

Exploring an extension of the Rayleigh distribution: statistical estimation and practical application

Harsh Tripathi^{†*} and Varun Agiwal[‡]

[†]*Symbiosis Statistical Institute, Symbiosis International (Deemed University), Pune, India*

[‡]*Indian Institute of Public Health, Hyderabad, India*

Email(s): rsearchstat21@gmail.com, varunagiwal.stats@gmail.com

Abstract. In this paper, we introduce a new extension of the Rayleigh distribution, termed the Rayleigh-Exponential distribution (Ra-Ex-D). The proposed distribution is capable of modeling a wide range of data behaviors and demonstrates greater flexibility compared to classical distributions. Several statistical properties of the proposed distribution are derived, including survival characteristics, mean deviation, entropy, and conditional moments, along with the necessary proofs. Model parameters are estimated using both classical and Bayesian approaches under different prior assumptions. A simulation study is conducted to assess the performance of the estimation methods, demonstrating that the proposed procedures yield reliable and consistent estimates. Furthermore, the practical applicability of the model is illustrated through the analysis of a real-life dataset, where the proposed distribution exhibits superior fitting performance compared to existing models.

Keywords: Bayesian estimation; Classical estimation; Moments; Rayleigh distribution.

1 Introduction

In recent decades, probability distributions have been widely used to solve various real-life problems. Many authors have developed new distributions based on transformations and families such as the alpha power transformation, Marshall–Olkin family, logarithmic family, exponentiated family, mixture of distributions, and the G-family, among others. These techniques typically add one or more parameters to a baseline model, thereby increasing the flexibility of the resulting distribution.

In recent years, several researchers have proposed probability distributions using transformation techniques and generalized families to further enhance modeling flexibility. Notable examples

*Corresponding author

Received: 17 May 2025 / Accepted: 5 June 2026

DOI: 10.22067/smeps.2026.93521.1047

include the Marshall–Olkin additive Weibull distribution (Affify et al., 2018), the Marshall–Olkin Pareto distribution (Bdair and Haj Ahmad, 2021), and the Marshall–Olkin Gompertz distribution (Eghwerido et al., 2021), as well as broader developments in Marshall–Olkin-type models (Jose and Paul, 2018; Gonzalez-Hernández et al., 2021). Additionally, transformation-based approaches such as the alpha power transformation (Dey et al., 2019; Mahdavi and Kundu, 2017; Shukla et al., 2022) and transmuted families (Tripathi and Chitra, 2022; Tripathi and Mishra, 2022) have contributed significantly to generating flexible distributions.

Other important contributions include generalized and hybrid families like the TX family Klakattawi et al. (2022), exponentiated classes of distributions (Cordeiro and Lemonte, 2013; Yadav et al., 2021), quasi and extended Xgamma distributions (Sen et al., 2019; Tripathi et al., 2022), and the odd generalized N–H family (Ahmad et al., 2020). These developments are supported by general frameworks for generating continuous distributions (Lee et al., 2013), underscoring the growing interest in flexible modeling techniques.

Despite these advancements, challenges remain in effectively modeling complex lifetime data, particularly in capturing diverse hazard rate behaviors. Recent studies, such as Hassan et al. (2025) and Alsulami and Alghamdi (2025), have proposed extensions of the Weibull distribution to improve flexibility and fitting performance. However, there is still a need for more adaptable models that can simultaneously accommodate a wide range of hazard rate shapes while maintaining tractable statistical properties. This motivates the development of the proposed model. Hence, Elgarhy and Hassan (2019) introduced the Exponentiated Weibull–Weibull (EWW) distribution, which provides a flexible framework for modeling lifetime data. Several extensions of classical lifetime distributions have been developed in the literature to provide greater flexibility in modeling different shapes of probability density and hazard rate functions. Among these models, the Rayleigh–Exponential distribution (Ra-Ex-D) can be considered as a special case of the EWW distribution. We explored the literature and found that no one has made effort regarding the development of Ra-Ex-D.

The main objective of this article is to study the theoretical and inferential properties of the Ra-Ex-D. Specifically, we derive several important statistical properties of the proposed model and investigate different parameter estimation methods under both classical and Bayesian frameworks. In addition, a simulation study is conducted to examine the performance and consistency of the proposed estimators under different parameter combinations. Finally, the applicability of the proposed distribution is illustrated using real-life datasets.

The probability density function (pdf) $f(x)$ and cumulative distribution function (cdf) $F(x)$ of Ra-Ex-D are given by

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

and

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

respectively. Survival function (SF) $S(x)$ and hazard rate function (HRF) $h(x)$ of proposed model are given below

$$S(x) = e^{-\alpha(e^{\lambda x} - 1)^2}, \quad (3)$$

and

$$h(x) = 2\alpha\lambda e^{\lambda x}(e^{\lambda x} - 1), \quad (4)$$

respectively.

Figure 1(a) shows the pdf of the proposed distribution, given in (1), indicating that the proposed model is unimodal. The SF, given in (3), and the cdf, given in (2), are plotted in Figures 1(b) and 1(c), respectively. Figure 1(d) presents the HRF of the proposed distribution, as given in (4), which exhibits both increasing and decreasing hazard rates.

