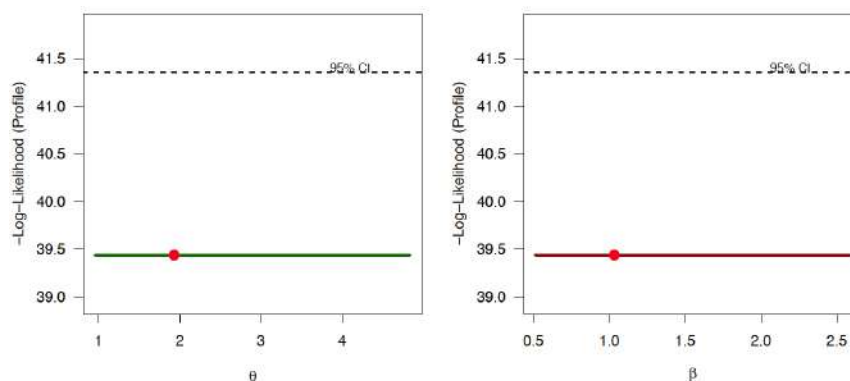


Figure 22: Log-likelihood Profile for Data II.

most of the distributions, suggesting a reasonable fit for Data I by several of the models. The P-P plots in Figure 23b for Data II, however, show that the data points deviate noticeably from the reference line for all distributions except the BE, with a few plots like NGLXTE showing particularly poor alignment. These

deviations indicate that none of the tested distributions provides a good fit for Data II.

The Quantile-Quantile (Q-Q) plots in Figure 24a show that the empirical quantiles for Data I generally align well with the theoretical quantiles for most of the

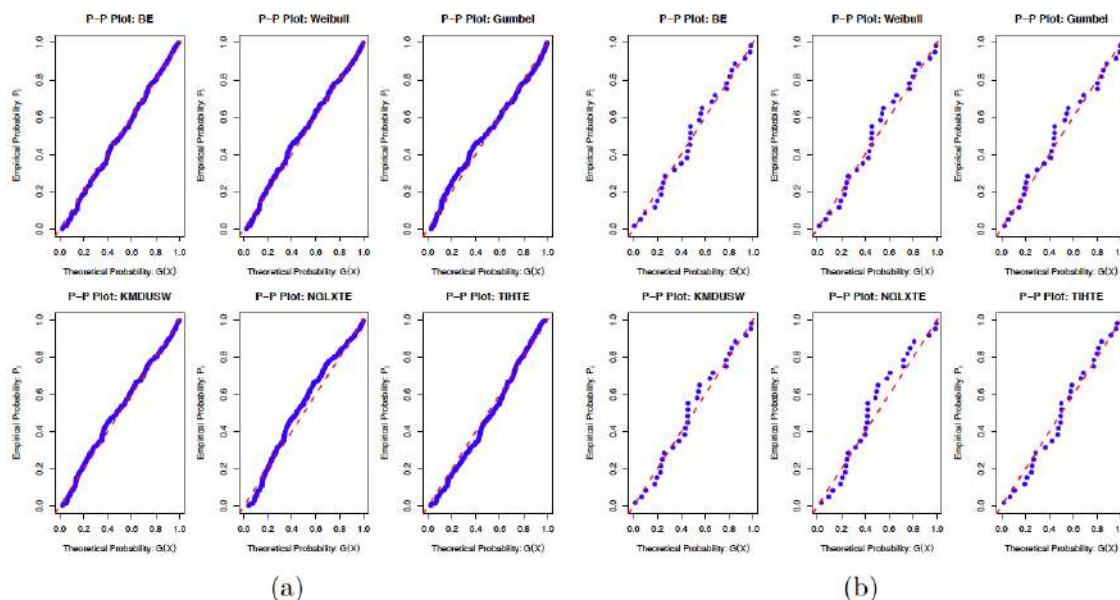


Figure 23: Probability-Probability (P-P) plots for (a) Data I (b) Data II.

