



Statistical Inference on Type II Topp Leone Generalized Inverted Exponential Distribution with Some Applications on Real Data

Zakeia A. Al-saiary^{1,*}, Sarah H. Al-jadaani¹

¹ *Department of Mathematics and Statistics, College of Science, University of Jeddah, Jeddah 22254, Saudi Arabia*

Abstract. In this article, different Bayes techniques are applied to derive the estimators of the additional shape parameter (α) of Type II Topp Leone Generalized Inverted Exponential (TI-ITLGIE) distribution in the case of complete samples. Two cases were adopted: (a) Bayesian estimation with non-informative priors (Jeffrey's prior and Quasi prior under three different risk functions: Relative Quadratic loss function (Rq), Precautionary loss function (P) and Entropy loss function (E)). (b) Bayesian estimation with informative prior depending on standard Bayesian technique is derived under three types of loss functions: Squared Error loss function (SE), Linear Exponential loss function (LINEX), and General Entropy loss function (GE). Mont Carlo simulation study was conducted to find the estimators and compare between these techniques and the Non-Bayesian techniques. Some results are given and showed the best estimation method. Numerical computations and real data sets are given to illustrate these results.

2020 Mathematics Subject Classifications: 62F15, 62C10, 62P20

Key Words and Phrases: Bayesian estimation, generalized inverted exponential distribution, Type II Topp-Leone, loss function, Jeffrey's prior and Quasi prior, Mont Carlo simulation

1. Introduction

Some statistical models are still of interest to statisticians. One from these models was Generalized Inverted Exponential (GIE) distribution witch proposed by [1]. It is a flexible model because it has various shapes of the hazard function. There are many generalizations of GIE distribution has discussed by a statisticians, see [2],[3]. Type II Topp Leone generalized inverted exponential distribution is a generalization of GIE distribution witch deduced by [4]. They derived the MLE of its parameters. They proved that the TIITLGIE distribution is more flexible than the baseline GIE distribution. There are some researchers studied the estimation parameters of GIED. [5] used different methods

*Corresponding author.

DOI: <https://doi.org/10.29020/nybg.ejpam.v18i4.5860>

Email addresses: zaalseare@uj.edu.sa (Z. Al-saiary),
2000277@uj.edu.sa (S. Al-jadaani)

of estimation to estimate the unknown parameters of GIE such as MLE, least square and weighted least square methods. Else, [6] derived Bayesian estimators for the shape parameter of GIE distribution using normal, Lindley's, and Tierney and Kadane's approximation methods. They applied the Monte Carlo simulations for comparison between various estimates. Moreover, [7] estimated the maximum likelihood and Bayes methods of generalized inverted exponential distribution. [8] obtained the MLEs and asymptotic confidence intervals for the unknown parameters of the generalized inverted exponential distribution. On the other hand, [9] proposed a new continuous distribution called Topp Leone distribution, this distribution is attractive as a generator. Also, many authors were interested in this distribution. The TL distribution had not received much attention until by [10] discovered it. The Topp Leone (TL-G) family was proposed by [11]. They introduced the general expression for density and distribution functions for this family. Furthermore, [12] deduced The power inverted Topp Leone power series (PITLPS) class is a novel three-parameter version of the power inverted Topp Leone (PITL) distribution. This compounding process allows for the creation of flexible classes of distributions with strong physical interpretations in fields such as biology and engineering. [13] proposed Topp-Leone Kumaraswamy Marshall-Olkin. They obtained the Shannon entropy, order statistics, the quantile power series, and several associated measures and functions. In this article we will find the Bayesian estimators of the shape parameter of the TIITL-GIE distribution using different techniques. We will compare between those techniques and the MLE using Monte Carlo simulations and Mathematica program depends on the mean square error and the bias. Finally, verify these results using real data sets from different fields.

The cumulative distribution function (CDF) and probability density function (PDF) of TIITLGIE distribution with three parameters $(\alpha, \lambda, \theta)$ is introduced by [4]

$$F(x) = 1 - [1 - [1 - [1 - e^{-\frac{\lambda}{x}}]^{\theta}]^2]^{\alpha}, x > 0, \lambda, \theta, \alpha > 0, \quad (1)$$

and

$$f(x) = \frac{2\alpha\theta\lambda}{x^2} e^{-\frac{\lambda}{x}} [1 - e^{-\frac{\lambda}{x}}]^{\theta-1} [1 - [1 - e^{-\frac{\lambda}{x}}]^{\theta}] [1 - [1 - [1 - e^{-\frac{\lambda}{x}}]^{\theta}]^2]^{\alpha-1}, x > 0, \lambda, \theta, \alpha > 0, \quad (2)$$

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

The quantile function of random variable X of TIITL-GIED is given by:

$$x = Q(u) = \frac{-\lambda}{\log[1 - [1 - [1 - [1 - u]^{\frac{1}{\alpha}}]^{\frac{1}{2}}]^{\frac{1}{\theta}}]}, \quad 0 < u < 1. \quad (3)$$

2. Maximum Likelihood Estimator for shape parameter α

Let $x_1, x_2, \dots, x_n$ be a random sample of size n from TIITLGIE distribution the likelihood function $L(\underline{x}, \alpha, \theta, \lambda)$ is written as follows:

$$L(\underline{x}) = (2\alpha\theta\lambda)^n \prod_{i=1}^n \frac{1}{x_i^2} e^{\sum_{i=1}^n \frac{-\lambda}{x_i}} \prod_{i=1}^n (1 - e^{\frac{-\lambda}{x_i}})^{\theta-1} \prod_{i=1}^n [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}] \times \prod_{i=1}^n [1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{\alpha-1}. \quad (4)$$

Therefore, the log-likelihood function is given as follows

$$\begin{aligned} \log(L) = & n\log(2) + n\log(\alpha) + n\log(\theta) + n\log(\lambda) - 2 \sum_{i=1}^n \log(x_i) - \sum_{i=1}^n \frac{\lambda}{x_i} + (\theta - 1) \\ & \times \sum_{i=1}^n \log(1 - e^{\frac{-\lambda}{x_i}}) + \sum_{i=1}^n \log(1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}) + (\alpha - 1) \sum_{i=1}^n \log(1 - (1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta})^2). \end{aligned}$$

Thus on solving $\frac{\partial}{\partial \alpha}(\log(L)) = 0$, we get the ML estimator of shape parameter as

$$\hat{\alpha} = \left[-\frac{1}{n} \sum_{i=1}^n \log(1 - (1 - (1 - e^{\frac{-\lambda}{x_i}})^{\hat{\theta}})^2) \right]^{-1} \quad (5)$$

3. Bayesian Estimators of the TIITL-GIE distribution for shape parameter α

In this section, Bayesian techniques are discussed for estimating the shape parameter α of TIITLGIE distribution under the assumption that the both parameters λ and θ are known based on the complete samples. These estimators are derived using two cases: In the first case we used Bayesian estimation with non-informative priors (Jeffrey's prior and Quasi prior under three different risk functions: Relative Quadratic loss function (Rq), Precautionary loss function (P) and Entropy loss function (E)). In the second case, we derive the standard Bayesian estimators with Gamma prior under three types of loss functions: squared error loss function, linear exponential loss function and general entropy loss function.

3.1. Case 1: Bayesian estimation with non-informative priors

(i) Posterior Distribution under Jeffrey's prior

Let $x_1, x_2, \dots, x_n$ be a random sample of size n having TIITL-GIED, the Jeffrey's

non-informative prior for parameter α is given by

$$g_1(\alpha) \propto \frac{1}{\alpha}, \quad \alpha > 0. \quad (6)$$

By using the Bayes theorem, we get

$$\pi_1(\alpha|\underline{x}) \propto L(x|\alpha)g_1(\alpha). \quad (7)$$

By substituting Equation (4) and Equation (6) in Equation (7), we have

$$\pi_1(\alpha|\underline{x}) = A_1 \alpha^{n-1} \exp[-\alpha \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{-1}].$$

$$\pi_1(\alpha|\underline{x}) = A_1 \alpha^{n-1} e^{-\alpha T}, \quad (8)$$

where

$$T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{-1}. \quad (9)$$

And A_1 is the normalizing constant and

$$A_1^{-1} = \int_0^{\infty} \alpha^{n-1} e^{-\alpha T} d\alpha,$$

$$A_1^{-1} = \frac{\Gamma(n)}{T^n}.$$

Therefore

$$A_1 = \frac{T^n}{\Gamma(n)}.$$

Using the value of A_1 in Equation (8), we have

$$\pi_1(\alpha|\underline{x}) = \frac{T^n}{\Gamma(n)} \alpha^{n-1} e^{-\alpha T}. \quad (10)$$

Now, we will estimate (α) with Jeffrey's prior under three Loss Functions as following:

- The Relative Quadratic (Rq) Loss Function was introduced by [14] is given by

$$r(\hat{\alpha}) = \left(\frac{\hat{\alpha} - \alpha}{\alpha}\right)^2 \quad (11)$$

- The precautionary (P) Loss Function was introduced by [15] is given by

$$r(\hat{\alpha}) = \frac{(\hat{\alpha} - \alpha)^2}{\hat{\alpha}} \quad (12)$$

- The Entropy (E) Loss Function, see [16]

$$r(\hat{\alpha}) = \left(\frac{\hat{\alpha}}{\alpha} - \log \left(\frac{\hat{\alpha}}{\alpha} \right) - 1 \right), \quad (13)$$

(a) **Estimation under Relative Quadratic Loss Function**

Theorem 1. Assuming the Rq loss function, the Bayesian estimator of α , if λ and θ are known, is of the form

$$\hat{\alpha}_{Rq} = \frac{n-2}{T} \quad (14)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})\theta]^2]^{-1}$.