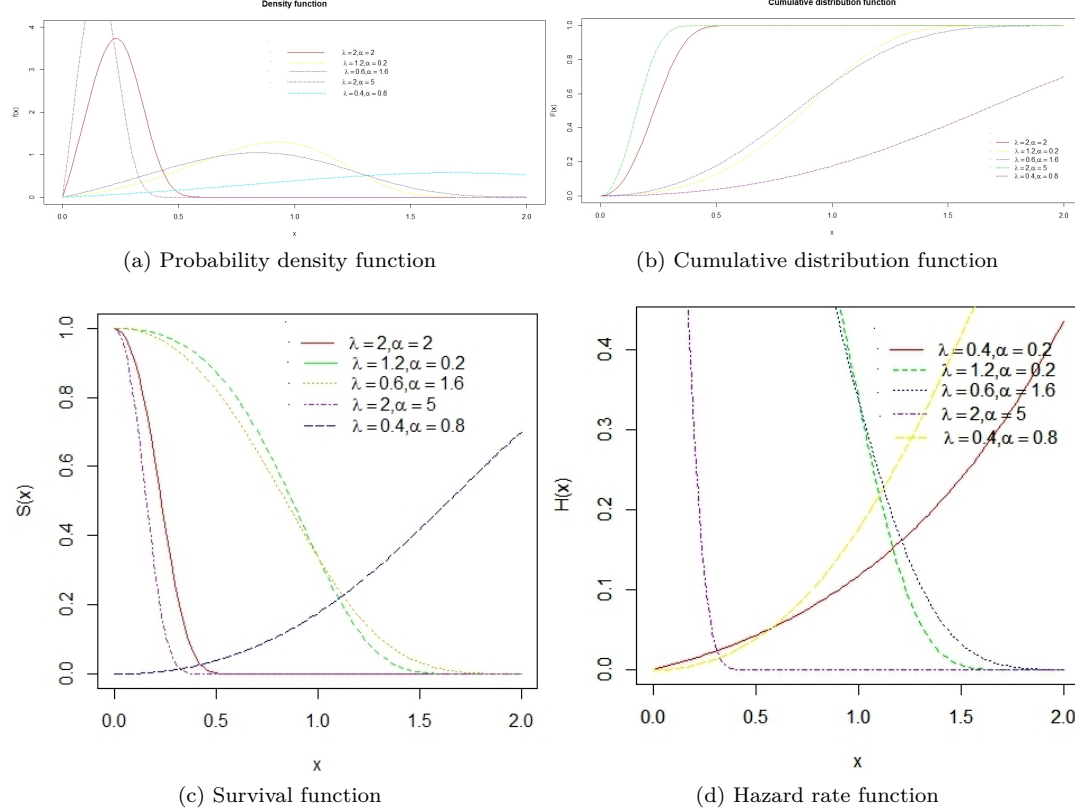


Figure 1: Plots of the various functions of the Ra-Ex-distribution for different parameter values.

The remainder of the paper is organized as follows. Section 2 presents the statistical properties of the proposed distribution. Classical estimation methods are discussed in detail in Section 3. Bayesian estimation is outlined in Section 4. Section 5 presents a simulation study of the considered estimation methods. The application of the proposed probability distribution is discussed in Section 6. Concluding remarks on the Ra-Ex-D are provided in Section 7.

2 Statistical properties

In this section, we derive the statistical properties of the proposed Ra-Ex-D, viz., moments, conditional moments, mean deviation, and entropy.

2.1 Moments

Moments are an important property of a probability distribution, as they help measure skewness and kurtosis. Hence, moments are useful for describing the nature of a dataset. Let X be a random variable following the Ra-Ex-D, with pdf given in (1). Then, the r -th moment about the origin is given by

$$E(X^r) = 2\alpha\lambda \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \frac{\lambda^i}{i!} \frac{(-1)^{j+k}}{j!} \alpha^j \binom{2j}{k} \frac{\Gamma(r+i+1)}{(\lambda k - 2\lambda j)^{r+i+1}} (2^i - 1), \quad r \geq 1.$$

2.2 Conditional moment

Assuming that the lifetime of the units under study follows our proposed model and exceeds a certain threshold x , one may be interested in the conditional moments. Moreover, these are useful for determining the mean deviation, as well as Bonferroni and Lorenz curves. Let X be a random variable following the Ra-Ex-D distribution as in (2); then the expression for the conditional moment is given by:

$$E(X^r | X > x) = \frac{2\alpha\lambda}{1 - F(x)} \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \frac{\lambda^i}{i!} \frac{(-1)^{j+k}}{j!} \alpha^j \binom{2j}{k} \frac{\Gamma(r+i+1, (\lambda k - 2\lambda j)x)}{(\lambda k - 2\lambda j)^{r+i+1}} (2^i - 1).$$

2.3 Mean deviation

Mean deviation (MD) provides the scatterness from mean. Mathematical expression of MD is given below

$$\begin{aligned} \text{MD} &= \int_0^{\infty} |x - \mu| f(x) dx, \\ &= \int_0^{\mu} (\mu - x) f(x) dx + \int_{\mu}^{\infty} (x - \mu) f(x) dx \\ &= 2\mu F(\mu) - 2\mu + 2 \int_{\mu}^{\infty} x f(x) dx, \end{aligned}$$

where

$$\int_{\mu}^{\infty} x f(x) dx = 2\alpha\lambda \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \frac{\lambda^i}{i!} \frac{(-1)^{j+k}}{j!} \alpha^j \binom{2j}{k} \frac{\Gamma(i+2, (\lambda k - 2\lambda j)\mu)}{(\lambda k - 2\lambda j)^{i+2}} (2^i - 1)$$

and μ is the mean.

2.4 Entropy

Entropy measures the uncertainty or information content of a probability distribution and finds applications in various fields, including physics, statistics, reliability analysis, and kernel density estimation. In the context of lifetime distributions, entropy provides insight into the variability and uncertainty associated with the underlying random variable.

The mathematical expression of the generalized entropy (GE) is given by

$$GE = \frac{v_{\eta} \mu^{-\eta} - 1}{\eta(\eta - 1)}; \quad \eta \neq 0, 1,$$

where $v_\eta = \int_0^\infty x^\eta f(x) dx$ represents the η th moment of the distribution and μ denotes the mean. For the proposed Ra-Ex-D distribution, substituting the pdf given in (1) into the above expression, we have

$$v_\eta = 2\alpha\lambda \int_0^\infty x^\eta e^{\lambda x} u e^{-\alpha u^2} dx,$$

where $u = (e^{\lambda x} - 1)$. Substituting this expression for v_η into the generalized entropy formula gives the entropy equation for the proposed distribution.

