
Research article

A generalized Lindley distribution: Properties, estimation and applications under progressively type-II censored samples

Haiping Ren^{1,*} and Jiajie Shi²

¹ Teaching Department of Basic Subjects, Jiangxi University of Science and Technology, Nanchang, 330013, China

² College of Science, Jiangxi University of Science and Technology, Ganzhou, 341000, China

* **Correspondence:** Email: chinarhp@163.com; Tel: +8615979014223.

Abstract: A two-parameter generalized Lindley distribution is proposed, whose probability density function contains two models that are particularly well-suited for lifetime studies: the inverted J-type and the unimodal left-leaning type. Some of the main numerical characteristics of the proposed distribution were investigated. Additionally, based on progressively type-II censored samples, three estimation methods were used to estimate the parameters, reliability function, and hazard function of the distribution: maximum product spacing estimation, maximum likelihood estimation, and Bayesian estimation. Bayesian estimators were obtained under squared error and general entropy loss functions, and the Metropolis-Hastings algorithm was used to obtain Bayesian estimates. In addition to point estimation, corresponding asymptotic confidence intervals and the highest posterior density intervals were also studied. Through Monte Carlo simulation, we measured the performance of the three estimators using four criteria. Finally, the ability of the proposed distribution to accurately fit the data was demonstrated using a real dataset.

Keywords: progressively type-II censored samples; maximum product spacing estimation; Bayesian estimation; Metropolis-Hasting algorithm; asymptotic confidence interval; highest posterior density interval

Mathematics Subject Classification: 62E10; 62F10

1. Introduction

Lifetime models have found widespread application in statistical modeling across various scientific and engineering domains. Lindley distribution, as one of the classical distributions, was first proposed by Lindley [1]. Lindley distribution is highly flexible and has a wide range of applications, such as in the field of medicine, astrophysics, and reliability engineering. However, it has strong limitations in processing complex data, such as skewed and multi-peak data. Based on this, many researchers improved the original Lindley distribution by adding parameters. With the help of the power exponentiated family of distributions, Rajitha and Akhlnath [2] added two parameters to the original Lindley distribution. They called it PEL distribution and notes its higher flexibility compared to the original model. Fatehi and Chhaya [3] extended the Lindley distribution into the extended odd Weibull-Lindley family. Ashour and Eltehiwy [4] introduced a three-parameter exponentiated power Lindley (EPL) distribution by extending the two-parameter power Lindley distribution. Alizadeh et al. [5] defined a four-parameter exponentiated power Lindley power series distribution on the EPL distribution and found that the newly proposed model provided a better fit than the original Lindley distribution to real datasets.

In most real-life testing and reliability experiments, it is seldom possible to wait until all test samples fail; in other words, it is difficult for investigators to observe the lifetime of all items under test, so experimental data obtained often contain censored data. The development and replacement of censored samples have been the focus of many researchers [6–8]. Type-I and type-II censored schemes, as classical methods in right examination, can only move the unit point at the end of the experiment, which lacks some flexibility [9]. The progressively type-II censored scheme is popular for its flexibility, whose various properties and applications have been extensively studied. Balakrishnan et al. [10] discussed the maximum likelihood estimation and the corresponding interval estimation of extreme value distributions under progressively type-II censored samples. Based on progressively type-II censored samples, Seo et al. [11] studied the hierarchical Bayesian estimation of the unknown parameters of a lifetime distribution with a bathtub-shaped failure rate function. Alshenawy et al. [12] used the maximum likelihood estimation method and the maximum product spacing method to estimate the parameters of the extended odd Weibull exponential distribution under progressively type-II censored samples. Their study further delved into the construction of both asymptotic and bootstrap confidence intervals for the said parameters.

This paper aims to introduce a variant of the Lindley distribution, referred to as the extension of the generalized Lindley (NGL) distribution. This extension is developed under progressively type-II censored samples with the objective of broadening the applicability of the traditional Lindley distribution. The NGL distribution is highly flexible, featuring many variants of the Lindley distribution and exponential function. Another purpose of this paper is to evaluate the estimator performance of the NGL distribution in preparation for the subsequent processing of real datasets.

The remainder of this paper is organized as follows: In Section 2, we introduce the NGL distribution and its basic properties. The numerical characteristics of the proposed distribution are investigated in Section 3. We study three estimators of this distribution in Section 4, obtain the corresponding point and interval estimates, and perform Monte Carlo simulations in Section 5. In Section 6, we examine the practical application of the proposed distribution using a real dataset. In

Section 7, the findings of this paper are summarized, and future research priorities are indicated.

2. Generalized Lindley distribution

In this section, we will introduce the NGL distribution, whose probability density function (PDF) and cumulative distribution function (CDF) are

$$f(x; \lambda, \theta) = \lambda e^{-\lambda x} (1 - \theta + \lambda \theta x), x > 0, \lambda > 0, 0 < \theta < 1, \quad (1)$$

$$F(x; \lambda, \theta) = 1 - (1 + \lambda \theta x) e^{-\lambda x}, x > 0, \lambda > 0, 0 < \theta < 1. \quad (2)$$

Here, λ and θ are the scale and shape parameters of the $NGL(\lambda, \theta)$, respectively.

- If $\theta = 1/(\lambda + 1)$, then the CDF is the Lindley distribution, i.e.,

$$F(x; \lambda) = 1 - \frac{\lambda + 1 + \lambda}{\lambda + 1} e^{-\lambda x}, x > 0, \lambda > 0. \quad (3)$$

- If $\theta = \beta/(\lambda + \beta)$, then the CDF is a two-parameter Lindley distribution [13], i.e.,

$$F(x; \lambda, \beta) = 1 - (1 + \frac{\beta \lambda}{\beta + \lambda} x) e^{-\lambda x}, x > 0, \lambda > 0, \beta > 0. \quad (4)$$

- If $\theta \rightarrow 0^+$, then the CDF is exponential distribution, i.e.,

$$F(x; \lambda) = 1 - e^{-\lambda x}, x > 0, \lambda > 0. \quad (5)$$

- If $\theta = 1/(\eta + 1)$, then the CDF is a two-parameter Lindley distribution [14], i.e.,

$$F(x; \lambda, \eta) = 1 - (1 + \frac{\lambda}{1 + \eta} x) e^{-\lambda x}, x > 0, \lambda > 0, \eta > -1. \quad (6)$$

The survival function (SF) and hazard rate function (HRF) of the NGL distribution are:

$$S(x; \lambda, \theta) = 1 - F(x; \lambda, \theta) = (1 + \lambda \theta x) e^{-\lambda x}, \quad (7)$$

$$h(x; \lambda, \theta) = \frac{f(x; \lambda, \theta)}{1 - F(x; \lambda, \theta)} = \lambda - \frac{\lambda \theta}{1 + \lambda \theta x}. \quad (8)$$

There are two main models for the NGL distribution: the inverted J-type and the unimodal and left-leaning (see Figure 1). This type of model is very suitable for the study of product life and species abundance distribution. θ , the shape parameter, has a large influence on the probability density distribution: when θ increases, the PDF image changes and shows a constant tendency to shift to the right. λ also has a certain impact on the NGL distribution; with an increase in λ , the peak value of PDF increases, and the image presents a left-leaning trend.