Proof. The risk function of the estimator α under the Rq loss function in Equation (11) is given by

$$r(\hat{\alpha}) = \int_0^\infty \left(\frac{\hat{\alpha}-\alpha}{\alpha} \right)^2 \pi_1(\alpha|\underline{x}) d\alpha, \quad (15)$$

By substituting Equation (10) in Equation (15), we get

$$r(\hat{\alpha}) = \frac{T^n}{\Gamma(n)} \int_0^\infty \left(\frac{\hat{\alpha}-\alpha}{\alpha} \right)^2 \alpha^{n-1} e^{-\alpha T} d\alpha,$$

by setting $z = \alpha T$

$$r(\hat{\alpha}) = \frac{1}{\Gamma(n)} \left[\hat{\alpha}^2 T^2 \int_0^\infty z^{n-3} e^{-z} dz - 2\hat{\alpha} T \int_0^\infty z^{n-2} e^{-z} dz + \int_0^\infty z^{n-1} e^{-z} dz \right]$$

after solving the integral, we get

$$r(\hat{\alpha}) = \left[\frac{\hat{\alpha}^2 T^2}{(n-1)(n-2)} - \frac{2\hat{\alpha} T}{n-1} + 1 \right] \quad (16)$$

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\frac{\hat{\alpha}^2 T^2}{(n-1)(n-2)} - \frac{2\hat{\alpha} T}{n-1} + 1 \right] = 0,$$

$$\hat{\alpha}_{Rq} = \frac{n-2}{T}.$$

Where T is defined in Equation (9)

Hence, the theorem is proved.

(b) **Estimation under Precautionary Loss Function**

Theorem 2. *Assuming the P loss function, the Bayesian estimator of α , if λ and θ are known, is of the form*

$$\hat{\alpha}_p = \frac{(n(n+1))^{\frac{1}{2}}}{T} \quad (17)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{-\frac{\lambda}{x_i}})^{\theta}]^2]^{-1}$.

Proof. The risk function of the estimator α under the P loss function in Equation (12) is given by

$$r(\hat{\alpha}) = \int_0^{\infty} \frac{(\hat{\alpha} - \alpha)^2}{\hat{\alpha}} \pi_1(\alpha|\underline{x}) d\alpha. \quad (18)$$

By substituting Equation (10) in Equation (18), we get

$$r(\hat{\alpha}) = \frac{T^n}{\Gamma(n)} \int_0^{\infty} \frac{(\hat{\alpha} - \alpha)^2}{\hat{\alpha}} \alpha^{n-1} e^{-\alpha T} d\alpha,$$

$$r(\hat{\alpha}) = \frac{T^n}{\Gamma(n)} \left[\int_0^{\infty} \hat{\alpha} \alpha^{n-1} e^{-\alpha T} d\alpha - 2 \int_0^{\infty} \alpha^n e^{-\alpha T} d\alpha + \frac{1}{\hat{\alpha}} \int_0^{\infty} \alpha^{n+1} e^{-\alpha T} d\alpha \right]$$

by setting $z = \alpha T$

$$r(\hat{\alpha}) = \frac{1}{\Gamma(n)} \left[\hat{\alpha} \int_0^{\infty} z^{n-1} e^{-z} dz - 2 \frac{1}{T} \int_0^{\infty} z^n e^{-z} dz + \frac{1}{\hat{\alpha} T^2} \int_0^{\infty} z^{n+1} e^{-z} dz \right]$$

after solving the integral, we get

$$r(\hat{\alpha}) = \left[\hat{\alpha} - 2 \frac{n}{T} + \frac{n(n+1)}{\hat{\alpha} T^2} \right] \quad (19)$$

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\hat{\alpha} - 2 \frac{n}{T} + \frac{n(n+1)}{\hat{\alpha} T^2} \right] = 0,$$

$$\hat{\alpha}_p = \frac{(n(n+1))^{\frac{1}{2}}}{T}.$$

Where T is defined in Equation (9)

Hence, the theorem is proved.

(c) Estimation under Entropy Loss Function

Theorem 3. *Assuming the E loss function, the Bayesian estimator of α , if*

λ and θ are known, is of the form

$$\hat{\alpha}_E = \frac{(n-1)}{T} \quad (20)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{-1}$.

Proof. The risk function of the estimator α under the E loss function in Equation (13) is given by

$$r(\hat{\alpha}) = \int_0^{\infty} \left(\frac{\hat{\alpha}}{\alpha} - \log\left(\frac{\hat{\alpha}}{\alpha}\right) - 1 \right) \pi_1(\alpha|\underline{x}) d\alpha, \quad (21)$$

By substituting Equation (10) in Equation (21), we get

$$\begin{aligned} r(\hat{\alpha}) &= \frac{T^n}{\Gamma(n)} \int_0^{\infty} \left[\frac{\hat{\alpha}}{\alpha} - \log\left(\frac{\hat{\alpha}}{\alpha}\right) - 1 \right] \alpha^{n-1} e^{-\alpha T} d\alpha \\ &= \frac{T^n}{\Gamma(n)} \left[\hat{\alpha} \int_0^{\infty} \alpha^{n-2} e^{-\alpha T} d\alpha - \log(\hat{\alpha}) \int_0^{\infty} \alpha^{n-1} e^{-\alpha T} d\alpha \right. \\ &\quad \left. + \int_0^{\infty} \log(\alpha) \alpha^{n-1} e^{-\alpha T} d\alpha - \int_0^{\infty} \alpha^{n-1} e^{-\alpha T} d\alpha \right] \end{aligned}$$

by setting $z = \alpha T$

$$\begin{aligned} r(\hat{\alpha}) &= \frac{1}{\Gamma(n)} \left[\hat{\alpha} T \int_0^{\infty} z^{n-2} e^{-z} dz - \log(\hat{\alpha}) \int_0^{\infty} z^{n-1} e^{-z} dz \right. \\ &\quad \left. + \int_0^{\infty} \log\left(\frac{z}{T}\right) z^{n-1} e^{-z} dz - \int_0^{\infty} z^{n-1} e^{-z} dz \right] \\ &= \frac{1}{\Gamma(n)} \left[\hat{\alpha} T \int_0^{\infty} z^{n-2} e^{-z} dz - \log(\hat{\alpha}) \int_0^{\infty} z^{n-1} e^{-z} dz \right. \\ &\quad \left. + \int_0^{\infty} \log(z) z^{n-1} e^{-z} dz - \int_0^{\infty} \log(T) z^{n-1} e^{-z} dz \right. \\ &\quad \left. - \int_0^{\infty} z^{n-1} e^{-z} dz \right], \end{aligned}$$

After solving the integral, we get

$$r(\hat{\alpha}) = \frac{\hat{\alpha} T}{n-1} - \log(\hat{\alpha}) + \psi(n) - \log(T) - 1$$

where $\psi(n) = \frac{\Gamma'(n)}{\Gamma(n)}$ called Psi (or digamma) function.

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\frac{T}{n-1} - \log(\hat{\alpha}) + \psi(n) - \log(T) - 1 \right] = 0,$$

$$\hat{\alpha}_E = \frac{(n-1)}{T}.$$

Where T is defined in Equation (9)

Hence, the theorem is proved.

(ii) Posterior Distribution under Quasi prior

The Quasi prior for parameter α is given by

$$g_2(\alpha) \propto \frac{1}{\alpha^d} \quad \alpha, d > 0. \quad (22)$$

By using the Bayes theorem, we get

$$\pi_2(\alpha|\underline{x}) \propto L(x|\alpha)g_2(\alpha). \quad (23)$$

By substituting Equation (4) and Equation (22) in Equation (23), we have

$$\pi_2(\alpha|\underline{x}) = A_2 \alpha^{n-d} \exp[-\alpha \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{-1}].$$

$$\pi_2(\alpha|\underline{x}) = A_2 \alpha^{n-d} e^{-\alpha T}, \quad (24)$$

where T is defined in Equation (9).