3 Classical estimation

In this section, we discuss the classical estimation of the parameters. To this end, we employ the following classical methods of estimation: maximum likelihood estimation (MLE), least squares estimation (LSE), weighted least squares estimation (WLSE), Cramér–von Mises estimation (CME), and maximum product spacing estimation (MPSE).

3.1 MLE

Suppose that $X_1, X_2, \dots, X_n$, are independently and identically distributed (i.i.d.) random variables following the proposed distribution with the pdf given in (1). Then, the likelihood function is given by

$$L(\alpha, \lambda | x) = 2^n \alpha^n \lambda^n e^{\lambda \sum_{i=1}^n x_i} \prod_{i=1}^n (e^{\lambda x_i} - 1) e^{-\alpha \sum_{i=1}^n (e^{\lambda x_i} - 1)^2},$$

and the corresponding log-likelihood function can be written as

$$\log L(\alpha, \lambda | x) = n(\log 2 + \log \alpha + \log \lambda) + \lambda \sum_{i=1}^n x_i + \sum_{i=1}^n \log(e^{\lambda x_i} - 1) - \alpha \sum_{i=1}^n (e^{\lambda x_i} - 1)^2. \quad (5)$$

To obtain the maximum likelihood estimates of the parameters, we maximize the log-likelihood function with respect to α and λ . Differentiating equation (5) with respect to α and λ yields the log-likelihood equations:

$$\frac{\partial \log L(\alpha, \lambda | x)}{\partial \alpha} = \frac{1}{\alpha} - \sum_{i=1}^n (e^{\lambda x_i} - 1)^2, \quad (6)$$

$$\frac{\partial \log L(\alpha, \lambda | x)}{\partial \lambda} = \frac{1}{\lambda} + \sum_{i=1}^n x_i + \sum_{i=1}^n \frac{x_i e^{\lambda x_i}}{(e^{\lambda x_i} - 1)} - 2\alpha \sum_{i=1}^n x_i e^{\lambda x_i} (e^{\lambda x_i} - 1). \quad (7)$$

From equation (6), the MLE of α for a given λ is obtained as

$$\hat{\alpha} = \frac{1}{\sum_{i=1}^n (e^{\lambda x_i} - 1)^2}. \quad (8)$$

Substituting (8) into (7), the MLE of λ can be obtained by solving the following non-linear equation:

$$\frac{\partial \log L(\alpha, \lambda | x)}{\partial \lambda} = \frac{1}{\lambda} + \sum_{i=1}^n x_i + \sum_{i=1}^n \frac{x_i e^{\lambda x_i}}{(e^{\lambda x_i} - 1)} - 2 \frac{\sum_{i=1}^n x_i e^{\lambda x_i} (e^{\lambda x_i} - 1)}{\sum_{i=1}^n (e^{\lambda x_i} - 1)^2} = 0. \quad (9)$$

Since equation (9) is not in closed form, numerical techniques such as the Newton–Raphson method are employed to solve this equation and obtain the MLEs of α and λ .

3.2 MPSE

This method of estimation was proposed by [Cheng and Amin \(1983\)](#). Let $X_1, X_2, \dots, X_n$ be a random sample from a distribution with pdf $f(x, \alpha, \lambda)$. Furthermore, let $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ denote the corresponding order statistics. Then the spacings D_i are defined as

$$D_1 = F(x_{(1)}, \alpha, \lambda), \quad D_{n+1} = 1 - F(x_{(n)}, \alpha, \lambda),$$

and

$$D_i = F(x_{(i)}, \alpha, \lambda) - F(x_{(i-1)}, \alpha, \lambda), \quad i = 2, 3, \dots, n,$$

where $\sum_{i=1}^{n+1} D_i = 1$. This method consists of finding the values of α and λ that maximize the geometric mean G of the spacings, i.e.,

$$G = \left\{ \prod_{i=1}^{n+1} D_i \right\}^{\frac{1}{n+1}},$$

or equivalently,

$$\log(G) = \frac{1}{n+1} \sum_{i=1}^{n+1} \log D_i,$$

$$\log(G) = \frac{1}{n+1} \sum_{i=1}^{n+1} \log [F(x_{(i)}, \alpha, \lambda) - F(x_{(i-1)}, \alpha, \lambda)].$$

The basic concept behind this estimation method is that the differences between the cdf at neighboring points should be identically distributed. The MPSEs can be obtained by solving the following non-linear equations:

$$\frac{\partial \log(G)}{\partial \alpha} = \sum_{i=1}^{n+1} \frac{F'_\alpha(x_{(i)}, \alpha, \lambda) - F'_\alpha(x_{(i-1)}, \alpha, \lambda)}{F(x_{(i)}, \alpha, \lambda) - F(x_{(i-1)}, \alpha, \lambda)},$$

$$\frac{\partial \log(G)}{\partial \lambda} = \sum_{i=1}^{n+1} \frac{F'_\lambda(x_{(i)}, \alpha, \lambda) - F'_\lambda(x_{(i-1)}, \alpha, \lambda)}{F(x_{(i)}, \alpha, \lambda) - F(x_{(i-1)}, \alpha, \lambda)},$$

where

$$F'_\alpha(x, \alpha, \lambda) = \left(e^{\lambda x} - 1 \right)^2 e^{-\alpha(e^{\lambda x} - 1)^2}$$

and

$$F'_\lambda(x, \alpha, \lambda) = 2\alpha\lambda e^{\lambda x} \left(e^{\lambda x} - 1 \right) e^{-\alpha(e^{\lambda x} - 1)^2}.$$