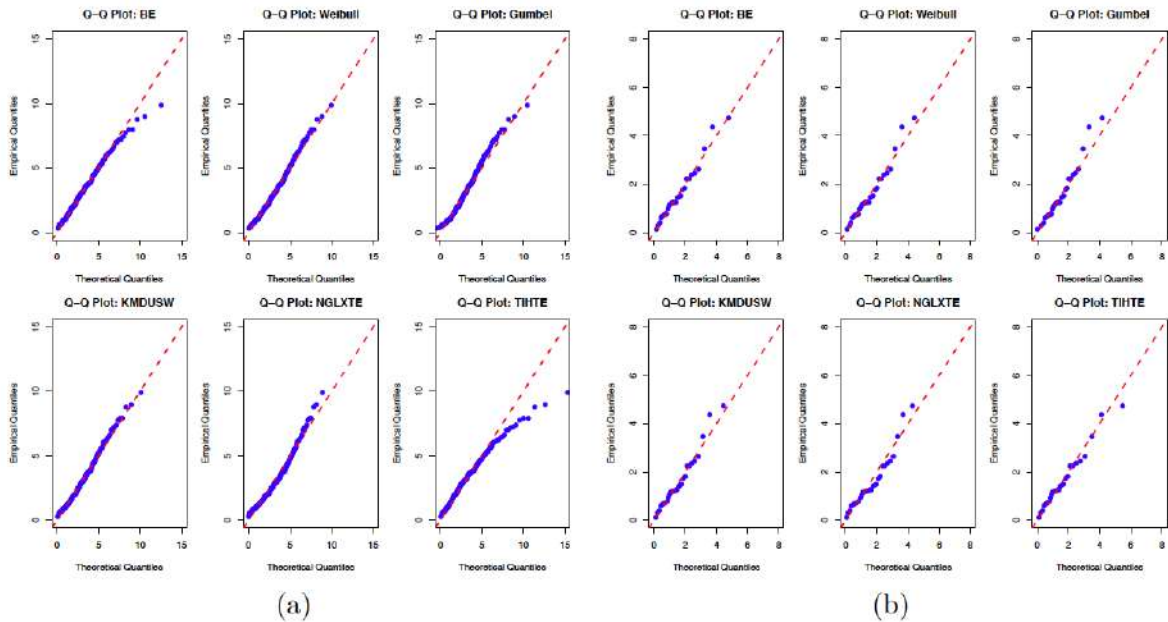


Figure 24: Quantile-Quantile (Q-Q) plots for (a) Data I (b) Data II.

distributions, indicating that several of the models could be plausible fits for this data. The Q-Q plots in Figure 24b for Data II, however, display noticeable curvature and deviation from the dashed 45° line for all the distributions, particularly at the tails. This systematic lack of alignment confirms that the tested distributions, which were also poorly supported by the P-P plots in Figure 24b, are not appropriate models for Data II.

8. CONCLUDING REMARKS

This research successfully introduced and characterized the Bilal-G family of distributions, a significant advancement in the development of flexible statistical models capable of capturing complex data characteristics that traditional distributions fail to address. By utilizing the Bilal distribution as a generator, we derived the new family's key theoretical properties and specifically focused on the two-parameter Bilal-Exponential (BE) distribution. The analysis confirmed the BE distribution's versatility, showing its ability to model diverse shapes and exhibit a consistently increasing failure rate, which was validated by the TTT plots of the real-world data used for application. The rigorous simulation study provided critical insights into the best-performing estimation methods; among the fifteen tested non-Bayesian estimators, the Minimum Spacing Linex Distance (MSLNDE) method demonstrated superior performance, consistently outranking all others in terms of minimizing bias and error across various sample sizes, thus providing a definitive recommendation for future parameter estimation involving the BE model.

Most importantly, the application to two distinct real datasets unequivocally established the BE distribution as the superior model fit compared to five well-known competing distributions, evidenced by its lowest AIC, BIC, and goodness-of-fit statistics. In summary, the Bilal-G family, through its BE sub-model, provides a powerful and practical new tool for researchers needing greater modeling flexibility in applied statistics.

8.1. Limitations and Future Work

While this study successfully introduced the Bilal-G family of distributions and provided a rigorous assessment of the Bilal-Exponential (BE) sub-model's properties, applicability, and estimation efficiency, it is important to acknowledge certain boundaries of the current work that naturally suggest avenues for future research.

1. **Scope of Estimation:** The comprehensive simulation study was confined exclusively to non-Bayesian estimation techniques. While the Maximum Likelihood and Minimum Spacing estimators proved superior, the potential robustness and performance improvements offered by Bayesian methods remain unexplored.
2. **Univariate Focus:** The development and application of the Bilal-G family were restricted to the univariate setting, which limits its utility in modeling complex systems where components or events are inherently dependent.