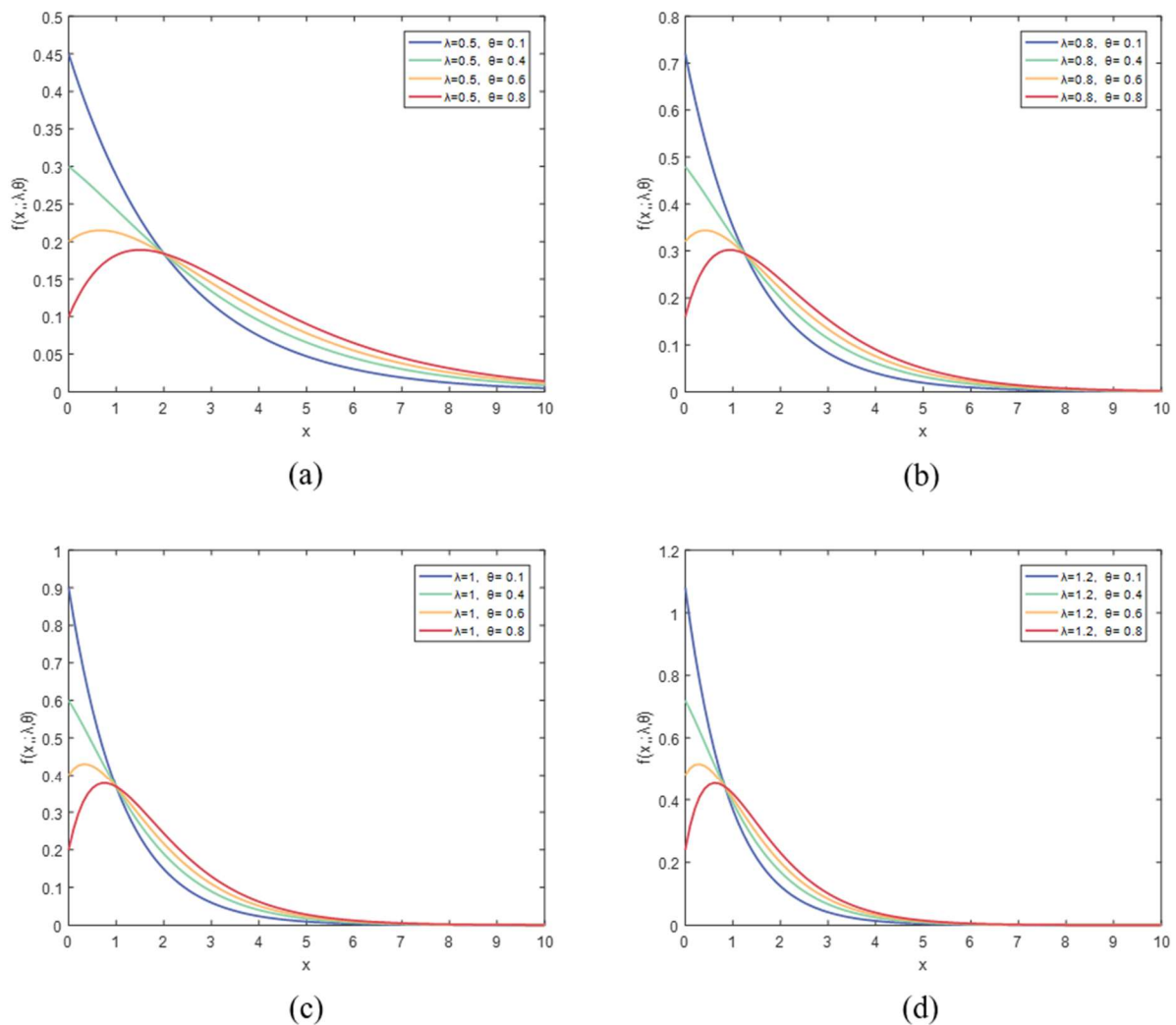


Figure 1. PDF curves of the NGL distribution for different shape and scale parameters.

The HRF curves also have three shapes: increasing, constant, and upside-down bathtub (see Figure 2). With an increase in λ and θ , the slope of (0,2) will increase continuously along with the decrease of peak value.

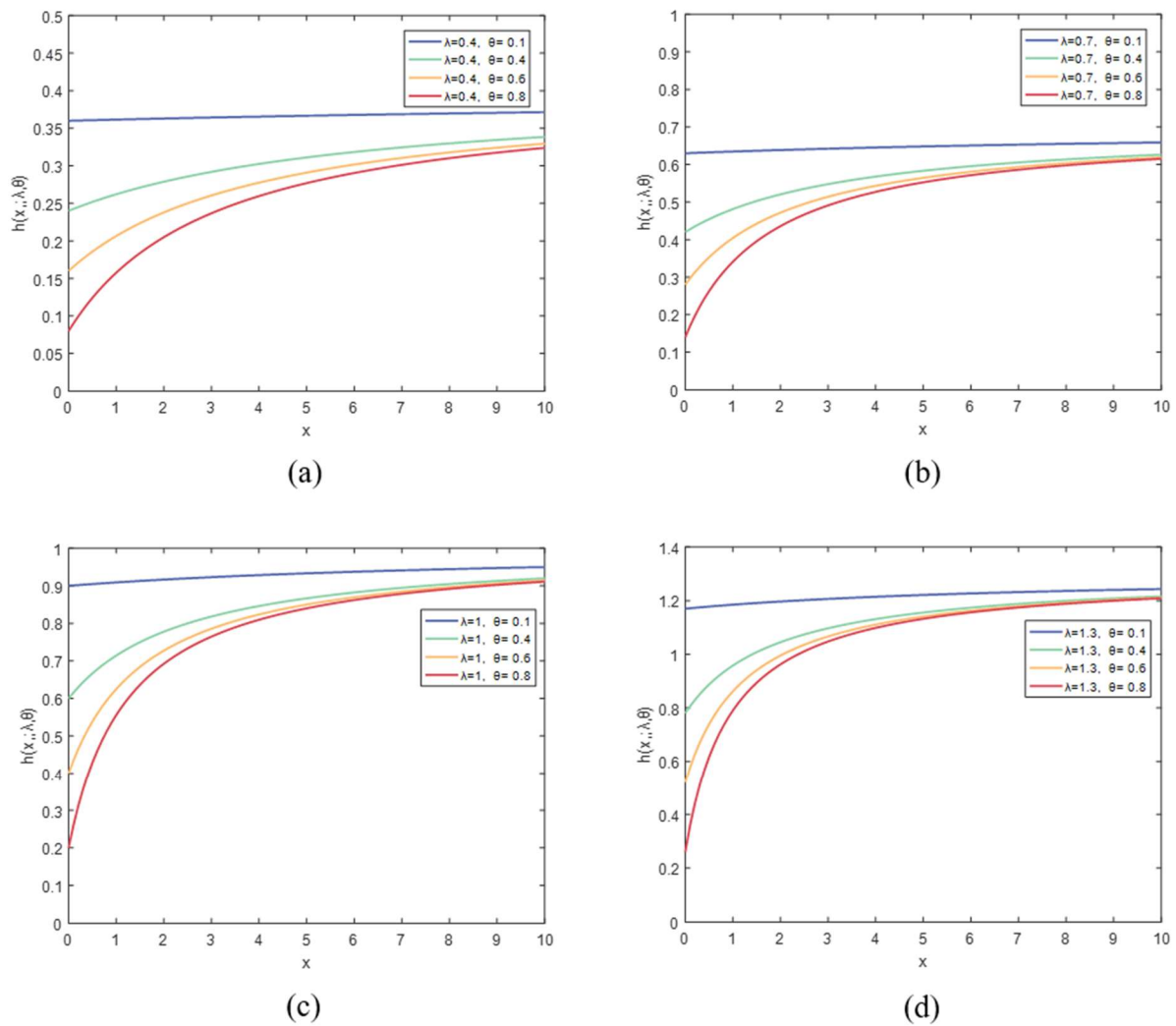


Figure 2. HRF curves of the NGL distribution for different shape and scale parameters.

3. Statistical properties of NGL distribution

In this section, we discuss some important statistical properties of NGL distribution, such as moments, kurtosis and skewness, quantile functions, order statistical functions, and other statistical properties.

3.1. Moments and related measures for the NGL distribution

As one of the most important digital features, the r -th moment plays an important role in both application and theory. It can be represented by the following formula

$$\mu_r = \int_0^\infty x^r e^{-\lambda x} (1 - \theta + \lambda \theta x) dx = (1 - \theta) \lambda^{-r-1} \Gamma(r+1) + \theta \lambda^{-r-1} \Gamma(r+2). \quad (9)$$

The first four moments can be easily calculated:

$$\mu_1 = (1 + \theta)\lambda^{-2}, \mu_2 = (2 + 4\theta)\lambda^{-3}, \mu_3 = (6 + 18\theta)\lambda^{-4}, \mu_4 = (24 + 96\theta)\lambda^{-5}. \quad (10)$$

Also, the variance of X is

$$\text{Var}(X) = (2\lambda + 4\theta\lambda - 1 - 2\theta - \theta^2)\lambda^{-4}. \quad (11)$$

3.2. Coefficient of skewness and kurtosis for the NGL distribution

The coefficient of kurtosis and skewness are important for describing the tail shape, peak degree, and asymmetry of probability distributions [15]. Let NCS stand for coefficient of skewness and NCK for coefficient of kurtosis. According to Eqs (10) and (11), NCS and NCK can be obtained as:

$$\text{NCS} = \frac{\mu_3 - 3\mu_1\mu_2 + 2(\mu_1)^3}{[V(X)]^{3/2}} = 2\lambda \cdot \frac{4\lambda + 11\theta\lambda + \theta^2\lambda - 3 - 9\theta - 6\theta^2}{(2\lambda + 4\theta\lambda - 1 - 2\theta - \theta^2)^{3/2}}, \quad (12)$$

$$\text{NCK} = \frac{\mu_4 - 4\mu_1\mu_3 + 6(\mu_1)^2\mu_2 - 3(\mu_1)^4}{[V(X)]^2} = \frac{24\lambda^3(1+4\theta) - 24\lambda^2(1+3\theta)(\theta+1) + 12\lambda(1+\theta)^2(1+2\theta) - 3(1+\theta)^4}{(2\lambda + 4\theta\lambda - 1 - 2\theta - \theta^2)^2}. \quad (13)$$

The coefficient of kurtosis and skewness of the NGL distribution show a certain regularity. In the NGL distribution, the coefficient of kurtosis and skewness are usually negatively correlated with θ and positively correlated with λ . This statistical property suggests that the thickness and asymmetry of the NGL distribution tail tend to decrease as θ increases, while an increase in λ increases these characteristics. Table 1 and Figure 3 illustrate different values of the coefficient of kurtosis and skewness.