3.3 LSE and WLSE

Let $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ be the ordered sample of size n drawn from the pdf given in (1). Then the expectation of the empirical cdf is given by

$$E[F(X_{(i)})] = \frac{i}{n+1}, \quad i = 1, 2, 3, \dots, n.$$

The least squares estimates (LSEs) of α and λ are obtained by minimizing

$$LS(\alpha, \lambda) = \sum_{i=1}^n \left(F(x_{(i)}, \alpha, \lambda) - \frac{i}{n+1} \right)^2 = \sum_{i=1}^n \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{i}{n+1} \right)^2,$$

where we have used the fact that $F(x, \alpha, \lambda) = 1 - \alpha(e^{\lambda x} - 1)^2$ for the proposed distribution. The LSEs $\hat{\alpha}$ and $\hat{\lambda}$ can be found by solving the following equations:

$$\begin{aligned} \frac{\partial LS(\alpha, \lambda)}{\partial \alpha} &= 2 \sum_{i=1}^n F'_\alpha(x, \alpha, \lambda) \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{i}{n+1} \right) = 0, \\ \frac{\partial LS(\alpha, \lambda)}{\partial \lambda} &= 2 \sum_{i=1}^n F'_\lambda(x, \alpha, \lambda) \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{i}{n+1} \right) = 0, \end{aligned}$$

where $F'_\alpha(x, \alpha, \lambda)$ and $F'_\lambda(x, \alpha, \lambda)$ are as defined in the previous subsection. These equations can be solved numerically using methods such as the Newton–Raphson method.

The WLSEs of α and λ are obtained by minimizing, with respect to α and λ , the following function:

$$WLS(\alpha, \lambda) = \sum_{i=1}^n \frac{(n+1)^2(n+2)}{i(n-i+1)} \left(F(x_{(i)}, \alpha, \lambda) - \frac{i}{n+1} \right)^2,$$

which can be expressed as

$$WLS(\alpha, \lambda) = \sum_{i=1}^n \frac{(n+1)^2(n+2)}{i(n-i+1)} \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{i}{n+1} \right)^2.$$

Thus, the WLSEs are determined by solving the equations $\frac{\partial WLS(\alpha, \lambda)}{\partial \alpha} = 0$ and $\frac{\partial WLS(\alpha, \lambda)}{\partial \lambda} = 0$ simultaneously. These equations are analogous to those presented for the LSE, with the weight sequence incorporated appropriately.

3.4 CME

Another well-known estimation method is the Cramér–von Mises estimation (CME), introduced by [Darling \(2019\)](#). In the context of the proposed model, the CMEs of α and λ are obtained by minimizing the following function with respect to α and λ :

$$C(\alpha, \lambda) = \frac{1}{12n} + \sum_{i=1}^n \left(F(x_{(i)}, \alpha, \lambda) - \frac{2i-1}{2n} \right)^2,$$

which can be expressed as

$$C(\alpha, \lambda) = \frac{1}{12n} + \sum_{i=1}^n \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{2i-1}{2n} \right)^2.$$

Thus, the CMEs are obtained by solving the following equations simultaneously: $\frac{\partial C(\alpha, \lambda)}{\partial \alpha} = 0$ and $\frac{\partial C(\alpha, \lambda)}{\partial \lambda} = 0$, where

$$\frac{\partial C(\alpha, \lambda)}{\partial \alpha} = 2 \sum_{i=1}^n F'_\alpha(x, \alpha, \lambda) \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{2i-1}{2n} \right),$$

$$\frac{\partial C(\alpha, \lambda)}{\partial \lambda} = 2 \sum_{i=1}^n F'_\lambda(x, \alpha, \lambda) \left(1 - \alpha \left(e^{\lambda x_{(i)}} - 1 \right)^2 - \frac{2i-1}{2n} \right),$$

and $F'_\alpha(x, \alpha, \lambda)$ and $F'_\lambda(x, \alpha, \lambda)$ are as given above.

This section presented classical parameter estimation methods for the proposed distribution. The MLE involves solving nonlinear equations, which require iterative numerical optimization such as the Newton–Raphson method, leading to moderate computational complexity. The LSE and WLSE methods are computationally simpler but may be less efficient. Although more computationally intensive, Bayesian estimation was also implemented to provide more robust estimates, particularly for small sample sizes or complex likelihood surfaces.

4 Bayesian estimation

In the Bayesian framework, a probability model is fitted not only using the likelihood function of the observations but also incorporating prior beliefs about the unknown parameters in the form of prior distributions. This yields model results based on the combined information from the likelihood and the prior, summarized in the posterior distribution. For the proposed distribution, the likelihood function is

$$L(\alpha, \lambda \mid x) = 2^n \alpha^n \lambda^n e^{\lambda \sum_{i=1}^n x_i} \prod_{i=1}^n \left(e^{\lambda x_i} - 1 \right) e^{-\alpha \sum_{i=1}^n (e^{\lambda x_i} - 1)^2}. \quad (10)$$

From a strict Bayesian perspective, it is not possible to definitively state that one prior is more relevant or better than another. Therefore, priors are broadly classified into two categories: informative and non-informative, each of which may be further subdivided. An informative prior is used when sufficient information about the parameters is available, meaning complete prior distributional information, such as conjugate priors or gamma priors, is at hand. Otherwise, a non-informative prior, which contains little or no prior information about the parameter(s), is considered (e.g., uniform prior or improper prior). Here, we employ the gamma prior distribution for our analysis because it offers greater flexibility, respects the positivity constraint, and exhibits conjugacy-like behavior that facilitates posterior computation. The forms of the priors are

$$\pi(\alpha) = \frac{q^p}{\Gamma(p)} \alpha^{p-1} \exp(-q\alpha); \quad p, q > 0, \quad (11)$$

and

$$\pi(\lambda) = \frac{s^r}{\Gamma(r)} \lambda^{r-1} \exp(-s\lambda); \quad r, s > 0, \quad (12)$$

where (p, q, r, s) are the hyperparameters. Note that if the hyperparameters are set to zero ($p = q = r = s = 0$), the gamma prior reduces to a non-informative prior of the form

$$\pi(\alpha) = \frac{1}{\alpha}, \quad (13)$$

and

$$\pi(\lambda) = \frac{1}{\lambda}. \quad (14)$$

The joint posterior density function of the parameters given the observed random variables X is obtained as

$$\pi(\alpha, \lambda \mid x) \propto L(\alpha, \lambda \mid x) \pi(\alpha) \pi(\lambda).$$