3. Regression Modeling: The practical application was focused solely on fitting the distribution to complete, non-censored data. The potential of the BE distribution to serve as a baseline model in regression or survival analysis frameworks for handling time-to-event data with covariates was not addressed.

8.2. Future Work

Building upon the robust theoretical foundation established here, the following directions are recommended for subsequent research:

1. Bayesian Inference: Develop and evaluate Bayesian estimation methodologies for the parameters of the BE distribution, incorporating various prior distributions (e.g., non-informative, gamma, or inverted-gamma priors) and utilizing Markov Chain Monte Carlo (MCMC) techniques for practical implementation. Multivariate Extensions: Derive and analyze bivariate and multivariate extensions of the Bilal-G family to facilitate the modeling of correlated data structures in fields such as engineering, finance, and biostatistics.
2. Bilal-G Regression: Propose and test a Bilal-G regression model suitable for analyzing censored data, which would significantly enhance the family's applicability in survival analysis and reliability studies.
3. Alternative Baseline Distributions: Explore the application of the Bilal-G generator with other baseline distributions, particularly those that are discrete (e.g., Bilal-Poisson) or possess unique characteristics such as heavy tails or zero-inflation, to further diversify its modeling power.
4. Alternative Estimation Techniques: Investigate the performance of other robust estimation methods not covered here, such as L-Moments or modified maximum likelihood methods, to ensure stability under different data conditions.

REFERENCES

- [1] Johnson NL. Systems of frequency curves generated by methods of translation. In: *Biometrika* 1949; 36(1/2): 149-176. <https://doi.org/10.2307/2332539>
- [2] Singh SK, Maddala GS. A function for the size distribution of expenditure. In: *Econometrica* 1976; 44: 963-970.
- [3] Mudholkar GS, Srivastava DK. Exponentiated Weibull family for analyzing bathtub failure-rate data. In: *IEEE Transactions on Reliability* 1993; 42(2): 299-302. <https://doi.org/10.1109/24.229504>
- [4] Marshall AW, Olkin I. A new method for adding a parameter to a family of distributions with application to the exponential and Weibull families. In: *Biometrika* 1997; 84(3): 641-652. <https://doi.org/10.1093/biomet/84.3.641>
- [5] Eugene N, Lee C, Famoye F. Beta-normal distribution and its applications. In: *Communications in Statistics-Theory and Methods* 2002; 31(4): 497-512. <https://doi.org/10.1081/STA-120003130>
- [6] Cordeiro GM, Castro MD. A new family of generalized distributions. In: *Journal of Statistical Computation and Simulation* 2011; 81(7): 883-898. <https://doi.org/10.1080/00949650903530745>
- [7] Alzaatreh A, Lee C, Famoye F. A new method for generating families of continuous distributions. In: *Metron* 2013; 71(1): 63-79. <https://doi.org/10.1007/s40300-013-0007-y>
- [8] Weibull W. A statistical distribution function of wide applicability. In: *Journal of applied mechanics* 1951. url: <https://hal.science/hal-03112318/document>.
- [9] Lomax KS. Business failures: Another example of the analysis of failure data. In: *Journal of the American Statistical Association* 1954; 49(268): 847-852. <https://doi.org/10.1080/01621459.1954.10501239>
- [10] Lindley DV. Fiducial distributions and Bayes' theorem. In: *Journal of the Royal Statistical Society. Series B (Methodological)* 1958; 102-107. url: <https://www.jstor.org/stable/2983909>.
- [11] Obulezi OJ, Oguadimma EE, Igboke CP, Chikwelu CU, Salem GS. A New Unit Family of Distributions with Applications to Public Health, Environmental, and Financial Data. In: *Appl Math* 2025; 19(6): 1437-1461. <https://doi.org/10.18576/amis/190617>
- [12] Obulezi OJ, Obiora-Iloanoa HO, Osujia GA, Kayidb M, Balogunc OS. Weibull Sine Generalized Distribution Family: Fundamental Properties, Sub-model, Simulations, with Biomedical Applications. In: *Electronic Journal of Applied Statistical Analysis* 2025; 18(01): 183-212. <https://doi.org/10.1285/i20705948v18n1p183>
- [13] Obisue EI, Okoro CN, Obulezi OJ, Abdalla GSS, Almetwally EM, Faal A, Elgarhy M. Kavya-Manoharan DUS Family of Distributions with Diverse Applications. In: *Journal of Statistics Applications & Probability* 2025; 14(6). <https://doi.org/10.18576/jasp/140608>
- [14] Mir AA, Ur Rasool S, Ahmad SP, Bhat AA, Jawa TM, Sayed-Ahmed N, Tolba AH. A robust framework for probability distribution generation: Analyzing structural properties and applications in engineering and medicine. In: *Axioms* 2025; 14(4): 281. <https://doi.org/10.3390/axioms14040281>
- [15] Orji GO, Etaga HO, Almetwally EM, Igboke CP, Aguwa OC, Obulezi OJ. A new odd reparameterized exponential transformed-x family of distributions with applications to public health data. In: *Innovation in Statistics and Probability* 2025; 1(1): 88-118. <https://doi.org/10.64389/isp.2025.01107>
- [16] El-Saeed AR, Obulezi OJ, El-Raouf MMA. Type II heavy tailed family with applications to engineering, radiation biology and aviation data. In: *Journal of Radiation Research and Applied Sciences* 2025; 18(3): 101547. <https://doi.org/10.1016/j.jrras.2025.101547>
- [17] Elgarhy M, Abdalla GSS, Obulezi OJ, Almetwally EM. A tangent DUS family of distributions with applications to the Weibull model. In: *Scientific African* 2025; e02843. <https://doi.org/10.1016/j.sciaf.2025.e02843>
- [18] Oramulu DO, Alsadat N, Kumar A, Bahloul MM, Obulezi OJ. Sine generalized family of distributions: Properties, estimation, simulations and applications. In: *Alexandria Engineering Journal* 2024; 109: 532-552. <https://doi.org/10.1016/j.aej.2024.09.001>