Table 1. The coefficient of kurtosis and skewness of NGL distribution under given parameter values.

θ	λ	NCS	NCK
0.1	0.9	2.0001	8.8471
0.3		1.8457	7.8450
0.5		1.7213	7.1333
0.8		1.6875	6.7500
0.1	1	1.9582	8.6941
0.3		1.7860	7.6311
0.5		1.6198	6.7959
0.8		1.4519	6.1212
0.1	1.2	2.0342	9.0202
0.3		1.8506	7.8409
0.5		1.6577	6.8325
0.8		1.4134	5.8320
0.1	1.5	2.2550	10.1405
0.3		2.0647	8.7803
0.5		1.8590	7.5600
0.8		1.5860	6.2452

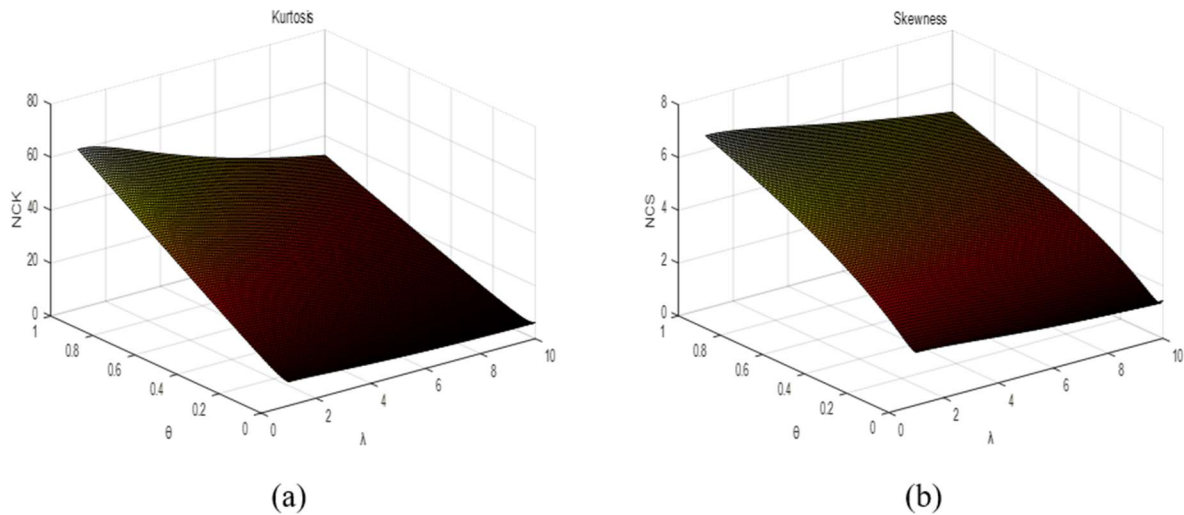


Figure 3. Plots of the coefficient of kurtosis and skewness of the NGL distribution.

3.3. Moment-generating function for the NGL distribution

Using Eq (9), the moment-generating function of the NGL distribution can be obtained as

$$\begin{aligned}
 M(t) &= \int_0^{\infty} e^{tx} f(x) dx = \sum_{r=0}^{\infty} \frac{t^r}{r!} \int_0^{\infty} x^r f(x) dx \\
 &= \sum_{r=0}^{\infty} \frac{t^r}{r!} [(1-\theta)\lambda^{-r-1}\Gamma(r+1) + \theta\lambda^{-r-1}\Gamma(r+2)].
 \end{aligned} \tag{14}$$

3.4. Quantile function for the NGL distribution

The quantile function can be obtained from the inverse function of CDF. That is, $x = F^{-1}(R)$. By applying Eq (2), we have

$$1 - (1 + \lambda\theta x)e^{-\lambda x} = R \tag{15}$$

Let $t = (1 - R)e^{\lambda x}$, then $x = \frac{1}{\lambda} \ln \frac{t}{1-R}$. Hence

$$\ln \frac{t}{1-R} = \frac{1}{\theta} (t - 1) \Rightarrow \frac{t}{1-R} = e^{\frac{t}{\theta}} / e^{\frac{1}{\theta}} \Rightarrow -\frac{t}{\theta} e^{-\frac{t}{\theta}} = \frac{R-1}{\theta e^{\frac{1}{\theta}}} \Rightarrow -\frac{t}{\theta} = W\left(\frac{R-1}{\theta e^{\frac{1}{\theta}}}\right).$$

Then

$$t = -\theta W\left(\frac{R-1}{\theta e^{\frac{1}{\theta}}}\right).$$

Here, W is Lambert W function. According to the basic properties of the Lambert W function, the

quantile function of the NGL distribution can be obtained as

$$x = \frac{-\theta W\left(\frac{e^{-1/\theta}(t-1)}{\theta}\right) - 1}{\lambda\theta}. \quad (16)$$

3.5. Order statistics of the NGL distribution

Order statistics are an important analytical tool for identifying outliers. The earliest failure time of a product can be expressed by the minimum order statistic (the smallest observed value in the sample), and the longest life of a product can be estimated by the maximum order statistic (the largest observed value in the sample). Let $(X_1, X_2, \dots, X_n)$ be a random sample from the NGL distribution, which is sorted from smallest to largest as $(X_{(1)}, X_{(2)}, \dots, X_{(n)})$, where the i -th order statistic is $X_{(i)}$. The PDF of the i -th order statistic $X_{(i)}$ is $f_{X_{(i)}}(x) = \frac{n!}{(i-1)!(n-i)!} F(x)^{i-1} [1 - F(x)]^{n-i} f(x)$. By inserting Eqs (1) and (2), the PDF of i -th order statistic $X_{(i)}$ of the NGL distribution is

$$f_{X_{(i)}}(x) = \frac{n!}{(i-1)!(n-i)!} [1 - (1 + \lambda\theta x)e^{-\lambda x}]^{i-1} [(1 + \lambda\theta x)e^{-\lambda x}]^{n-i} \lambda e^{-\lambda x} (1 - \theta + \lambda\theta x), x > 0, \lambda > 0, 0 < \theta < 1. \quad (17)$$

The PDF of the minimum order statistics $X_{(1)}$ of the NGL distribution can be obtained as

$$f_{X_{(1)}}(x) = n\lambda(1 + \lambda\theta x)^{n-1} e^{-n\lambda x} (1 - \theta + \lambda\theta x). \quad (18)$$

The PDF of the maximum order statistics $X_{(n)}$ of the NGL distribution can be obtained as

$$f_{X_{(n)}}(x) = n\lambda[1 - (1 + \lambda\theta x)e^{-\lambda x}]^{n-1} (1 - \theta + \lambda\theta x)e^{-\lambda x}. \quad (19)$$

4. Parameter estimation of the NGL distribution

4.1. Maximum product spacing estimation

Maximum product spacing (MPS) estimation is a robust method. It uses an optimization algorithm to find the corresponding parameter values that maximize the product spacing of the parametric functions [16]. Let $(X_{1:m:n}, X_{2:m:n}, \dots, X_{m:m:n})$ be the progressively type-II censored sample of the NGL distribution, and $x_{1:m:n}, x_{2:m:n}, \dots, x_{m:m:n}$ is the observation of the sample $(X_{1:m:n}, X_{2:m:n}, \dots, X_{m:m:n})$. For convenience, in the following discussion, we always set $x_i = x_{i:m:n}$. The product spacing function under the progressively type-II censored scheme is [16]:

$$G(\lambda, \theta) = \prod_{i=1}^{m+1} [F(x_i; \lambda, \theta) - F(x_{i-1}; \lambda, \theta)] \prod_{i=1}^m S(x_i; \lambda, \theta)^{P_i}. \quad (20)$$