Under the informative prior, the joint posterior density function is formed using equations (10)–(12)

$$\pi(\alpha, \lambda | x) \propto 2^n \alpha^{n+p-1} \lambda^{n+r-1} e^{\lambda(\sum_{i=1}^n x_i - s)} \prod_{i=1}^n (e^{\lambda x_i} - 1) e^{-\alpha(\sum_{i=1}^n (e^{\lambda x_i} - 1)^2 + q)}.$$

Under the non-informative prior, the joint posterior density function is formed using equations (10), (13), and (14)

$$\pi(\alpha, \lambda | x) \propto 2^n \alpha^{n-1} \lambda^{n-1} e^{\lambda \sum_{i=1}^n x_i} \prod_{i=1}^n (e^{\lambda x_i} - 1) e^{-\alpha \sum_{i=1}^n (e^{\lambda x_i} - 1)^2}, \quad (15)$$

In the Bayesian framework, a loss function is used to define the statistical risk (error) associated with parameter estimation. It is a function of the true and estimated values, and the goal is to select the estimator that minimizes the risk for each parameter. Here, we consider the squared error loss function (SELF) as a symmetric loss function to facilitate better understanding. The Bayes estimate under this loss function is the posterior mean, given by

$$L_{\text{SELF}}(\Theta, \hat{\Theta}) = (\hat{\Theta} - \Theta)^2.$$

Bayes estimators of any parametric function, say $\phi(\Theta)$, under SELF are defined as

$$\phi_{\text{self}}^*(\Theta | \text{data}) = E_{\pi}[\phi(\Theta) | \text{data}] \propto \int_{\Theta} \phi(\Theta) \pi(\Theta | \text{data}) d\Theta.$$

Due to the high-dimensional integration involved in the joint posterior distributions, the expression for the Bayes estimator cannot be solved analytically. Therefore, we adopt the most popular Markov Chain Monte Carlo (MCMC) methodology. In this methodology, a sequence of random samples is generated from the posterior distribution of interest. Here, we use hybrid Gibbs sampling methods to obtain the Bayes estimates based on the full conditional posterior distributions. Using the joint posterior density function given in (15), the expressions for the conditional posterior pdfs of α and λ can be written as

$$\pi(\alpha | x, \lambda) \propto \alpha^{n+p-1} e^{-\alpha(\sum_{i=1}^n (e^{\lambda x_i} - 1)^2 + q)},$$

and

$$\pi(\lambda | x, \alpha) \propto \lambda^{n+r-1} e^{\lambda(\sum_{i=1}^n x_i - s)} \prod_{i=1}^n (e^{\lambda x_i} - 1) e^{-\alpha \sum_{i=1}^n (e^{\lambda x_i} - 1)^2}.$$

We can directly generate samples from the posterior distribution of α because its conditional posterior follows a gamma distribution with shape parameter $n + p$ and rate (or scale) parameter $\sum_{i=1}^n (e^{\lambda x_i} - 1)^2 + q$. However, we cannot directly generate samples from the posterior distribution of λ because it does not correspond to any well-known distributional form. Hence, the Metropolis–Hastings (MH) algorithm is used to produce the posterior samples. The steps of the hybrid Gibbs sampling approach are as follows:

Step 1. Start with an initial value x^* such that $h(x^*) > 0$.

Step 2. Set $t = 1$.

Step 3. Propose a move y with conditional probability density given x^* , denoted by $g(x^*, y)$ (the candidate proposal density).

Step 4. Calculate the Hastings ratio:

$$\rho(x^*, y) = \frac{h(y)g(y, x^*)}{h(x^*)g(x^*, y)},$$

and if $g(x^*, y) = g(y, x^*)$ (symmetric proposal), the Hastings ratio reduces to

$$\rho(x^*, y) = \frac{h(y)}{h(x^*)}.$$

Step 5. Accept the proposed move y with probability $\gamma = \min[1, \rho(x^*, y)]$ and reject with probability $1 - \gamma$.

Step 6. Repeat Steps 2–4 for $t = 10000$ iterations and record the sequence of parameter observations. Finally, obtain the Bayes estimates under SELF.

5 Simulation study

In this section, we present a numerical study to compare the behavior of the different estimation methods described above. Since closed-form solutions are not available, numerical optimization techniques are employed for classical estimators, while Bayesian estimates are obtained through posterior simulation. For this purpose, different combinations of parameter values and various sample sizes are considered to draw meaningful inferences. Random samples from the proposed distribution are generated using the inverse CDF method, which involves the following steps:

1. Set n , α , and λ .
2. Generate $U_l \sim U(0, 1)$ for $l = 1, 2, \dots, n$.
3. Compute $x = (x_1, x_2, \dots, x_n)$ using U such that

$$x_l = F^{-1}(U_l) = \frac{1}{\lambda} \log \left[1 + \left(-\frac{\log(1 - U_l)}{\alpha} \right)^{\frac{1}{2}} \right].$$

4. Compute the estimates using the generated samples.

We consider sample sizes $n = (20, 50, 100)$ and four sets of parameter values:

$$(\alpha, \lambda) = \{(0.2, 0.5), (0.5, 1.5), (1, 1), (1.5, 1.2)\}$$

to evaluate the performance of the estimators. These values are selected to represent increasing, decreasing, and bathtub-shaped hazard functions. Bayesian estimation is implemented using hybrid Gibbs sampling and the Metropolis–Hastings algorithm with 10000 iterations, including a burn-in period of 2000 iterations. All computations are performed using R-4.5.3. The hyperparameters of the gamma prior are chosen such that the prior means equal the true parameter values and the prior variances are equal to 1; these estimates (Bayes informative prior) are denoted as “Bayes IP”. The non-informative prior estimates are denoted as “Bayes NIP”. We obtain the parameter estimates under classical and Bayesian methods, reporting the average estimates (AV) and mean

Table 1: AV and MSE of the proposed model for classical and Bayesian estimates with $\alpha = 0.2$ and $\lambda = 0.5$.