And A_2 is the normalizing constant and

$$A_2^{-1} = \int_0^{\infty} \alpha^{n-d} e^{-\alpha T} d\alpha,$$

$$A_2^{-1} = \frac{\Gamma(n-d+1)}{T^{n-d+1}}.$$

Therefore

$$A_2 = \frac{T^{n-d+1}}{\Gamma(n-d+1)}.$$

Using the value of A_2 in Equations (24), we have

$$\pi_2(\alpha|\underline{x}) = \frac{\alpha^{n-d} T^{n-d+1}}{\Gamma(n-d+1)} e^{-\alpha T}. \quad (25)$$

(i) **Estimation under Relative Quadratic Loss Function**

Theorem 4. Assuming the Rq loss function, the Bayesian estimator of α , if λ and θ is known, are of the form

$$\hat{\alpha}_{Rq} = \frac{n-d-1}{T} \quad (26)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{\frac{-\lambda}{x_i}})^{\theta}]^2]^{-1}$.

Proof. The risk function of the estimator α under the Rq loss function in Equation (11) is given by

$$r(\hat{\alpha}) = \int_0^{\infty} \left(\frac{\hat{\alpha}-\alpha}{\alpha}\right)^2 \pi_2(\alpha|\underline{x}) d\alpha, \quad (27)$$

By substituting Equation (25) in Equation (27), we get

$$\begin{aligned} r(\hat{\alpha}) &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \int_0^{\infty} \left(\frac{\hat{\alpha}-\alpha}{\alpha}\right)^2 \alpha^{n-d} e^{-\alpha T} d\alpha \\ &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \int_0^{\infty} \left[\hat{\alpha}^2 \alpha^{n-d-2} - 2\hat{\alpha}\alpha^{n-d-1} + \alpha^{n-d} \right] e^{-\alpha T} d\alpha \end{aligned}$$

By setting $z = \alpha T$

$$r(\hat{\alpha}) = \frac{T^{n-d}}{\Gamma(n-d+1)} \int_0^{\infty} \left[\hat{\alpha}^2 \frac{z^{n-d-2}}{T^{n-d-2}} - 2\hat{\alpha} \frac{z^{n-d-1}}{T^{n-d-1}} + \frac{z^{n-d}}{T^{n-d}} \right] e^{-z} dz$$

After solving the integral, we get

$$r(\hat{\alpha}) = \left[\frac{\hat{\alpha}^2 T^2}{(n-d)(n-d-1)} - \frac{2\hat{\alpha}T}{n-d} + 1 \right] \quad (28)$$

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\frac{\hat{\alpha}^2 T^2}{(n-d)(n-d-1)} - \frac{2\hat{\alpha}T}{n-d} + 1 \right] = 0,$$

$$\hat{\alpha}_{Rq} = \frac{n-d-1}{T}.$$

Where T is defined in Equation (9).

Hence, the theorem is proved.

(ii) **Estimation under Precautionary Loss Function**

Theorem 5. Assuming the P loss function, the Bayesian estimator of α , if λ and θ are known, is of the form

$$\hat{\alpha}_P = \frac{((n-d+2)(n-d+1))^{\frac{1}{2}}}{T} \quad (29)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{-\frac{\lambda}{x_i}})^{\theta}]^2]^{-1}$.

Proof. The risk function of the estimator α under the P loss function in Equation (12) is given by

$$r(\hat{\alpha}) = \int_0^{\infty} \frac{(\hat{\alpha} - \alpha)^2}{\hat{\alpha}} \pi_2(\alpha|\underline{x}) d\alpha, \quad (30)$$

By substituting Equation (25) in Equation (30), we get

$$\begin{aligned} r(\hat{\alpha}) &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \int_0^{\infty} \frac{(\hat{\alpha} - \alpha)^2}{\hat{\alpha}} \alpha^{n-d} e^{-\alpha T} d\alpha \\ &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \left[\hat{\alpha} \int_0^{\infty} \alpha^{n-d} e^{-\alpha T} d\alpha - 2 \int_0^{\infty} \alpha^{n-d+1} e^{-\alpha T} d\alpha \right. \\ &\quad \left. + \frac{1}{\hat{\alpha}} \int_0^{\infty} \alpha^{n-d+2} e^{-\alpha T} d\alpha \right] \end{aligned}$$

By setting $z = \alpha T$

$$\begin{aligned} r(\hat{\alpha}) &= \frac{1}{\Gamma(n-d+1)} \left[\hat{\alpha} \int_0^{\infty} z^{n-d} e^{-z} dz - 2 \frac{1}{T} \int_0^{\infty} z^{n-d+1} e^{-z} dz \right. \\ &\quad \left. + \frac{1}{\hat{\alpha} T^2} \int_0^{\infty} z^{n-d+2} e^{-z} dz \right] \end{aligned}$$

After solving the integral, we get

$$r(\hat{\alpha}) = \left[\hat{\alpha} - 2 \frac{n-d+1}{T} + \frac{(n-d+2)(n-d+1)}{\hat{\alpha} T^2} \right]$$

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\hat{\alpha} - 2 \frac{n-d+1}{T} + \frac{(n-d+2)(n-d+1)}{\hat{\alpha} T^2} \right] = 0,$$

$$\hat{\alpha}_P = \frac{((n-d+2)(n-d+1))^{\frac{1}{2}}}{T}.$$

Where T is defined in Equation (9).

Hence, the theorem is proved.

(iii) Estimation under Entropy Loss Function

Theorem 6. Assuming the E loss function, the Bayesian estimator of α , if λ and θ are known, is of the form

$$\hat{\alpha}_E = \frac{(n-d)}{T} \quad (31)$$

Where $T = \sum_{i=1}^n \log[1 - [1 - (1 - e^{-\frac{\lambda}{x_i}})^{\theta}]^2]^{-1}$.

Proof. The risk function of the estimator α under the E loss function in Equation (13) is given by

$$r(\hat{\alpha}) = \int_0^{\infty} \left(\frac{\hat{\alpha}}{\alpha} - \log\left(\frac{\hat{\alpha}}{\alpha}\right) - 1 \right) \pi_2(\alpha|\underline{x}) d\alpha, \quad (32)$$

By substituting Equation (25) in Equation (32).

$$\begin{aligned} r(\hat{\alpha}) &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \int_0^{\infty} \left[\frac{\hat{\alpha}}{\alpha} - \log\left(\frac{\hat{\alpha}}{\alpha}\right) - 1 \right] \alpha^{n-d} e^{-\alpha T} d\alpha \\ &= \frac{T^{n-d+1}}{\Gamma(n-d+1)} \left[\hat{\alpha} \int_0^{\infty} \alpha^{n-d-1} e^{-\alpha T} d\alpha - \log(\hat{\alpha}) \int_0^{\infty} \alpha^{n-d} e^{-\alpha T} d\alpha \right. \\ &\quad \left. + \int_0^{\infty} \log(\alpha) \alpha^{n-d} e^{-\alpha T} d\alpha - \int_0^{\infty} \alpha^{n-d} e^{-\alpha T} d\alpha \right] \end{aligned}$$

By setting $z = \alpha T$

$$\begin{aligned} r^*(\hat{\alpha}) &= \frac{1}{\Gamma(n-d+1)} \left[\hat{\alpha} T \int_0^{\infty} z^{n-d-1} e^{-z} dz - \log(\hat{\alpha}) \int_0^{\infty} z^{n-d} e^{-z} dz \right. \\ &\quad \left. + \int_0^{\infty} \log\left(\frac{z}{T}\right) z^{n-d} e^{-z} dz - \int_0^{\infty} z^{n-d} e^{-z} dz \right], \\ &= \frac{1}{\Gamma(n-d+1)} \left[\hat{\alpha} T \int_0^{\infty} z^{n-d-1} e^{-z} dz - \log(\hat{\alpha}) \int_0^{\infty} z^{n-d} e^{-z} dz \right. \\ &\quad \left. + \int_0^{\infty} \log(z) z^{n-d} e^{-z} dz - \int_0^{\infty} \log(T) z^{n-d} e^{-z} dz - \int_0^{\infty} z^{n-d} e^{-z} dz \right], \end{aligned}$$

After solving the integral, we get

$$r(\hat{\alpha}) = \frac{\hat{\alpha} T}{n-d} - \log(\hat{\alpha}) + \psi(n-d+1) - \log(T) - 1,$$

where $\psi(n-d+1) = \frac{\Gamma'(n-d+1)}{\Gamma(n-d+1)}$ called Psi (or digamma) function.

The optimal estimator is obtained by solving $\frac{\partial r(\hat{\alpha}, T)}{\partial \hat{\alpha}} = 0$

$$\frac{\partial}{\partial \hat{\alpha}} \left[\frac{T}{n-d} - \log(\hat{\alpha}) + \psi(n-d+1) - \log(T) - 1 \right] = 0,$$

$$\hat{\alpha}_E = \frac{(n-d)}{T}.$$

Where T is defined in Equation (9).