- [19] Mohammad S, Gaire AK. Exponential Arctan-G Family of Distribution With Properties and Applications. In: Journal of Probability and Statistics 2025; 2025(1): 3643496. <https://doi.org/10.1155/jpas/3643496>
- [20] Kumar A, Yadav AS. A new class of ArcSin Topp- Leone probability distribution: classical estimation and its application to real data. In: Life Cycle Reliability and Safety Engineering 2025; 1-20. <https://doi.org/10.1007/s41872-025-00349-y>
- [21] Alkhairy I, Nagy M, Muse AH, Hussam E. The Arctan-X Family of Distributions: Properties, Simulation, and Applications to Actuarial Sciences. In: Complexity 2021; 2021(1): 4689010. <https://doi.org/10.1155/2021/4689010>
- [22] Christophe CC. Theory on a new bivariate trigonometric Gaussian distribution. In: Innovation in Statistics and Probability 2025; 1(2): 1-17. <https://doi.org/10.64389/isp.2025.01223>
- [23] Mousa MN, Moshref ME, Youns N, Mansour MMM. Inference under Hybrid Censoring for the Quadratic Hazard Rate Model: Simulation and Applications to COVID-19 Mortality. In: Modern Journal of Statistics 2025; 2(1): 1-31. <https://doi.org/10.64389/mjs.2026.02113>
- [24] Ragab I, Elgarhy M. Type II half logistic Ailamujia distribution with numerical illustrations to medical data. In: Computational Journal of Mathematical and Statistical Sciences 2025; 4(2): 379-406. <https://doi.org/10.21608/cjmss.2025.346849.1095>
- [25] Orji GO, Etaga HO, Almetwally EM, Igbokwe CP, Aguwa OC, Obulezi OJ. A New Odd Reparameterized Exponential Transformed-X Family of Distributions with Applications to Public Health Data. In: Innovation in Statistics and Probability 2025; 1(1): 88-118. <https://doi.org/10.64389/isp.2025.01107>
- [26] Gemeay AM, Moakofi T, Balogun OS, Ozkan E, Md. Hossain M. Analyzing Real Data by a New Heavy-Tailed Statistical Model. In: Modern Journal of Statistics 2025; 1(1): 1-24. <https://doi.org/10.64389/mjs.2025.01108>
- [27] Hassan A, Metwally DS, Elgarhy M, Gemeay AM. A new probability continuous distribution with different estimation methods and application. In: Computational Journal of Mathematical and Statistical Sciences 2025; 4(2): 512-532.
- [28] Abd-Elrahman AM. Utilizing ordered statistics in lifetime distributions production: a new lifetime distribution and applications. In: Journal of probability and statistical science 2013; 11(2): 153-164.
- [29] Fisher RA. On an absolute criterion for fitting frequency curves. In: Messenger of Mathematics 1912; 41: 155-156.
- [30] Fisher RA. On the mathematical foundations of theoretical statistics. In: Philosophical transactions of the Royal Society of London. Series A, Containing Papers of a Mathematical or Physical Character 1922; 222(594-604): 309-368. <https://doi.org/10.1098/rsta.1922.0009>
- [31] Anderson TW, Darling DA. Asymptotic theory of certain "goodness of fit" criteria based on stochastic processes. In: The Annals of Mathematical Statistics 1952; 193-212. url: <https://www.jstor.org/stable/2236446>
- [32] Anderson TW, Darling DA. A test of goodness of fit. In: Journal of the American Statistical Association 1954; 49(268) 765-769. <https://doi.org/10.1080/01621459.1954.10501232>