P	n		MLE	MPSE	LSE	WLSE	CME	Bayes IP	Bayes NIP
$\alpha = 0.2$	20	AV	0.3669	0.3705	0.3830	0.3300	0.3515	0.2832	0.1809
		MSE	0.4402	0.4650	0.4717	0.3867	0.4136	0.0211	0.0477
	50	AV	0.2589	0.3429	0.3401	0.2744	0.3052	0.2501	0.1887
		MSE	0.2514	0.3311	0.3111	0.1016	0.1779	0.0053	0.0074
	100	AV	0.2271	0.3062	0.2768	0.2365	0.2454	0.1882	0.2126
		MSE	0.0640	0.0780	0.1101	0.0595	0.0680	0.0002	0.0008
$\lambda = 0.5$	20	AV	0.5874	0.4308	0.4671	0.4804	0.4908	0.5067	0.5519
		MSE	0.1273	0.0551	0.0529	0.0528	0.0557	0.0023	0.0059
	50	AV	0.5335	0.4583	0.4779	0.4881	0.4916	0.4692	0.5175
		MSE	0.0696	0.0451	0.0476	0.0467	0.0526	0.0016	0.0039
	100	AV	0.5148	0.4731	0.4863	0.4947	0.4975	0.4903	0.4915
		MSE	0.0558	0.0435	0.0456	0.0465	0.0489	0.0001	0.0002

Table 2: AV and MSE of the proposed model for classical and Bayesian estimates with $\alpha = 0.5$ and $\lambda = 1.5$.

P	n		MLE	MPSE	LSE	WLSE	CME	Bayes IP	Bayes NIP
$\alpha = 0.5$	20	AV	0.6023	0.6707	0.6858	0.5669	0.4580	0.5785	0.5820
		MSE	0.6237	0.6679	0.6780	0.6055	0.6307	0.0248	0.0396
	50	AV	0.5430	0.5929	0.5877	0.5168	0.4659	0.4486	0.4719
		MSE	0.5690	0.5970	0.5837	0.5688	0.5908	0.0095	0.0172
	100	AV	0.5290	0.4823	0.5463	0.5045	0.4905	0.5381	0.4951
		MSE	0.4758	0.5384	0.5682	0.4370	0.5075	0.0033	0.0050
$\lambda = 1.5$	20	AV	1.6401	1.4433	1.4351	1.4552	1.5199	1.4860	1.3840
		MSE	1.0252	0.7076	0.7598	0.7314	0.8281	0.0252	0.0280
	50	AV	1.6034	1.4673	1.4330	1.4596	1.5141	1.5497	1.4046
		MSE	0.8416	0.5783	0.6615	0.6030	0.7072	0.0157	0.0234
	100	AV	1.5576	1.4925	1.4544	1.4948	1.4955	1.4990	1.4618
		MSE	0.6565	0.5045	0.5766	0.5742	0.6017	0.0034	0.0092

Table 3: AV and MSE of the proposed model for classical and Bayesian estimates with $\alpha = 1$ and $\lambda = 1$.

P	n		MLE	MPSE	LSE	WLSE	CME	Bayes IP	Bayes NIP
$\alpha = 1$	20	AV	1.4134	1.3021	0.9073	1.1235	0.9521	0.9242	1.1296
		MSE	0.6144	0.5981	0.6596	0.5626	0.6286	0.1357	0.1827
	50	AV	1.1459	1.1278	1.1960	1.1129	1.1716	1.0442	1.0863
		MSE	0.4195	0.4005	0.5543	0.3531	0.4251	0.0648	0.0139
	100	AV	1.1099	0.9210	1.0563	0.9722	0.9104	1.0012	1.0158
		MSE	0.2589	0.2318	0.4303	0.2086	0.3459	0.0131	0.0519
$\lambda = 1$	20	AV	1.2231	0.9109	1.0425	1.0371	1.0997	0.9318	1.1682
		MSE	0.4349	0.2228	0.1621	0.1362	0.2174	0.0210	0.0377
	50	AV	1.0743	0.8708	1.0177	1.0257	1.0650	1.0371	1.0601
		MSE	0.1533	0.1403	0.1292	0.0944	0.1366	0.0056	0.0064
	100	AV	1.0377	0.9121	1.0033	1.0162	1.0409	0.9591	0.9922
		MSE	0.0760	0.0792	0.0732	0.0575	0.0800	0.0028	0.0046

Table 4: AV and MSE of the proposed model for classical and Bayesian estimates with $\alpha = 1.5$ and $\lambda = 1.2$.

P	n		MLE	MPSE	LSE	WLSE	CME	Bayes IP	Bayes NIP
$\alpha = 1.5$	20	AV	1.6979	1.1828	1.1693	1.2095	1.0294	1.3572	1.3708
		MSE	0.6917	0.5852	0.6105	0.5712	0.6401	0.2525	0.3985
	50	AV	1.6211	1.3461	1.3180	1.3397	1.2232	1.4129	1.3758
		MSE	0.5796	0.5253	0.5520	0.5267	0.6079	0.1418	0.1728
	100	AV	1.5478	1.4653	1.4521	1.4617	1.4466	1.4822	1.4412
		MSE	0.4866	0.4398	4.8382	0.4577	0.4768	0.0202	0.0367
$\lambda = 1.2$	20	AV	1.5947	1.1616	1.3208	1.3114	1.4281	1.1335	1.0945
		MSE	0.4278	0.4465	0.3450	0.2964	0.4155	0.0187	0.0297
	50	AV	1.3246	1.0645	1.2457	1.2431	1.3250	1.1673	1.3005
		MSE	0.2834	0.3005	0.2324	0.2039	0.2548	0.0165	0.0272
	100	AV	1.2613	1.0861	1.2071	1.2336	1.2651	1.2241	1.2181
		MSE	0.1544	0.2076	0.1884	0.1437	0.1783	0.0121	0.0201

squared errors (MSE). For the different parameter combinations with varying n , Tables 1–4 record the AV and MSE.