Hence, the theorem is proved.

3.2. Case 2: standard Bayes Estimation

Suppose α is unknown and λ, θ are known. And α has Gamma prior distribution $\alpha \sim \text{Gamma}(a, b)$; thus, the prior for α is given by:

$$\pi(\alpha) = \frac{b^a}{\Gamma(a)} \alpha^{a-1} e^{-b\alpha}, \quad \alpha > 0, a, b > 0. \quad (33)$$

By combining (4) and (33) the posterior distribution of the unknown parameter α is given by:

$$\pi^*(\alpha|\underline{x}) = k \alpha^{n+a-1} e^{-\alpha[b - \sum_{i=1}^n \log[1 - [1 - [1 - e^{-\frac{\lambda}{x_i}}]^\theta]^2]} \quad (34)$$

where k is the normalizing constant, can be written as

$$k^{-1} = \int_0^\infty \alpha^{n+a-1} e^{-\alpha[b - \sum_{i=1}^n \log[1 - [1 - [1 - e^{-\frac{\lambda}{x_i}}]^\theta]^2]} d\alpha. \quad (35)$$

Now, we will estimate (α) with Gamma prior under three Loss Functions as following:

(i) Squared Error loss function (SE) when α is unknown

The SE loss function of $\varphi(\alpha)$ can be defined as:

$$\hat{\varphi}_{SE}(\alpha) = E(\varphi(\alpha)|\underline{x}) = \int \varphi(\alpha) \pi^*(\alpha|\underline{x}) d\alpha, \quad (36)$$

then the Bayes estimator of $\varphi(\alpha)$ under SE loss function, denoted by $\hat{\varphi}_{SE}(\alpha)$, and

can be found using Equations (34) and (36).

$$\hat{\varphi}_{SE}(\alpha) = k \int_0^\infty \varphi(\alpha) \alpha^{n+a-1} e^{-\alpha[b - \sum_{i=1}^n \log[1 - [1 - [1 - e^{-\frac{\lambda}{x_i}}]^\theta]^2]} d\alpha. \quad (37)$$

Where k^{-1} is defined in equation (35).

(ii) **LINEX loss function (LINEX) when α is unknown**

The LINEX loss function proposed by [17] as follows:

$$\hat{\varphi}_{LINEX}(\alpha) = -\frac{1}{\tau} \ln [E_\alpha(e^{-\tau\varphi(\alpha)} | \underline{x})] = -\frac{1}{\tau} \ln \left[\int e^{-\tau\varphi(\alpha)} \pi^*(\alpha | \underline{x}) d\alpha \right], \quad (38)$$

then the Bayes estimator of $\varphi(\alpha)$ under the LINEX loss function, denoted by $\hat{\varphi}_{LINEX}(\alpha)$, can be found by using Equations (34) and (38).

$$\hat{\varphi}_{LINEX}(\alpha) = -\frac{1}{\tau} \ln \left[k \int_0^\infty e^{-\tau\varphi(\alpha)} \alpha^{n+a-1} e^{-\alpha[b - \sum_{i=1}^n \log[1 - [1 - [1 - e^{-\frac{\lambda}{x_i}}]^\theta]^2]} d\alpha \right], \quad (39)$$

Where k^{-1} is defined in equation (35).

(iii) **General Entropy loss function (GE) when α is unknown**

The GE loss function for $\varphi(\alpha)$ can be defined by [16] as follows:

$$\hat{\varphi}_{GE}(\alpha) = [E_\alpha(\varphi(\alpha)^{-q} | \underline{x})]^{-\frac{1}{q}} = \left[\int \varphi(\alpha)^{-q} \pi^*(\alpha | \underline{x}) d\alpha \right]^{-\frac{1}{q}}, \quad (40)$$

then the Bayes estimator of $\varphi(\alpha)$ under the GE loss function, denoted by $\hat{\varphi}_{GE}(\alpha)$, can be found by using Equations (34) and (40).

$$\hat{\varphi}_{GE}(\alpha) = \left[k \int_0^\infty \varphi(\alpha)^{-q} \alpha^{n+a-1} e^{-\alpha[b - \sum_{i=1}^n \log[1 - [1 - [1 - e^{-\frac{\lambda}{x_i}}]^\theta]^2]} d\alpha \right]^{-\frac{1}{q}}, \quad (41)$$

Where k^{-1} is defined in equation (35).

The Bayes estimators of α under SE, LINEX, GE loss functions are found numerically by the NIntegrate function via Mathematica 11.3

4. Simulation Study of Bayes and ML Estimators

In this section, numerical results of Bayes and ML estimators for parameter (α) of TIITL-GIED have been presented using Monte Carlo simulation. Comparing between

these estimates are studied based on FindRoot and NIntegrate functions in Mathematica (V.11.3) program. Also, the performance of the parameters is determined according to the bias and mean squared error (MSE), that given, respectively, by

$$\text{Bias}(\hat{\beta}) = E(\hat{\beta}) - \beta, \quad (42)$$

and

$$\text{MSE}(\hat{\beta}) = E(\hat{\beta} - \beta)^2. \quad (43)$$

The following steps are used to accomplish this process.

- (i) A sample of size n from TIITL-GIED is generated using Equation (3) for given values of the parameters.
- (ii) Based on ML and Bayes methods and for different initial values of parameters, the ML and Bayesian estimates of parameter α under Rq, P and E loss functions is computed by solving the Equations (14), (17), (20), (26), (29) and (31).
- (iii) Based on ML and standard Bayes and for given values of hyperparameters ($a = 1, b = 1$) and ($a = 1, b = 2$), the ML and Bayesian estimates of parameter α is computed under SE, LINEX and GE loss functions by solving the Equations (37), (39) and (41).
- (iv) Based on ML and Bayes methods and for different initial values of parameters, the ML and Bayesian estimates of parameter α under all loss functions that we include in this study is computed by solving the Equations (14), (17), (20), (26), (29), (31), (37), (39) and (41).
- (v) The proposed values of all MLE and Bayes of the shape parameters are used.
- (vi) The performance of the parameters is determined according to the Bias and MSE using Equations (42) and (43).
- (vii) The above steps are repeated 1000 times.

Table 1: The ML and Bayesian estimates of the unknown parameter α under Jeffery and Quasi priors for the initial values ($\alpha = 1, \theta = 1, \lambda = 1$).

n	Parameters	ML Estimates (Bias) (MSE)	Jeffery's prior			Quasi prior					
			$\hat{\alpha}_{Rq}$	$\hat{\alpha}_P$	$\hat{\alpha}_E$	$\hat{\alpha}_{Rq}$		$\hat{\alpha}_P$		$\hat{\alpha}_E$	
						D=0.3	D=1	D=0.3	D=1	D=0.3	D=1
10	α	1.09649	0.87719	1.15000	0.98684	0.95394	0.87719	1.22684	1.15000	1.06359	0.98684
		0.09649	-0.12281	0.15000	-0.01316	-0.04606	-0.12281	0.022684	0.15000	0.06359	-0.01316
		0.14407	0.10133	0.17073	0.10933	0.10412	0.10133	0.22016	0.17073	0.13084	0.10933
30	α	1.03119	0.96244	1.04823	0.99681	0.98650	0.96244	1.07229	1.04823	1.02087	0.99681
		0.03119	-0.03756	0.04823	-0.00319	-0.01350	-0.03756	0.07230	0.04823	0.02087	-0.00319
		0.03826	0.03390	0.04086	0.03486	0.03431	0.03390	0.04555	0.04086	0.03699	0.03486
50	α	1.02204	0.98116	1.03221	1.00160	0.99547	0.98116	1.04652	1.03221	1.01591	1.00160
		0.02204	-0.01884	0.03221	0.00160	-0.00453	-0.01884	0.04652	0.03221	0.01591	0.00160
		0.02232	0.02048	0.02331	0.02098	0.02074	0.02048	0.02506	0.02331	0.02182	0.02098
70	α	1.01599	0.98696	1.02322	1.00147	0.99712	0.98696	1.03338	1.02322	1.01163	1.00147
		0.01599	-0.01304	0.02322	0.00147	-0.00288	-0.01304	0.03338	0.02322	0.01163	0.00147
		0.01559	0.01464	0.01609	0.01490	0.01477	0.01464	0.01697	0.01609	0.01533	0.01490

Table 2: The ML and Bayesian estimates of the unknown parameter α under Jeffery and Quasi priors for the initial values ($\alpha = 2, \theta = 2, \lambda = 2$).