- [33] Choi K, Bulgren WG. An estimation procedure for mixtures of distributions. In: Journal of the Royal Statistical Society: Series B (Methodological) 1968; 30(3): 444-460. <https://doi.org/10.1111/j.2517-6161.1968.tb00743.x>
- [34] Swain JJ, Venkatraman S, Wilson JR. Least-squares estimation of distribution functions in Johnson's translation system. In: Journal of Statistical Computation and Simulation 1988; 29(4): 271-297. <https://doi.org/10.1080/00949658808811068>
- [35] Mukhtar MS, El-Morshedy M, Eliwa MS, Yousof HM. Expanded Fr chet model: mathematical properties, copula, different estimation methods, applications and validation testing. In: Mathematics 2020; 8(11): 1949. <https://doi.org/10.3390/math8111949>
- [36] Aguilar GAS, Moala FA, Cordeiro GM. Zero-truncated poisson exponentiated gamma distribution: Application and estimation methods. In: Journal of Statistical Theory and Practice 2019; 13: 1-20. <https://doi.org/10.1007/s42519-019-0059-2>
- [37] Kao JHK. Computer methods for estimating Weibull parameters in reliability studies. In: IRE Transactions on Reliability and Quality Control 1958; 15-22. <https://doi.org/10.1109/IRE-PGRQC.1958.5007164>
- [38] Torabi H. A general method for estimating and hypotheses testing using spacings. In: Journal of Statistical Theory and Applications 2008; 8(2): 163-168.
- [39] Ahmad Z, Almaspoor Z, Khan F, El-Morshedy M. On predictive modeling using a new flexible Weibull distribution and machine learning approach: Analyzing the COVID-19 data. In: Mathematics 2022; 10(11): 1792. <https://doi.org/10.3390/math10111792>
- [40] Prabhakar Murthy DN, Min Xie, and Renyan Jiang. Weibull models. John Wiley & Sons, 2004.
- [41] Al-Omari AI, Dobbah SA. On the mixture of Shanker and gamma distributions with applications to engineering data. In: Journal of Radiation Research and Applied Sciences 2023; 16(1): 100533. <https://doi.org/10.1016/j.jrras.2023.100533>
- [42] Obulezi OJ, Obiora-Ilouno HO, Osuji GA, Kayid M, Balogun OS. A new family of generalized distributions based on logistic-x transformation: sub-model, properties and useful applications. In: Research in Statistics 2025; 3(1): 2477232. <https://doi.org/10.1080/27684520.2025.2477232>
- [43] Aforka KF, Semary HE, Onyeagu SI, Etaga HO, Obulezi OJ, Al-Moisheer AS. A New Exponentiated Power Distribution for Modeling Censored Data with Applications to Clinical and Reliability Studies. In: Symmetry 2025; 17(7): 1153. <https://doi.org/10.3390/sym17071153>
- [44] Weibull W. A statistical theory of strength of materials. In: IVB-Handl 1939. url: <https://cir.nii.ac.jp/all?q=IVB-Handl>.
- [45] Gumbel EJ. Statistics of Extremes. Columbia University Press 1958.
- [46] Nwankwo BC, Obiora-Ilouno HO, Almulhim FA, Mustafa MSA, Obulezi OJ. Group acceptance sampling plans for type-I heavy-tailed exponential distribution based on truncated life tests. In: AIP Advances 2024; 14(3): <https://doi.org/10.1063/5.0194258>

Received on 05-10-2025

Accepted on 06-11-2025

Published on 24-11-2025

<https://doi.org/10.6000/1929-6029.2025.14.65>  2025 Udobi *et al.*

This is an open-access article licensed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the work is properly cited.