Using Eq (2), the product spacing function of the NGL distribution can be obtained as

$$G(\lambda, \theta|x) = \prod_{i=1}^{m+1} [(1 + \lambda\theta x_{i-1})e^{-\lambda x_{i-1}} - (1 + \lambda\theta x_i)e^{-\lambda x_i}] \prod_{i=1}^m [(1 + \lambda\theta x_i)e^{-\lambda x_i}]^{P_i} \quad (21)$$

and the log-product spacing function is given by

$$\ln G(\lambda, \theta|x) = \sum_{i=1}^{m+1} \ln[(1 + \lambda\theta x_{i-1})e^{-\lambda x_{i-1}} - (1 + \lambda\theta x_i)e^{-\lambda x_i}] + \sum_{i=1}^m P_i[\ln(1 + \lambda\theta x_i)\lambda x_i]. \quad (22)$$

$$\text{Let } g(\lambda, \theta|x) = \ln G(\lambda, \theta|x) = H(\lambda, \theta|x) + M(\lambda, \theta|x),$$

where

$$H(\lambda, \theta|x) = \sum_{i=1}^{m+1} \ln[(1 + \lambda\theta x_{i-1})e^{-\lambda x_{i-1}} - (1 + \lambda\theta x_i)e^{-\lambda x_i}] \quad (23)$$

$$M(\lambda, \theta|x) = \sum_{i=1}^m P_i[\ln(1 + \lambda\theta x_i) - \lambda x_i] \quad (24)$$

MPS estimators of λ and θ , denoted by $\hat{\lambda}_{MPS}$ and $\hat{\theta}_{MPS}$, are the solutions of Eqs (25) and (26).

$$\frac{\partial g(\lambda, \theta|x)}{\partial \lambda} = \frac{\partial H(\lambda, \theta|x)}{\partial \lambda} + \frac{\partial M(\lambda, \theta|x)}{\partial \lambda} = 0, \quad (25)$$

$$\frac{\partial g(\lambda, \theta|x)}{\partial \theta} = \frac{\partial H(\lambda, \theta|x)}{\partial \theta} + \frac{\partial M(\lambda, \theta|x)}{\partial \theta} = 0. \quad (26)$$

Here,

$$\frac{\partial H(\lambda, \theta|x)}{\partial \theta} = \sum_{i=1}^{m+1} \left[\frac{\theta x_{i-1}}{1 + \lambda\theta x_{i-1}} e^{-\lambda x_{i-1}} - \lambda x_i e^{-\lambda x_i} \right], \quad (27)$$

$$\frac{\partial H(\lambda, \theta|x)}{\partial \lambda} = \sum_{i=1}^{m+1} \left\{ \left[\frac{\theta x_{i-1}}{1 + \lambda\theta x_{i-1}} - x_{i-1} \ln(1 + \lambda\theta x_{i-1}) \right] e^{-\lambda x_{i-1}} + [x_i + \lambda\theta x_i^2 - \theta x_i] e^{-\lambda x_i} \right\}, \quad (28)$$

$$\frac{\partial M(\lambda, \theta|x)}{\partial \theta} = \sum_{i=1}^m \frac{P_i \lambda x_i}{1 + \lambda\theta x_i}, \quad (29)$$

$$\frac{\partial M(\lambda, \theta|x)}{\partial \lambda} = \sum_{i=1}^m P_i \left(\frac{\theta x_i}{1 + \lambda\theta x_i} - x_i \right). \quad (30)$$

$\hat{\lambda}_{MPS}$ and $\hat{\theta}_{MPS}$ cannot be obtained directly from the above equations, so we used the Newton Raphson algorithm to obtain the approximate maximum product spacing estimates of these parameters. The steps are as follows:

- (i) Establish the corresponding log-product spacing function.
- (ii) Compute the gradient vector and Hessian matrix for the log-product spacing function:

$$\nabla g(\lambda, \theta|x) = \left(\frac{\partial g(\lambda, \theta|x)}{\partial \lambda}, \frac{\partial g(\lambda, \theta|x)}{\partial \theta} \right), H(\lambda, \theta|x) = \begin{pmatrix} \frac{\partial^2 g(\lambda, \theta|x)}{\partial \lambda^2} & \frac{\partial^2 g(\lambda, \theta|x)}{\partial \lambda \partial \theta} \\ \frac{\partial^2 g(\lambda, \theta|x)}{\partial \theta \partial \lambda} & \frac{\partial^2 g(\lambda, \theta|x)}{\partial \theta^2} \end{pmatrix}.$$

- (iii) Select the appropriate initial value $(\lambda^{(0)}, \theta^{(0)})^T$.

- (iv) The parameters are updated by Newton iteration formula:

$$\begin{bmatrix} \lambda^{(k+1)} \\ \theta^{(k+1)} \end{bmatrix} = \begin{bmatrix} \lambda^{(k)} \\ \theta^{(k)} \end{bmatrix} - \begin{pmatrix} \frac{\partial^2 g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \lambda^2} & \frac{\partial^2 g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \lambda \partial \theta} \\ \frac{\partial^2 g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \theta \partial \lambda} & \frac{\partial^2 g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \theta^2} \end{pmatrix} \begin{bmatrix} \frac{\partial g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \lambda} \\ \frac{\partial g(\lambda^{(k)}, \theta^{(k)}|x)}{\partial \theta} \end{bmatrix}. \quad (31)$$

- (v) Repeat (iv) until the change in parameter values between the two iterations is less than a very

- small threshold: $|\lambda^{(k+1)} - \lambda^{(k)}| < \varepsilon, |\theta^{(k+1)} - \theta^{(k)}| < \varepsilon$
- (vi) The final estimated parameters can be obtained and are noted by $\hat{\lambda}_{MPS} = \lambda^{(k+1)}, \hat{\theta}_{MPS} = \theta^{(k+1)}$.

4.2. Maximum likelihood estimation

Under the progressively type-II censored sample, the likelihood function is ([18])

$$L(\lambda, \theta|x) = c \prod_{i=1}^m f(x_i) S(x_i; \lambda, \theta)^{P_i}, \quad (32)$$

Here, $c = n(n - P_1 - 1)(n - P_1 - P_2 - 2) \cdots (n - \sum_{i=1}^{m-1} (P_i + 1))$. Using Eqs (1) and (7), the likelihood function of the NGL distribution can be represented as

$$L(\lambda, \theta|x) = c \lambda^m \prod_{i=1}^m (1 - \theta + \lambda \theta x_i)(1 + \lambda \theta x_i)^{P_i} \exp(-\lambda x_i(1 + P_i)). \quad (33)$$

The log-likelihood function of the NGL distribution is

$$l(\lambda, \theta|x) = \ln c + m \ln \lambda - \lambda \sum_{i=1}^m x_i(1 + P_i) + \sum_{i=1}^m \ln(1 - \theta + \lambda \theta x_i) + \sum_{i=1}^m P_i \ln(1 + \lambda \theta x_i). \quad (34)$$

The ML estimator of λ and θ can be obtained by solving the following equations:

$$\frac{\partial l(\lambda, \theta|x)}{\partial \lambda} = \frac{m}{\lambda} - \sum_{i=1}^m x_i(1 + P_i) + \sum_{i=1}^m \frac{\theta x_i}{1 - \theta + \lambda \theta x_i} + \sum_{i=1}^m \frac{P_i \theta x_i}{1 + \lambda \theta x_i} = 0, \quad (35)$$

$$\frac{\partial l(\lambda, \theta|x)}{\partial \theta} = \sum_{i=1}^m \frac{\lambda x_i - 1}{1 - \theta + \lambda \theta x_i} + \sum_{i=1}^m \frac{P_i \lambda x_i}{1 + \lambda \theta x_i} = 0. \quad (36)$$

Newton Raphson algorithm to obtain the corresponding result. Due to the complexity of Eqs (35) and (36), it is difficult to judge the existence and uniqueness of $\hat{\lambda}_{ML}$ and $\hat{\theta}_{ML}$ by conventional numerical methods. In this paper, a graphical tool, the counter diagram, is employed, as referenced in Alotaibi et al. [19]. Setting the real values $(\lambda, \theta) = (0.5, 0.75)$, $(m, n) = (50, 100)$, $P = (0 * 49, 50)$ and generating progressively type-II censored samples, the counter diagram of $\hat{\lambda}_{ML}$ and $\hat{\theta}_{ML}$, which is obtained from Eq (34), is shown in Figure 4. It shows that the likelihood function obtained has only one obvious peak, that is, there is uniqueness.