n	Parameters	ML Estimates (Bias) (MSE)	Jeffery's prior			Quasi prior					
			$\hat{\alpha}_{Rq}$	$\hat{\alpha}_P$	$\hat{\alpha}_E$	$\hat{\alpha}_{Rq}$		$\hat{\alpha}_P$		$\hat{\alpha}_E$	
						D=0.3	D=1	D=0.3	D=1	D=0.3	D=1
10	α	2.21363	1.77091	2.32168	1.99227	1.92586	1.77091	2.47680	2.32168	2.14722	1.99227
		0.21363	-0.22909	0.32168	-0.00773	-0.07414	-0.22909	0.47680	0.32168	0.14722	-0.00773
		0.61962	0.41983	0.73486	0.46498	0.43994	0.41983	0.94590	0.73486	0.56173	0.46498
30	α	2.06463	1.92698	2.09875	1.99580	1.97516	1.92698	2.14693	2.09875	2.04398	1.99580
		0.06463	-0.07302	0.09875	0.00420	-0.02484	-0.07302	0.14693	0.09875	0.04398	0.00420
		0.14499	0.12799	0.15525	0.13160	0.12949	0.12799	0.17385	0.15525	0.13994	0.13160
50	α	2.03165	1.95038	2.05186	1.99101	1.97882	1.95038	2.08031	2.05186	2.01946	1.99101
		0.03165	-0.04962	0.05186	-0.00899	-0.02118	-0.04962	0.08031	0.05186	0.01946	-0.00899
		0.08375	0.07872	0.08709	0.07955	0.07895	0.07872	0.09321	0.08709	0.08214	0.07955
70	α	2.03025	1.97224	2.04470	2.00124	1.99254	1.97224	2.06500	2.04470	2.02155	2.00124
		0.03025	-0.02776	0.04470	0.00124	-0.00746	-0.02776	0.06500	0.04470	0.02155	0.00124
		0.06242	0.05881	0.06438	0.05976	0.05930	0.05881	0.06786	0.06438	0.06145	0.05976

Table 3: The ML and Bayesian estimates of the unknown parameter α via the standard Bayes techniques for the initial values ($\alpha = 1, \theta = 1, \lambda = 1$) and ($a=1, b=1$).

n	Parameters	ML Estimates (Bias) (MSE)	Bayes Estimates						
			SE (Bias) (MSE)	LINEX (Bias) (MSE)			GE (Bias) (MSE)		
				$\tau=0.001$	$\tau=3$	$\tau=5$	q = -1	q = 3	q = -3
10	α	1.10296	1.08190	1.08190	0.98133	0.86996	1.08190	0.88153	1.17752
		0.10296	0.08190	0.08190	-0.01867	-0.13004	0.08190	-0.11847	0.17752
		0.15134	0.10980	0.10980	0.06902	0.05968	0.10980	0.08247	0.15362
30	α	1.03901	1.03660	1.03660	1.00245	0.95656	1.03660	0.96932	1.06967
		0.03901	0.03660	0.03660	0.00245	-0.04344	0.03660	-0.03068	0.06967
		0.03769	0.03477	0.03477	0.02918	0.02608	0.03477	0.03018	0.04046
50	α	1.02008	1.01926	1.01926	0.99907	0.97070	1.01926	0.97917	1.03913
		0.02008	0.01928	0.01928	-0.00093	-0.02930	0.01928	-0.02083	0.03913
		0.02128	0.02036	0.02036	0.01844	0.01730	0.02036	0.01888	0.02230
70	α	1.01138	1.01101	1.01101	0.99669	0.97620	1.01101	0.98247	1.02519
		0.01138	0.01101	0.01101	-0.00331	-0.02380	0.01101	-0.01755	0.02519
		0.01516	0.01471	0.01471	0.01378	0.01325	0.01471	0.01408	0.01563

Table 4: The ML and Bayesian estimates of the unknown parameter α via the standard Bayes techniques for the initial values ($\alpha = 0.7, \theta = 0.5, \lambda = 0.8$) and ($a=1, b=2$).

n	Parameters	ML Estimates (Bias) (MSE)	Bayes Estimates						
			SE (Bias) (MSE)	LINEX (Bias) (MSE)			GE (Bias) (MSE)		
				$\tau=0.001$	$\tau=3$	$\tau=5$	q = -1	q = 3	q = -3
10	α	0.76742	0.72191	0.72191	0.67520	0.61879	0.72191	0.58822	0.78572
		0.06742	0.02191	0.02191	-0.02480	-0.08121	0.02191	-0.11178	0.08572
		0.07785	0.04536	0.04536	0.03452	0.03044	0.04536	0.04229	0.06052
30	α	0.72302	0.71162	0.71162	0.69528	0.67259	0.71162	0.66545	0.73434
		0.02302	0.01162	0.01162	-0.00472	-0.02741	0.01162	-0.03455	0.03434
		0.01982	0.01690	0.01690	0.01524	0.01402	0.01690	0.01585	0.01903
50	α	0.71283	0.70656	0.70656	0.69679	0.68279	0.70656	0.67876	0.72033
		0.01283	0.00656	0.00656	-0.00321	-0.01722	0.00656	-0.02124	0.02033
		0.00984	0.00900	0.00900	0.00847	0.00809	0.00900	0.00872	0.00972
70	α	0.70903	0.70469	0.70469	0.69770	0.68755	0.70469	0.68479	0.71457
		0.00903	0.00469	0.00469	-0.00230	-0.01245	0.00469	-0.01521	0.01457
		0.00705	0.00662	0.00662	0.00635	0.00614	0.00662	0.00647	0.00691

Table 5: The ML and Bayesian estimates of the unknown parameter α for the initial values ($\alpha = 1, \theta = 1, \lambda = 1$) and ($a=1, b=1$).

n	Parameters	ML Estimates (Bias) (MSE)	Jeffery's prior			Quasi prior						SE (Bias) (MSE)	LINEX (Bias) (MSE)			GE (Bias) (MSE)	
			$\hat{\alpha}_{Rq}$	$\hat{\alpha}_P$	$\hat{\alpha}_E$	$\hat{\alpha}_{Rq}$		$\hat{\alpha}_P$		$\hat{\alpha}_E$			$\tau=0.001$	$\tau=3$	$\tau=5$	q=-1	q=-3
						D=0.3	D=1	D=0.3	D=1	D=0.3	D=1						
10	α	1.12613	0.90090	1.18109	1.01352	0.97973	0.90090	1.26001	1.18109	1.09234	1.01352	1.09950	0.99464	0.87971	1.09950	0.89592	1.19673
		0.12613	-0.09910	0.18109	0.01352	-0.02027	-0.09910	0.26001	0.18109	0.09234	0.01352	0.09950	-0.00536	-0.12029	0.09950	-0.10408	0.19673
		0.20076	0.12813	0.23613	0.14991	0.14033	0.12813	0.29902	0.23613	0.18245	0.14991	0.13840	0.08356	0.06557	0.13840	0.09617	0.19096
30	α	1.02906	0.96046	1.04608	0.99476	0.98447	0.96046	1.07009	1.04608	1.01877	0.99476	1.02690	0.99340	0.94830	1.02690	0.96033	1.05975
		0.02907	-0.03954	0.04608	-0.00524	-0.00155	-0.03954	0.07009	0.04608	0.01877	-0.00524	0.02690	-0.00658	-0.05170	0.02690	-0.03970	0.05975
		0.03743	0.03343	0.03993	0.03421	0.03372	0.03343	0.04447	0.03993	0.03621	0.03421	0.03450	0.03450	0.02956	0.02710	0.03450	0.03110
50	α	1.01323	0.97270	1.02332	0.99297	0.98689	0.97270	1.03750	1.02332	1.00715	0.99297	1.01260	0.99263	0.96461	1.01260	0.97273	1.03230
		0.01323	-0.02730	0.02332	-0.00703	-0.01311	-0.02730	0.03750	0.02332	0.00715	-0.00703	0.01260	-0.00737	-0.03539	0.01260	-0.02727	0.03230
		0.02090	0.01984	0.02168	0.01995	0.01983	0.01984	0.02313	0.02168	0.02052	0.01995	0.01999	0.018353	0.01757	0.01999	0.01905	0.02166
70	α	1.01664	0.98760	1.02388	1.00212	0.99776	0.98760	1.03405	1.02388	1.01229	1.00212	1.01620	1.00175	0.98107	1.01620	0.98751	1.03045
		0.01665	-0.01240	0.02388	0.00212	-0.00224	-0.01240	0.03405	0.02388	0.01229	0.00212	0.01620	0.00175	-0.01893	0.01620	-0.01249	0.03045
		0.01457	0.01364	0.01507	0.01389	0.01377	0.01364	0.01594	0.01507	0.01432	0.01389	0.01412	0.01308	0.01239	0.01412	0.01324	0.01518

Table 6: The ML and Bayesian estimates of the unknown parameter α for the initial values ($\alpha = 0.7, \theta = 0.5, \lambda = 0.8$) and ($a=3, b=5$).