From the tables, we observe that the MSEs for α based on WLSE are smaller than those of the other classical estimators for all sample sizes. MPSE performs better for λ when $\alpha < 1$, while for $\alpha \geq 1$, WLSE performs better compared to other classical estimators in terms of MSE. In comparison with Bayesian methods, informative prior estimators yield smaller MSEs than other estimators for both parameters. We also see that the MSEs of classical estimators are larger than those of Bayesian estimators for all sample sizes and parameter combinations. Overall, the results indicate that Bayesian estimators consistently exhibit lower MSE compared to classical methods, particularly for small sample sizes. Among classical methods, MLE performs better than least squares-based estimators. It can also be observed that as the sample size increases, the difference in MSE between the two approaches decreases.

6 Real life example

In this section, we demonstrate the application of the proposed model to a real-world dataset. For this purpose, we consider a dataset representing the tensile strength (measured in GPa) of 69 carbon fibers tested under tension at a gauge length of 20 mm (Bader and Priest, 1982):

1.312, 1.314, 1.479, 1.552, 1.700, 1.803, 1.861, 1.865, 1.944, 1.958, 1.966, 1.997, 2.006, 2.021, 2.027, 2.055, 2.063, 2.098, 2.140, 2.179, 2.224, 2.240, 2.253, 2.270, 2.272, 2.274, 2.301, 2.301, 2.359, 2.382, 2.382, 2.426, 2.434, 2.435, 2.478, 2.490, 2.511, 2.514, 2.535, 2.554, 2.566, 2.570, 2.586, 2.629, 2.633, 2.642, 2.648, 2.684, 2.697, 2.726, 2.770, 2.773, 2.800, 2.809, 2.818, 2.821, 2.848, 2.880, 2.954, 3.012, 3.067, 3.084, 3.090, 3.096, 3.128, 3.233, 3.433, 3.585, 3.858

A descriptive summary of the considered dataset is presented in Table 5. First, we assess whether the dataset fits the Ra-Ex-D distribution. For this purpose, we use Akaike’s Information Criterion (AIC), Bayesian Information Criterion (BIC), the Kolmogorov–Smirnov (K-S) statistic, and the corresponding p -value. A model with the lowest K-S value and the largest p -value indicates the best fit. For the proposed model, the values of the log-likelihood (L-L), AIC, BIC, K-S statistic, and p -value are reported in Table 6. From Table 6, we can readily verify that the Ra-Ex-D distribution fits the dataset better than the Lindley distribution (LD), Akash distribution (AKD), exponential distribution (ED), and generalized exponential distribution (GED). Figure 2 displays the histogram with fitted density, the empirical CDF plot, and the probability–probability (P-P) plot for the considered dataset. From Figure 2, we conclude that the Ra-Ex-D distribution provides a satisfactory fit. In certain scenarios, the Ra-Ex-D distribution may be preferred over other distributions for problem-solving. The estimated parameter values along with their standard deviations (SD) under different estimators are presented in Table 7.

Table 5: Descriptive summary for the considered dataset.

Minimum	Maximum	Q_1	Q_2	Q_3	Skewness	Kurtosis	Mean
1.312	3.858	2.098	2.478	2.773	0.102	3.2253	2.455

where Q_1 , Q_2 and Q_3 are the first quartile, median and third quartile, respectively.

Table 6: Model fitting summary for the considered dataset.

Model	MLE	L-L	AIC	BIC	K-S	p -value
Ra-Ex-D	$\alpha = 0.0136, \lambda = 0.8407$	-55.4231	114.8463	119.3145	0.0905	0.6247
LD	$\lambda = 0.6535$	-119.3110	240.6220	242.8561	0.4005	<0.001
AKD	$\lambda = 0.9637$	-112.2662	226.5324	228.7665	0.3615	<0.001
ED	$\lambda = 0.4073$	-130.9789	263.9578	266.1919	0.4477	<0.001

Table 7: Classical and Bayes estimates for the parameters of the proposed model based on the real dataset.

Estimators		MLE	MPSE	LSE	WLSE	CME	Bayes NIP
α	Estimate	0.0136	0.0146	0.0158	0.0165	0.0151	0.0211
	SD	0.0118	0.0114	0.0109	0.0111	0.0106	0.0290
λ	Estimate	0.8407	0.9832	1.0801	1.0909	1.1126	0.7212
	SD	0.1306	0.1362	0.1598	0.1514	0.1451	0.1558

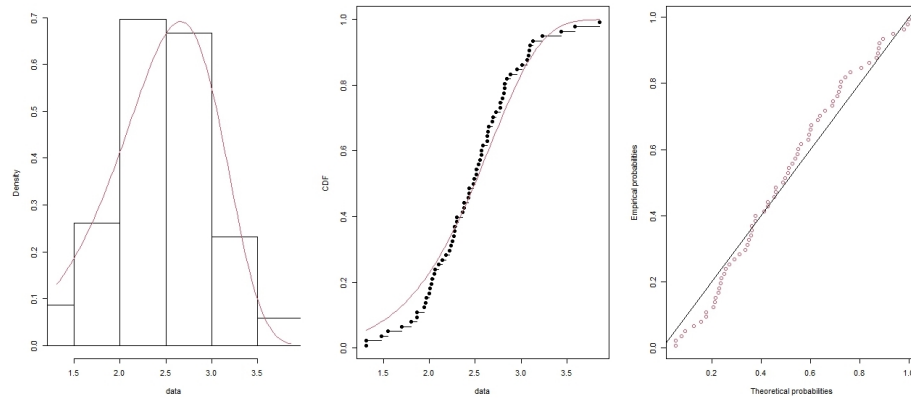


Figure 2: Histogram with fitted density, empirical CDF, and P-P plot.