Figure 4 shows the existence and uniqueness of ML estimates of λ and θ , with $(\hat{\lambda}_{ML}, \hat{\theta}_{ML}) = (0.855, 0.625)$. The ML estimation of $S(t)$ and $h(t)$ can be obtained by substituting $\hat{\lambda}_{ML}$ and $\hat{\theta}_{ML}$ [20]:

$$\hat{S}(t) = (1 + \hat{\lambda}_{ML} \hat{\theta}_{ML} t) e^{-\hat{\lambda}_{ML} t} \text{ and } \hat{h}(t) = \hat{\lambda}_{ML} - \frac{\hat{\lambda}_{ML} \hat{\theta}_{ML}}{1 + \hat{\lambda}_{ML} \hat{\theta}_{ML} t}.$$

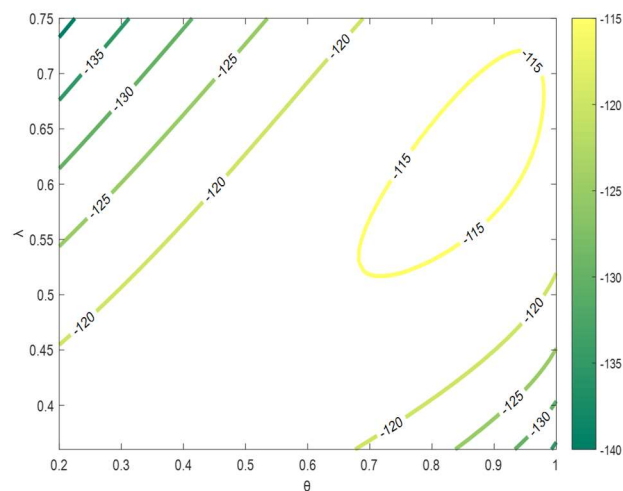


Figure 4. Counter diagram of $\hat{\lambda}_{ML}$ and $\hat{\theta}_{ML}$.

4.3. Asymptotic confidence interval

In the above analysis, we have obtained point estimates for λ , θ , $S(t)$, and $h(t)$. Point estimates provide a singular numerical approximation for unknown parameters, yet they fall short of encapsulating the full spectrum of uncertainty associated with these parameters. At this point, we need to transition from point estimation to interval estimation to get a more comprehensive understanding.

The ML estimator is asymptotically normal [21]. That is, $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, I^{-1}(\theta))$, where $\theta = (\lambda, \theta)$, and $I^{-1}(\theta)$ as the inverse of the information matrix for unknown parameters, and it can be obtained as follows:

$$I^{-1}(\hat{\theta}) \cong \left(\begin{array}{cc} -\frac{\partial^2 l(\theta|x)}{\partial \lambda^2} & -\frac{\partial^2 l(\theta|x)}{\partial \lambda \partial \theta} \\ -\frac{\partial^2 l(\theta|x)}{\partial \theta \partial \lambda} & -\frac{\partial^2 l(\theta|x)}{\partial \theta^2} \end{array} \right)^{-1} \bigg|_{\theta=\hat{\theta}} = \begin{pmatrix} I_{11} & I_{12} \\ I_{21} & I_{22} \end{pmatrix}^{-1}, \quad (37)$$

with the following elements

$$\frac{\partial^2 l(\theta|x)}{\partial \lambda^2} = -\frac{m}{\lambda^2} - \sum_{i=1}^m \frac{\theta^2 x_i^2}{(1-\theta+\lambda\theta x_i)^2} - \sum_{i=1}^m \frac{P_i \theta^2 x_i^2}{(1+\lambda\theta x_i)^2} \quad (38)$$

$$\frac{\partial^2 l(\theta|x)}{\partial \lambda \partial \theta} = \sum_{i=1}^m \frac{x_i}{(1-\theta+\lambda\theta x_i)^2} + \sum_{i=1}^m \frac{P_i x_i}{(1+\lambda\theta x_i)^2} = \frac{\partial^2 l(\theta|x)}{\partial \theta \partial \lambda} \quad (39)$$

$$\frac{\partial^2 l(\theta|x)}{\partial \theta^2} = -\sum_{i=1}^m \frac{(\lambda x_i - 1)^2}{(1-\theta+\lambda\theta x_i)^2} - \sum_{i=1}^m \frac{P_i \lambda^2 x_i^2}{(1+\lambda\theta x_i)^2}. \quad (40)$$

The $100(1 - \alpha)\%$ asymptotic confidence intervals (ACIs) for λ and θ are $\left(\hat{\lambda} \pm z_{\alpha/2} \sqrt{\text{Var}(\hat{\lambda})} \right)$

and $\left(\hat{\theta} \pm z_{\alpha/2} \sqrt{\text{Var}(\hat{\theta})}\right)$, respectively. Here, $z_{\alpha/2}$ represents the upper $\alpha/2$ percentage point of $N(0,1)$.

Similarly, the ACIs of $S(t)$ and $h(t)$ can also be constructed by calculating the corresponding variances, among which one of the most famous methods is Greene's delta method [22]. Under this method, the distribution of $\hat{S}(t)(\hat{h}(t))$ is approximately a normal distribution with mean $S(t)(h(t))$ and variance $\sigma_S^2 = \Delta_S I^{-1}(\lambda, \theta) \Delta_S^T|_{(\lambda=\hat{\lambda}, \theta=\hat{\theta})}$ ($\sigma_h^2 = \Delta_h I^{-1}(\lambda, \theta) \Delta_h^T|_{(\lambda=\hat{\lambda}, \theta=\hat{\theta})}$) [23], where $\Delta_S = \left(\frac{\partial S(x)}{\partial \lambda} \quad \frac{\partial S(x)}{\partial \theta}\right)$, $\Delta_h = \left(\frac{\partial h(x)}{\partial \lambda} \quad \frac{\partial h(x)}{\partial \theta}\right)$, with the following elements

$$\frac{\partial S(x)}{\partial \lambda} = (\theta - 1 - \lambda \theta x) x e^{-\lambda x}, \frac{\partial S(x)}{\partial \theta} = \lambda x e^{-\lambda x}, \frac{\partial h(x)}{\partial \lambda} = 1 - \frac{\theta}{(1 + \lambda \theta x)^2}, \frac{\partial h(x)}{\partial \theta} = -\frac{\lambda}{(\lambda \theta x + 1)^2}.$$

According to the above conclusions, the $100(1 - \alpha)\%$ ACIs of $S(t)$ and $h(t)$ for a given t are

$$\left(\hat{S}(t) \pm z_{\alpha/2} \sqrt{\sigma_S^2}\right), \left(\hat{h}(t) \pm z_{\alpha/2} \sqrt{\sigma_h^2}\right).$$

Since the asymptotic property of MPS estimator is similar to that of MLE estimator, this paper adopts the above method to obtain the interval estimation of MPS. For more details, see Ghosh and Jammalamadaka [24].