n	Parameters	ML Estimates (Bias) (MSE)	Jeffery's prior			Quasi prior						SE (Bias) (MSE)	LINEX (Bias) (MSE)			GE (Bias) (MSE)		
			$\hat{\alpha}_{Rq}$	$\hat{\alpha}_P$	$\hat{\alpha}_E$	$\hat{\alpha}_{Rq}$		$\hat{\alpha}_P$		$\hat{\alpha}_E$			$\tau=0.001$	$\tau=3$	$\tau=5$	q=-1	q=3	
						D=0.3	D=1	D=0.3	D=1	D=0.3	D=1							
10	α	0.78767	0.63013	0.82611	0.70891	0.68527	0.63013	0.88131	0.82611	0.76404	0.70891	0.71768	0.67896	0.63039	0.71768	0.60561	0.77159	
		0.08767	-0.06987	0.12611	0.00890	-0.01473	-0.06987	0.18131	0.12611	0.06404	0.00890	0.01768	-0.02104	-0.06961	0.01768	-0.09439	0.07159	
		0.08542	0.05463	0.10141	0.06304	0.05905	0.05463	0.13018	0.10141	0.07724	0.06304	0.02959	0.02384	0.02231	0.02959	0.02976	0.03897	
		0.72282	0.67463	0.73476	0.69872	0.69149	0.67463	0.75163	0.73476	0.71559	0.69872	0.70721	0.69209	0.67097	0.70721	0.66412	0.72844	
30	α	0.02282	-0.02537	0.03476	-0.00128	-0.00851	-0.02537	0.05163	0.03476	0.01559	-0.00128	0.00721	-0.00791	-0.02903	0.00721	-0.03588	0.02843	
		0.01918	0.01689	0.02049	0.01743	0.01715	0.01689	0.02284	0.02049	0.01853	0.01743	0.01398	0.01398	0.01282	0.01210	0.01398	0.01559	
		0.71537	0.68676	0.72249	0.70106	0.69677	0.68676	0.73251	0.72249	0.71108	0.70106	0.70677	0.69735	0.68384	0.70677	0.68001	0.72002	
		0.01537	-0.01324	0.02249	0.00106	-0.00323	-0.01324	0.03251	0.02249	0.01108	0.00106	0.00677	-0.00265	-0.01616	0.00677	-0.01999	0.02002	
50	α	0.01070	0.00982	0.01118	0.01005	0.00994	0.00982	0.01203	0.01118	0.01046	0.01005	0.00890	0.00890	0.00802	0.00890	0.00859	0.00959	
		0.71021	0.68992	0.71527	0.70007	0.69702	0.68992	0.72237	0.71527	0.70717	0.70007	0.70445	0.69766	0.68779	0.70445	0.68510	0.71406	
		0.01021	-0.01008	0.01527	0.00006	-0.00298	-0.01008	0.02237	0.01527	0.00717	0.00006	0.00445	0.00445	-0.00234	-0.01221	0.00445	-0.01490	0.01406
		0.00698	0.00659	0.00720	0.00668	0.00663	0.00659	0.00761	0.00720	0.00687	0.00668	0.00613	0.00589	0.00571	0.00613	0.00600	0.00648	
70	α																	

Table 7: The ML and Bayesian estimates of the unknown parameter α for the initial values ($\alpha = 0.7, \theta = 0.5, \lambda = 0.8$) and ($a=1, b=1$).

n	Parameters	ML Estimates (Bias) (MSE)	Jeffery's prior			Quasi prior						SE (Bias) (MSE)	LINEX (Bias) (MSE)			GE (Bias) (MSE)		
			$\hat{\alpha}_{Rq}$	$\hat{\alpha}_P$	$\hat{\alpha}_E$	$\hat{\alpha}_{Rq}$		$\hat{\alpha}_P$		$\hat{\alpha}_E$			$\tau=0.001$	$\tau=3$	$\tau=5$	$q=-1$	$q=3$	$q=-3$
						D=0.3	D=1	D=0.3	D=1	D=0.3	D=1							
10	α	0.76932	0.61545	0.80687	0.69239	0.60931	0.61545	0.80078	0.80687	0.74624	0.69239	0.77977	0.77977	0.72516	0.66040	0.77977	0.63536	0.84869
		0.06932	-0.08455	0.10687	-0.00762	-0.03069	-0.08455	0.16078	0.10687	0.04624	-0.00762	0.07977	0.07977	0.02516	-0.03960	0.07977	-0.06464	0.14869
30	α	0.07518	0.05219	0.08883	0.05706	0.05421	0.05219	0.11395	0.08883	0.06835	0.05706	0.06618	0.06618	0.04480	0.03196	0.06618	0.04389	0.09297
		0.72263	0.67445	0.73457	0.69854	0.69131	0.67445	0.75144	0.73457	0.71540	0.69854	0.72856	0.72856	0.71143	0.68770	0.72856	0.68129	0.75182
50	α	0.02263	-0.02555	0.03457	-0.00146	-0.00869	-0.02555	0.05144	0.03457	0.01540	-0.00146	0.02855	0.02855	0.01143	-0.01230	0.02855	-0.01871	0.05182
		0.01902	0.01678	0.02032	0.01730	0.01702	0.01678	0.02266	0.02032	0.01838	0.01730	0.01866	0.01866	0.01630	0.01423	0.01866	0.01595	0.02168
70	α	0.71513	0.68653	0.72225	0.70083	0.69654	0.68653	0.73226	0.72225	0.71084	0.70083	0.71893	0.71893	0.70879	0.69428	0.71893	0.69065	0.73295
		0.01513	-0.01347	0.02225	0.00083	-0.00346	-0.01347	0.03226	0.02225	0.018428	0.00083	0.01893	0.01893	0.00879	-0.00572	0.01893	-0.00935	0.03295
		0.01107	0.01017	0.01155	0.01041	0.01030	0.01017	0.01241	0.01155	0.01083	0.01041	0.01099	0.01099	0.00926	0.01099	0.00990	0.01214	
		0.70655	0.68636	0.71158	0.69646	0.69343	0.68636	0.71865	0.71158	0.70352	0.69646	0.70938	0.70938	0.70229	0.69199	0.70938	0.68935	0.71933
	α	0.00655	-0.01364	0.01158	-0.00354	-0.00657	-0.01364	0.01865	0.01158	0.00352	-0.00354	0.00938	0.00938	0.00229	-0.00800	0.00938	-0.01065	0.01933
		0.00739	0.00712	0.00758	0.00715	0.00712	0.00712	0.00795	0.00758	0.00729	0.00715	0.00734	0.00734	0.00697	0.00663	0.00734	0.00696	0.00783

4.1. Results and Remarks

The main remarks and findings of the two estimation methods from the above tables summed up as follows.

4.1.1. Results of Bayes and ML Estimates of α from Tables (1, 2).

- (i) The estimates of α get closer to the actual values with increasing sample size. Further, the MSEs and bias decrease and sometimes tend to zero.
- (ii) The MSEs of ML and Bayes estimates of α decrease as sample size increase.
- (iii) The Bayes estimate under Rq and E loss functions usually gives better estimation than ML estimate when $n = 10, 30, 50$ and 70 . It has the smallest MSEs.
- (iv) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of Rq loss function under Jeffery's prior and Rq loss function with ($d = 1$) under Quasi prior are very similar.
- (v) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of P loss function under Jeffery's prior and P loss function with ($d = 1$) under Quasi prior are very similar.
- (vi) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of E loss function under Jeffery's prior and E loss function with ($d = 1$) under Quasi prior are very similar.
- (vii) The ML estimate gives better results of MSEs than P loss function under Jeffery's prior and Quasi prior with ($d = 0.3, 1$)
- (viii) In general Rq loss function under Jeffery's prior and Rq loss function with ($d = 1$) under Quasi prior gives better result of MSEs as increase sample size.

4.1.2. Results of standard Bayes and ML Estimates of α from Tables (3, 4).

- (i) The MSEs of ML and Bayes estimates of α decrease as sample size increase.
- (ii) The Bayes estimate usually gives better estimation than ML estimate when $n = 10, 30, 50$ and 70 .
- (iii) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates under SE loss function, LINEX loss function with ($\tau = 0.001$), and GE loss function with ($q = -1$), are very similar.
- (iv) From table 3, the ML estimate gives better results of MSEs than GE loss function with ($q = -3$).
- (v) LINEX loss function gives better results of the MSEs as increase the value of τ . Also, GE loss function performs better as we increase the value of q .
- (vi) In general LINEX loss function with ($\tau = 5$) gives better results of MSE as increase sample size.

4.1.3. Results of Bayes and ML Estimates of α from Tables (5 to 7).

- (i) The MSEs of ML and Bayes estimates of α decrease as sample size increase.
- (ii) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of Rq loss function under Jeffery's prior and Rq loss function with ($d = 1$) under Quasi prior are very similar.
- (iii) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of P loss function under Jeffery's prior and P loss function with ($d = 1$) under Quasi prior are very similar.
- (iv) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates of E loss function under Jeffery's prior and E loss function with ($d = 1$) under Quasi prior are very similar.
- (v) The $\hat{\alpha}$, bias and MSEs values for Bayes estimates under SE loss function, LINEX loss function with ($\tau = 0.001$), and GE loss function with ($q = -1$), are very similar.
- (vi) LINEX loss function gives better results of the MSE as increase the value of τ . Also, GE loss function performs better as increase the value of q .
- (vii) In general LINEX loss function with ($\tau = 5$) gives better results of MSE as we increase sample size.