7 Conclusions

In this paper, we developed a new probability distribution as an extended version of the Rayleigh distribution, termed the Ra-Ex-D. The main theoretical contribution of this work lies in deriving explicit mathematical expressions for its probability density function, cumulative distribution function, survival function, hazard rate function, and several important statistical characteristics, including moments and related measures. The shapes of the probability density function and hazard rate function of the proposed distribution are discussed through graphical illustrations.

The inferential framework of the model was developed under both classical and Bayesian paradigms, allowing for a comprehensive comparison of estimation procedures. The simulation results demonstrated that the estimators perform satisfactorily, with the mean squared error decreasing as the sample size increases, thereby confirming the consistency and reliability of the estimation methods. Overall, the Bayesian estimates are found to be better compared to the classical estimates. Finally, a real-life data example is provided to illustrate the applicability of the proposed model.

Acknowledgements

Availability of data and material: The link of the data set used in this study is included within the paper and also cited in the paper accordingly.

References

- Affify, A. Z., Cordeiro, G. M., Yousof, H. M., Saboor, A., and Ortega, E. M. (2018). The Marshall-Olkin additive Weibull distribution with variable shapes for the hazard rate. *Haceteppe Journal of Mathematics and Statistics*, **47**, 365–381.
- Ahmad, Z., Elgarhy, M., Hamedani, G. G., and Butt, N. S. (2020). Odd generalized NH generated family of distributions with application to exponential model. *Pakistan Journal of Statistics and Operation Research*, 53–71.
- Alsulami, D., and Alghamdi, A. S. (2025). A new extended Weibull distribution: estimation methods and applications in engineering, physics, and medicine. *Mathematics*, **13**, 3262.
- Bdair, O., and Haj Ahmad, H. (2021). Estimation of the Marshall-Olkin Pareto distribution parameters: comparative study. *Revista Investigacion Operacional*, **42**.
- Bader, M. G., and Priest, A. M. (1982). Statistical aspects of fibre and bundle strength in hybrid composites. *Progress in Science and Engineering of Composites*, 1129–1136.
- Cheng, R. C. H., and Amin, N. A. K. (1983). Estimating parameters in continuous univariate distributions with a shifted origin. *Journal of the Royal Statistical Society: Series B (Methodological)*, **45**, 394–403.
- Cordeiro, G. M., and Lemonte, A. J. (2013). On the Marshall–Olkin extended weibull distribution. *Statistical Papers*, **54**, 333–353.
- Dey, S., Ghosh, I., and Kumar, D. (2019). Alpha-power transformed Lindley distribution: properties and associated inference with application to earthquake data. *Annals of Data Science*, **6**, 623–650.

- Darling, D. A. (1957). The kolmogorov-Smirnov, Cramer-von Mises tests. *The Annals of Mathematical Statistics*, **28**, 823–838.
- Eghwerido, J. T., Ogbo, J. O., and Omotoye, A. E. (2021). The Marshall-Olkin Gompertz distribution: properties and applications. *Statistica*, **81**, 183–215.
- Elgarhy, M., and Hassan, A. (2019). Exponentiated-Weibull distribution: statistical properties and applications. *Gazi University Journal of Science*, **32**, 616–635.
- Jose, K., and Paul, A. L. B. I. N. (2018). Marshall-Olkin exponential power distribution and its generalization: theory and applications. *IAPQR Transactions*, **43**.
- González-Hernández, I. J., Granillo-Macías, R., Rondero-Guerrero, C., and Simón-Marmolejo, I. (2021). Marshall-Olkin distributions: a bibliometric study. *Scientometrics*, **126**, 9005–9029.
- Hassan, A. S., Morgan, Y. S., and Abd-Allah, M. (2025). Parameter estimation and data analysis using a novel extension of Weibull distribution. *The Egyptian Statistical Journal*, **69**, 131–154.
- Klakattawi, H., Alsulami, D., Elaal, M. A., Dey, S., and Baharith, L. (2022). A new generalized family of distributions based on combining Marshal-Olkin transformation with TX family. *PloS one*, **17**, e0263673.
- Lee, C., Famoye, F., and Alzaatreh, A. Y. (2013). Methods for generating families of univariate continuous distributions in the recent decades. *Wiley Interdisciplinary Reviews: Computational Statistics*, **5**, 219–238.
- Mahdavi, A., and Kundu, D. (2017). A new method for generating distributions with an application to exponential distribution. *Communications in Statistics-Theory and Methods*, **46**, 6543–6557.
- Shukla, S., Yadav, A. S., Rao, G. S., Saha, M., and Tripathi, H. (2022). Alpha power transformed Xgamma distribution and applications to reliability, survival and environmental data. *Journal of Scientific Research*, **66**, 330–346.
- Sen, S., Afify, A. Z., Al-Mofleh, H., and Ahsanullah, M. (2019). The quasi xgamma-geometric distribution with application in medicine. *Filomat*, **33**, 5291–5330.
- Tripathi, H., Yadav, A. S., Saha, M., and Shukla, S. (2022). A new flexible extension of Xgamma distribution and its application to COVID-19 Data. *Nepal Journal of Mathematical Sciences*, **3**, 11–30.
- Tripathi, H., and Chitra. (2022). Transmuted extended xgamma-2 distribution and its statistical properties. *International Journal of Mechanical Engineering*, **7**.
- Tripathi, H., and Mishra, S. (2022). The transmuted inverse xgamma distribution and its statistical properties. *International Journal of Mechanical Engineering*, **7**.
- Yadav, A. S., Saha, M., Tripathi, H., and Kumar, S. (2021). The exponentiated XGamma distribution: a new monotone failure rate model and its applications to lifetime data. *Statistica*, **81**, 303–334.