4.4. Bayesian estimation

In this section, the Bayesian estimator is used to estimate the parameters of the NGL distribution, and the corresponding highest posterior density (HPD) intervals are considered. The idea of Bayesian estimation in this paper refers to several papers [25–28]. Assume that λ and θ are independent and obey gamma distributions. The prior PDFs of λ and θ are:

$$\pi(\lambda | \sigma_1, \omega_1) = \frac{\omega_1^{\sigma_1}}{\Gamma(\sigma_1)} \lambda^{\sigma_1-1} e^{-\omega_1 \lambda}, \lambda > 0, \sigma_1 > 0, \omega_1 > 0, \quad (41)$$

$$\pi(\theta | \sigma_2, \omega_2) = \frac{\omega_2^{\sigma_2}}{\Gamma(\sigma_2)} \theta^{\sigma_2-1} e^{-\omega_2 \theta}, \theta > 0, \sigma_2 > 0, \omega_2 > 0. \quad (42)$$

Based on these assumptions, the joint prior density function of λ and θ can be obtained

$$\pi(\lambda, \theta) \propto \lambda^{\sigma_1-1} \theta^{\sigma_2-1} e^{-\omega_1 \lambda - \omega_2 \theta}, \lambda, \theta > 0, \omega_1, \omega_2 > 0. \quad (43)$$

According to the likelihood function with the prior knowledge and Bayes' theorem, the posterior density distribution of the unknown parameters λ and θ can be obtained:

$$\pi(\lambda, \theta | x) = \frac{L(\lambda, \theta | x) \pi(\lambda, \theta)}{\int_0^\infty \int_0^\infty L(\lambda, \theta | x) \pi(\lambda, \theta) d\lambda d\theta}. \quad (44)$$

Using Eqs (33) and (47), the posterior density distribution can be expressed:

$$\pi(\lambda, \theta | x) = \frac{\lambda^{m+\sigma_1-1} \theta^{\sigma_2-1} \exp(-\omega_1 \lambda - \omega_2 \theta - \lambda \sum_{i=1}^m x_i (1+P_i)) \prod_{i=1}^m (1-\theta + \lambda \theta x_i) (1 + \lambda \theta x_i)^{P_i}}{\int_0^\infty \int_0^\infty \lambda^{m+\sigma_1-1} \theta^{\sigma_2-1} \exp(-\omega_1 \lambda - \omega_2 \theta - \lambda \sum_{i=1}^m x_i (1+P_i)) \prod_{i=1}^m (1-\theta + \lambda \theta x_i) (1 + \lambda \theta x_i)^{P_i} d\lambda d\theta}. \quad (45)$$

The loss function is a key factor to make decisions in Bayesian estimation. In this paper, the SE and GE loss function are used to measure overestimation and underestimation in the investigation. The SE and GE loss functions, as symmetric loss function and asymmetric loss function, respectively, have different measures of the importance of overestimation and underestimation. The Bayesian estimator represents the posterior mean in the case of the SE loss function, and overestimation and underestimation have equal weight. The SE loss is defined as [29]

$$L_S(\phi(\lambda, \theta), \hat{\phi}(\lambda, \theta)) = (\phi(\lambda, \theta) - \hat{\phi}(\lambda, \theta))^2, \quad (46)$$

and the corresponding Bayesian estimator is

$$\hat{\phi}_S(\lambda, \theta) = E[\phi(\lambda, \theta) | x] = \int_0^\infty \int_0^\infty \phi(\lambda, \theta) \pi(\lambda, \theta | x) d\lambda d\theta. \quad (47)$$

The GE loss, which has a different tendency to weight overestimation and underestimation, is defined as [30]

$$L_G(\phi(\lambda, \theta), \hat{\phi}(\lambda, \theta)) \propto \left(\frac{\hat{\phi}(\lambda, \theta)}{\phi(\lambda, \theta)} \right)^\gamma - \gamma \log \left(\frac{\hat{\phi}(\lambda, \theta)}{\phi(\lambda, \theta)} \right) - 1, \gamma \neq 0, \quad (48)$$

here, γ is the parameter of the degree of asymmetry. Under GE loss, the Bayesian estimator is

$$\hat{\phi}_G(\lambda, \theta) = [E(\{\phi(\lambda, \theta)\}^{-\gamma} | x)]^{-1/\gamma} = \left\{ \int_0^\infty \int_0^\infty [\phi(\lambda, \theta)]^{-\gamma} \pi(\lambda, \theta | x) d\lambda d\theta \right\}^{-1/\gamma}. \quad (49)$$

Obviously, the integral of Eqs (47) and (49) cannot be calculated directly. Therefore, the MCMC approach, which is a very popular method for estimating parameters, is employed to calculate the corresponding HPD intervals and the Bayesian estimates (BE) of λ and θ . The full conditional posterior distribution of λ and θ as a key factor of the MCMC method can be derived by Eq (45)

$$\pi_1(\lambda | \theta, x) \propto \lambda^{m+\sigma_1-1} \exp(-\omega_1 \lambda - \lambda \sum_{i=1}^m x_i (1+P_i)) \prod_{i=1}^m (1-\theta + \lambda \theta x_i) (1 + \lambda \theta x_i)^{P_i}, \quad (50)$$

and

$$\pi_2(\theta | \lambda, x) \propto \theta^{\sigma_2-1} \exp(-\omega_2 \theta) \prod_{i=1}^m (1-\theta + \lambda \theta x_i) (1 + \lambda \theta x_i)^{P_i}. \quad (51)$$

Because of the nonlinearity of the full conditional posterior distribution of λ and θ , the Metropolis-Hastings (M-H) algorithm is applied to obtain the unknown parameters of Bayesian estimation. We assume the normal distribution as the proposed distribution to obtain the Bayesian estimation and HPD intervals of λ , θ , $S(t)$ and $h(t)$. Follow the next steps to generate the MCMC sample:

- (i) Set $k = 1$.
- (ii) Set the initial values of (λ, θ) to $(\lambda^{(0)}, \theta^{(0)})$.
- (iii) Generate λ^* and θ^* from $N(\lambda^{(k-1)}, \sigma_\lambda^2)$ and $N(\theta^{(k-1)}, \sigma_\theta^2)$, respectively. When $\lambda^* \leq 0$ or $\theta^* \notin (0,1)$, repeat step (iii), where $\lambda^{(k-1)}$ and $\theta^{(k-1)}$ represent previous state, σ_λ^2 and σ_θ^2 represent the variance of the previous state.
- (iv) Define acceptance probability $\omega(\lambda_{k-1}, \lambda^*) = \min(1, \frac{\pi_1(\lambda^* | \theta^{(k-1)}, x)}{\pi_1(\lambda_{k-1} | \theta^{(k-1)}, x)})$ and $\omega(\theta_{k-1}, \theta^*) =$

- $\min(1, \frac{\pi_2(\theta^*|\lambda^{(k)}, x)}{\pi_2(\theta^{(k-1)}|\lambda^{(k)}, x)}).$
- (v) Generate $u_{(1)}$ and $u_{(2)}$ from the uniform distribution $U(0,1)$.
 - (vi) If $u_{(1)} \leq \omega(\lambda_{k-1}, \lambda^*)$, then $\lambda^{(k)} = \lambda^*$, otherwise $\lambda^{(k)} = \lambda^{(k-1)}$.
 - (vii) If $u_{(2)} \leq \omega(\theta_{k-1}, \theta^*)$, then $\theta^{(k)} = \theta^*$, otherwise $\theta^{(k)} = \theta^{(k-1)}$.
 - (viii) Calculate $S^{(k)}(t)$ and $h^{(k)}(t)$ according to the following formulas: $S^{(k)}(t) = (1 + \lambda^{(k)}\theta^{(k)}t)e^{-\lambda^{(k)}t}$ and $h^{(k)}(t) = \lambda^{(k)} - \frac{\lambda^{(k)}\theta^{(k)}}{1+\lambda^{(k)}\theta^{(k)}t}$, where $t > 0$.
 - (ix) Set $k = k + 1$.
 - (x) Repeat (iii)–(ix) L times to get $\{\lambda^{(k)}\}, \{\theta^{(k)}\}, \{h^{(k)}(t)\}$ and $\{S^{(k)}(t)\} (k = 1, 2, \dots, L)$, and discard the first L^* samples of $\{\lambda^{(k)}\}, \{\theta^{(k)}\}, \{h^{(k)}(t)\}$ and $\{S^{(k)}(t)\}$ to eliminate the influence of initial value selection.
 - (xi) Based on SE and GE loss functions, calculate the Bayesian estimation and HPD intervals of $\lambda, \theta, S(t)$ and $h(t)$. Take λ for example:
 - Compute the Bayesian estimate of $\lambda: \hat{\lambda}_S = \frac{1}{l} \sum_{i=L^*+1}^L \lambda^{(i)}$ and $\hat{\lambda}_G = \frac{1}{l} (\sum_{i=L^*+1}^L [\lambda^{(i)}]^{-\gamma})^{-1/\gamma}$, where $l = L - L^*$.
 - Create the HPD interval of λ [31]: Let $\lambda_{(L^*+1)}, \lambda_{(L^*+2)}, \dots, \lambda_{(L)}$ be the ascending values of $\lambda_{(L^*+1)}, \lambda_{(L^*+2)}, \dots, \lambda_{(L)}$, the $100(1 - \alpha)\%$ HPD interval of λ can be approximated to $(\lambda_{(k^*)}, \lambda_{(k^* + [(1-\alpha)l])})$, where $k^* \in \{L^* + 1, L^* + 2, \dots, L\}$ is selected according to the formula $\lambda_{(k^* + [(1-\alpha)l])} - \lambda_{(k^*)} = \min_{L^*+1 \leq k \leq L^* + [\alpha l]} (\lambda_{(k^* + [(1-\alpha)l])} - \lambda_{(k^*)})$, where $[x]$ is the downward integer of x , that is, the largest integer less than or equal to x .