5. Applications to Real Data

In this section, we applied three real data sets in different fields to demonstrate the utility and flexibility of using the TIITLGIED. The performance of the estimates is evaluated using some information criteria. The best technique is which has the lowest Akaike information criterion (AIC), log-likelihood (LL), Bayesian information criteria (BIC), consistent Akaike information criteria (CAIC), and Hannan-Quinn information criterion (HQIC) [18]. The formulas of these criteria are The three data sets are modeled with different distributions in some previous studies. These distributions are the Exponentiated Generalized Inverted Exponential Distribution (EGIED), Generalized Inverted Exponential Distribution (GIED) [19], Exponential distribution (EXPD), Topp Leone Generalized Standard Inverted Exponential Distribution (TIITLGSIED) [4] and Topp Leone Standard Inverted Exponential Distribution (TIITLGSIED) [4].

5.1. Data set 1

This data is used by [20]. It consists of nonlinear growth curves regarding cold tolerance for 3 plants each from 2 locations (Quebec and Mississippi) and 2 treatments (controlled and chilled) at 7 levels of CO₂ concentration. The variable CO₂ uptake rate is selected for model fitting. The dataset of 18 observations is recorded as:

10.79, 10.80, 10.86, 10.93, 10.99, 10.96, 10.98, 11.03, 11.08, 11.10, 11.19, 11.25, 11.40, 11.61, 11.69, 11.91, 12.07, and 12.32.

Table 8: MLEs of the model parameters and the statistics of the L.L, AIC, BIC, CAIC and HQIC for the data set 1

The model	The parameters estimate			L.L	AIC	BIC	CAIC	HQIC
	α	θ	λ					
TIITLGIED	1371.31	565.29	114.187	-13.785	33.569	36.241	35.284	33.938
GIED	—	53399.4	124.793	-20.185	44.3693	46.1501	45.1693	44.6149
EGIED	1.62187	3.6238	18.929	-49.593	105.187	107.858	106.901	105.555
EXPD	0.08869	—	—	-61.608	125.215	126.105	125.465	125.338
TIITLGSIED	8.88564	0.15370	—	-68.784	141.568	142.368	143.349	141.814
TIITLSIED	0.550591	—	—	-73.958	149.916	150.166	150.806	150.039

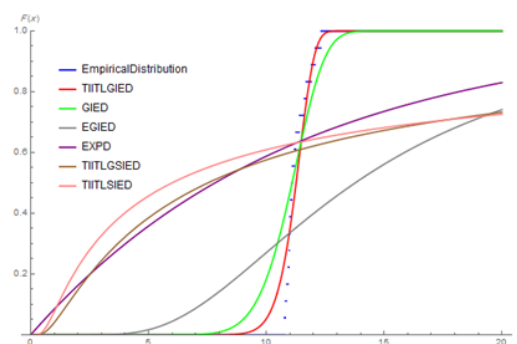
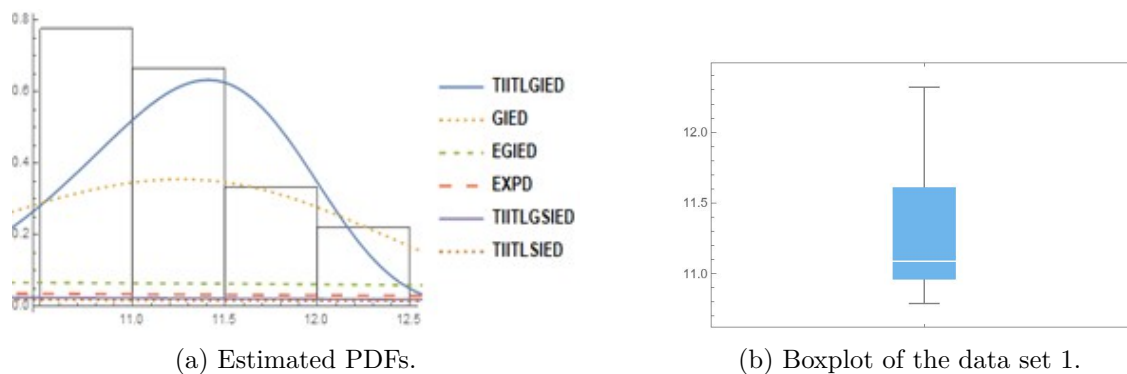


Figure 1: The empirical distribution and estimated CDF of models using data set 1.

Table 9: Descriptive statistics for data set 1

n	Minimum	Maximum	Mean	Median	Standard deviation	Kurtosis	Skewness
18	10.79	12.32	11.2656	11.09	0.4591	2.7667	0.9774



(a) Estimated PDFs.

(b) Boxplot of the data set 1.

Figure 2: Plots of the estimated PDF of the models and the Boxplot using data set 1.

5.2. Data set 2

This data represent current health expenditure in terms of economic growth rates in Saudi Arabia. See [21].

42.1159410, 44.6173573, 42.4937630, 39.7909737, 35.8400607, 34.1867280 36.1922503, 35.6228733, 29.7100425, 42.9041958, 36.4785600, 37.1177721, 40.1962376, 44.6568298, 52.2795486, 59.9834490, 58.3562946, 62.6256323 57.4845695, 56.8828773.

Table 10: MLEs of the model parameters and the statistics of the L.L, AIC, BIC, CAIC and HQIC for the data set 2

The model	The parameters estimate			L.L	AIC	BIC	CAIC	HQIC
	α	θ	λ					
TIITLGIED	4.95268	8.92284	132.01	-73.057	152.114	155.102	153.614	152.698
EGIED	1.19226	38.5462	176.175	-74.122	154.243	157.23	155.743	154.826
EXPD	0.08869	—	—	-55.587	117.174	119.845	118.888	117.542
TIITLGSIED	5.88639	0.12657	—	-111.94	227.893	229.885	228.599	228.282
TIITLSIED	0.32232	—	—	-118.56	239.123	240.119	239.346	239.318

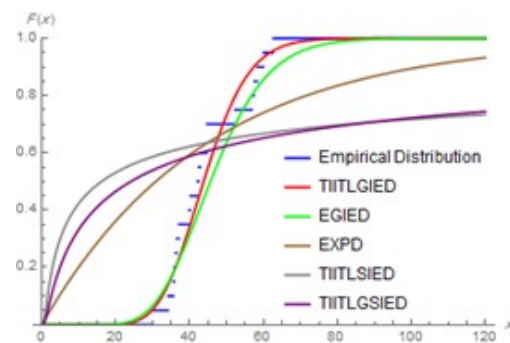


Figure 3: The empirical distribution and estimated CDF of models using data set 2.

Table 11: Descriptive statistics for data set 2

n	Minimum	Maximum	Mean	Median	Standard deviation	Kurtosis	Skewness
20	29.71	62.6256	44.4768	42.3049	9.9008	1.9749	0.5360

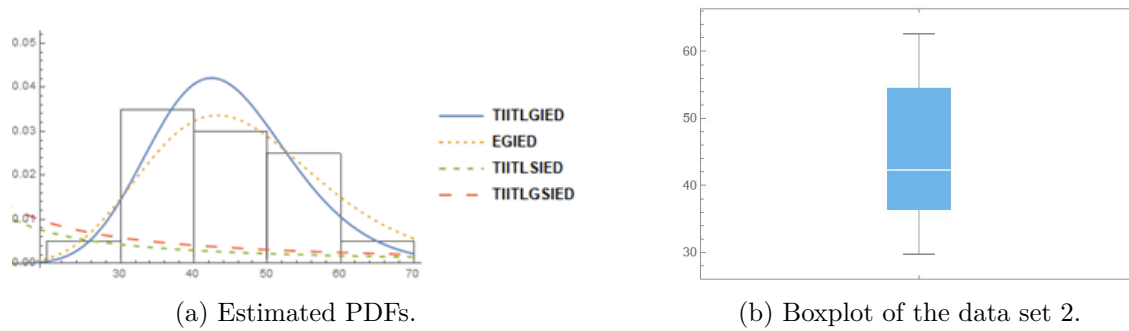


Figure 4: Plots of the estimated PDF of the models and the Boxplot using data set 2.

5.3. Data set 3

The following dataset is the susceptibility index (SI) of peppermint packets irradiated with gamma rays (6, 8, 10 KGy) or microwave rays (1, 2, 3 min) under choice packet test to *Stegobium paniceum* L. See [22]. The SI values were: 2.463, 2.53, 2.34, 2.16, 2.10, 2.13, 1.79, 1.86, 1.75, 1.45, 1.35, 1.37, 2.24, 2.32, 2.29, 2.02, 2.09, 2.15, 1.57, 1.48, and 1.47.