5. Monte Carlo simulation

In this section, a Monte Carlo simulation will be performed to demonstrate and compare the performance of the above estimators for the NGL distribution in parameter estimation.

5.1. Simulation plan

We simulate 1000 progressively type-II censored samples of the NGL(0.5,0.75) based on the parameter selection of n (Total number of samples), m (Number of valid samples), and P (Censored schemes). Meanwhile, in order to reasonably evaluate the estimates of $S(t)$ and $h(t)$, first obtain their true values at $t = 0.5$, which are 0.9248 and 0.1842, respectively. In addition, $n(= 50, 90)$ is determined and the valid sample proportion is used to determine that the value of m meets $m/n(= 60\%, 80\%)$. Also, the three progressively censored schemes used are shown in Table 2.

Next, the specific steps to generate progressively type-II censored samples are given [32]:

- (i) Generate m observations ς_i (for $i = 1, 2, \dots, m$) that follow uniform distribution $U(0,1)$.
- (ii) Give the corresponding value according to the specific censored schemes, set $\zeta_i = \varsigma_i^{1/(i + \sum_{k=m-i+1}^m P_k)}$ (for $i = 1, 2, \dots, m$).
- (iii) Set $\xi_i = 1 - \zeta_m \zeta_{m-1} \cdots \zeta_{m-i+1}$ (for $i = 1, 2, \dots, m$).
- (iv) Generate progressively type-II censored samples x_i of the NGL (0.5, 0.75), set

$$x_i = \left\lceil -0.75W\left(\frac{e^{-1/0.75}(\xi_i - 1)}{0.75}\right) - 1 \right\rceil / 0.375.$$

Table 2. Three kinds of censored schemes.

Scheme	
1	$P=(0 * (m - 1), n - m)$
2	$P=([(n - m)/2], 0 * (m - 2), n - m - [(n - m)/2])$
3	$P=(n - m, 0 * (m - 1))$

In Table 2, $P = (6, 0, 0, 0, 1)$ stands for $P = (6, 0 * 4, 1)$.

After obtaining progressively type-II censored samples, MPS estimates and ML estimates and 95% ACIs of λ , θ , $S(t)$ and $h(t)$ are calculated using Matlab R2016a, as well as Bayesian estimates based on SE and GE ($\gamma(= -2)$) loss function and HPD interval. The large 12,000 M-H samples are generated by M-H sampler, and then the first 2000 samples are deleted as fluctuation samples. Then we complete the Bayesian estimator by setting up two prior sets called *Prior - a*: $(\sigma_1, \sigma_2, \omega_1, \omega_2) = (8, 10, 10, 5)$ and *Prior - b*: $(\sigma_1, \sigma_2, \omega_1, \omega_2) = (4, 5, 5, 2.5)$. Repeat 1000 times to ensure the accuracy and stability of the estimation results. The evaluation criteria for verifying the reliability of the estimation method are shown in Table 3.

Table 3. Evaluation criteria of point estimation and interval estimation.

Name	Formula of errors
AE (average estimates)	$\bar{\hat{\lambda}} = \frac{1}{1000} \sum_{i=1}^{1000} \hat{\lambda}_i$
RMSE (root mean squared errors)	$RMSE_{\bar{\hat{\lambda}}} = \sqrt{\frac{1}{1000} \sum_{i=1}^{1000} (\hat{\lambda}_i - \lambda)^2}$
MRAB (mean relative absolute biases)	$MRAB_{\bar{\hat{\lambda}}} = \frac{1}{1000} \sum_{i=1}^{1000} \frac{1}{\lambda} \hat{\lambda}_i - \lambda $
ACL (average confidence lengths)	$ACL_{\lambda}^{(1-\alpha)\%} = \frac{1}{1000} \sum_{i=1}^{1000} (R_{\hat{\lambda}_i} - L_{\hat{\lambda}_i})$
CP (coverage percentages)	$CP_{\lambda}^{(1-\alpha)\%} = \frac{1}{1000} \sum_{i=1}^{1000} l_{(L_{\hat{\lambda}_i}; R_{\hat{\lambda}_i})}(\lambda)$

Where $l_A()$ represents the indicator function, and R_0 and L_0 denote the upper and lower bounds, respectively, for each $(1 - \alpha)\%$ ACI/HPD interval.

5.2. Simulation result

RMSE, MRAB, ACL, and CP of λ , θ , $S(t)$ and $h(t)$ estimates of various estimators are shown using heat maps in Figures 5–8. The corresponding specific numerical results are shown in the appendix [Tables 4–11]. Here, SE-Pa represents the estimate of *Prior - a* by Bayesian estimator based on SE loss, GE-Pb represents the estimate of *Prior - b* by Bayesian estimator based on GE

loss (for $\gamma(= -2)$), and ACL-MPS corresponds to the interval estimate of the MPS.

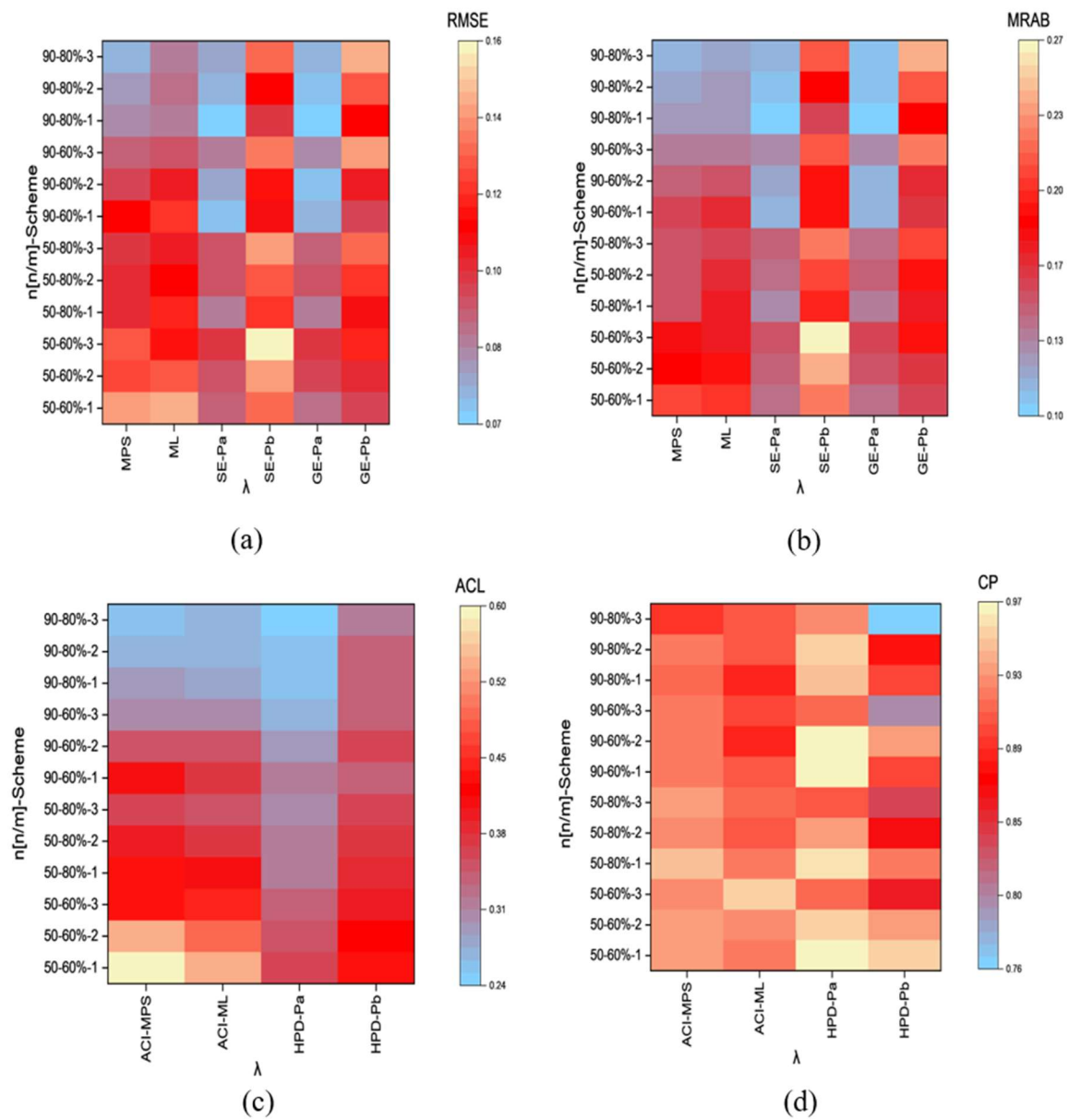


Figure 5. Heat maps of the associated estimates of λ .