Table 12: MLEs of the model parameters and the statistics of the L.L, AIC, BIC, CAIC and HQIC for the data set 3

The model	The parameters estimate			L.L	AIC	BIC	CAIC	HQIC
	α	θ	λ					
TIITLGIED	20.5142	5.22388	6.36555	-9.0119	24.0239	27.1575	25.4357	24.704
EGIED	1.70988	5.12732	4.33391	-23.048	52.0966	55.2302	53.5084	52.7767
EXPD	0.51316	—	—	-35.011	72.0211	73.0657	72.2317	72.2478
TIITLGSIED	13.3358	0.33297	—	-26.924	57.8489	59.9379	58.5155	58.3023
TIITLSIED	2.27419	—	—	-29.593	61.1858	62.2299	61.3959	61.4121

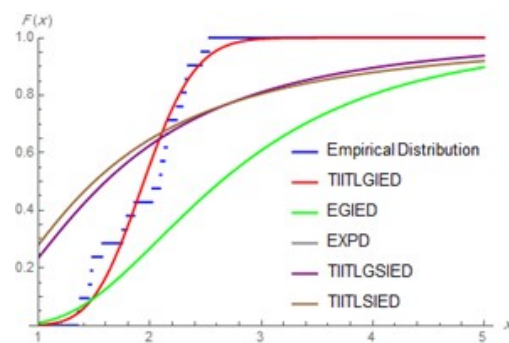
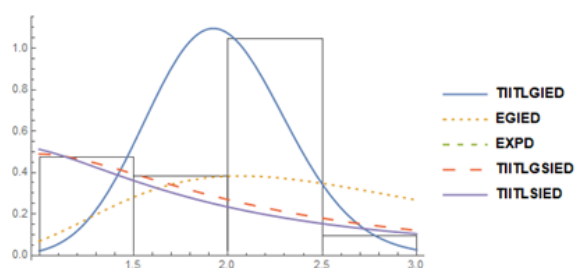


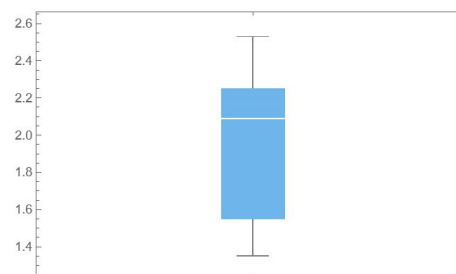
Figure 5: The empirical distribution and estimated CDF of models using data set 3.

Table 13: Descriptive statistics for data set 3

n	Minimum	Maximum	Mean	Median	Standard deviation	Kurtosis	Skewness
21	1.35	2.53	1.9487	2.09	0.3787	1.7344	-0.2584



(a) Estimated PDFs.



(b) Boxplot of the data set 3.

Figure 6: Plots of the estimated PDF of the models and the Boxplot using data set 3.

Tables (8,10,12) show that the parameters of the presented models were estimated along with the values of (LL, AIC, BIC, CAIC and HQIC) using three data sets. Thus, we observed from these tables that the TIITL-GIED gives the smallest values compared to other models values for these data sets. Therefore, the TIITL-GIED is the best fitting for this data. Further, Figures (1,3,5) and (2a, 4a, 6a) Show that the TIITL-GIED is the best fitting for this data comparing with the competing models.

6. Conclusions

In this study, the maximum likelihood and Bayes estimators of the additional shape parameter (α) of Type II Topp Leone Generalized Inverted Exponential Distribution were derived. We compared between these different methods. Hence, we concluded that the MSE and Bias of ML and Bayes estimates of α decreases as sample size increases for all methods and the estimate of α approximate to the initial value. In general LINEX loss function with ($\tau=5$) gives better result of MSE. Three real data sets were applied and the TIITLGED provided a better fit than other compared distributions.

References

- [1] A. M. Abouammoh and Arwa M. Alshingiti. Reliability estimation of generalized inverted exponential distribution. *Journal of Statistical Computation and Simulation*, 79(11):1301–1315, 2009.
- [2] H. Panahi. Exact confidence interval for the generalized inverted exponential distribution with progressively censored data. *Malaysian Journal of Mathematical Sciences*, 11(3):331–345, 2017.

- [3] Gui Wenhao and Guo Lei. Different estimation methods and joint confidence regions for the parameters of a generalized inverted family of distributions. *Hacettepe Journal of Mathematics and Statistics*, 47(1):203–221, 2018.
- [4] Sarah H. Al-Jadaani and Zakeia A. Al-Saiary. Type II Topp Leone Generalized Inverted Exponential Distribution with Real Data Applications. *Asian Journal of Probability and Statistics*, pages 37–51, 2022.
- [5] Sanku Dey and Mazen Nassar. Generalized inverted exponential distribution under constant stress accelerated life test: Different estimation methods with application. *Quality and Reliability Engineering International*, 36(4):1296–1312, 2020.
- [6] Rana A. Bakoban and Maha A. Aldahlan. Bayesian Approximation Techniques for the Generalized Inverted Exponential Distribution. *Intelligent Automation and Soft Computing*, 31(1):129–142, 2022.
- [7] Abdullah M. Almarashi. Inferences of generalized inverted exponential distribution based on partially constant-stress accelerated life testing under progressive Type-II censoring. *Alexandria Engineering Journal*, 63:223–232, 2023.
- [8] Mahmoud R. Mahmoud, Hiba Z. Muhammed, and Ahmed R. El-Saeed. Inference for generalized inverted exponential distribution under progressive Type-I censoring scheme in presence of competing risks model. *Sankhya A*, pages 1–34, 2021.
- [9] Chester W. Topp and Fred C. Leone. A family of J-shaped frequency functions. *Journal of the American Statistical Association*, 50(269):209–219, 1955.
- [10] Saralees Nadarajah and Samuel Kotz. Moments of some J-shaped distributions. *Journal of Applied Statistics*, 30(3):311–317, 2003.
- [11] Ali Al-Shomrani, Osama Arif, A. Shawky, Saman Hanif, and Muhammad Qaiser Shahbaz. Topp–Leone family of distributions: Some properties and application. *Pakistan Journal of Statistics and Operation Research*, pages 443–451, 2016.
- [12] Ahmed R. El-Saeed, Amal S. Hassan, Neema M. Elharoun, Aned Al Mutairi, Rana H. Khashab, and Said G. Nassr. A class of power inverted Topp-Leone distribution: Properties, different estimation methods & applications. *Journal of Radiation Research and Applied Sciences*, 16(4):100643, 2023.
- [13] Mintodê Nicodème Atchadé, Melchior AG N’bouké, Aliou Moussa Djibril, Aned Al Mutairi, Manahil SidAhmed Mustafa, Eslam Hussam, Hassan Alsuhabi, and Said G. Nassr. A new Topp-Leone Kumaraswamy Marshall-Olkin generated family of distributions with applications. *Heliyon*, 10(2), 2024.
- [14] Arnold Zellner. Bayesian and non-Bayesian analysis of the log-normal distribution and log-normal regression. *Journal of the American Statistical Association*, 66(334):327–330, 1971.
- [15] Jan Gerhard Norstrom. The use of precautionary loss functions in risk analysis. *IEEE Transactions on reliability*, 45(3):400–403, 1996.
- [16] R. Calabria and G. Pulcini. An engineering approach to Bayes estimation for the Weibull distribution. *Microelectronics Reliability*, 34(5):789–802, 1994.
- [17] Hal R. Varian. A Bayesian approach to real estate assessment. *Studies in Bayesian econometric and statistics in Honor of Leonard J. Savage*, pages 195–208, 1975.
- [18] Edward J. Hannan and Barry G. Quinn. The determination of the order of an au-

- toregression. *Journal of the Royal Statistical Society: Series B (Methodological)*, 41(2):190–195, 1979.
- [19] P. E. Oguntunde, A. Adejumo, and O. S. Balogun. Statistical properties of the exponentiated generalized inverted exponential distribution. *Applied Mathematics*, 4(2):47–55, 2014.
- [20] A. Gul, M. Mohsin, and Z. Almaspoor. A New Family of Distributions to Encounter the Extrapolating Issues. *Mathematical Problems in Engineering*, page 9397398, 2023.
- [21] M. A. Almuqrin and M. AbaOud. Bayesian and classical inference for a novel model: medical and economical real data analysis. *AIMS Mathematics*, 10(8):17334–17361, 2025.
- [22] M. M. Badr, J. A. Darwish, and F. Alkhathaami. Modeling radiation and engineering data using the exponentiated generalized Topp-Leone exponential model. *Journal of Radiation Research and Applied Sciences*, 18(1):101287, 2025.