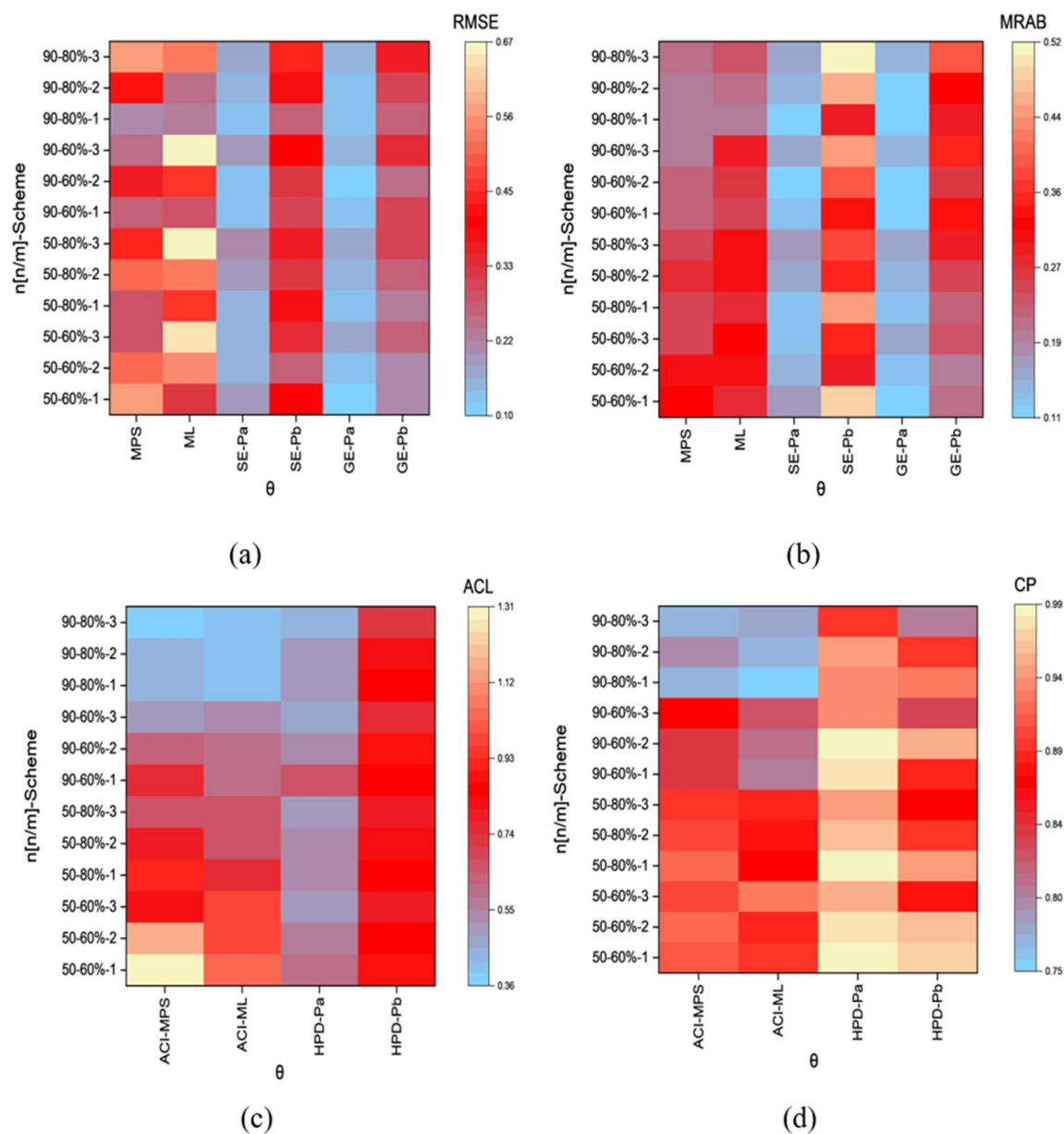


Figure 6. Heat maps of the associated estimates of θ .

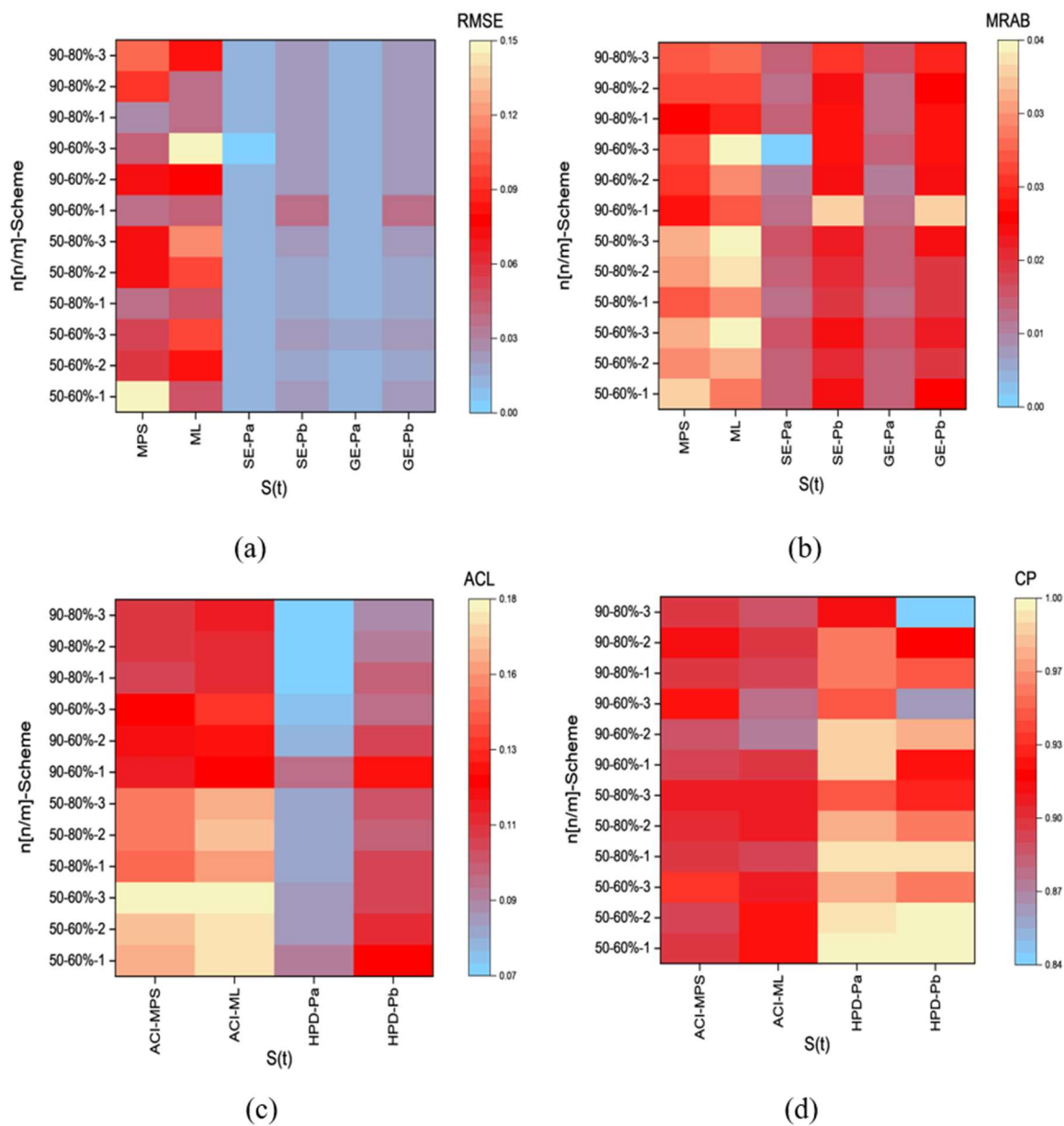


Figure 7. Heat maps of the associated estimates of $S(t)$.

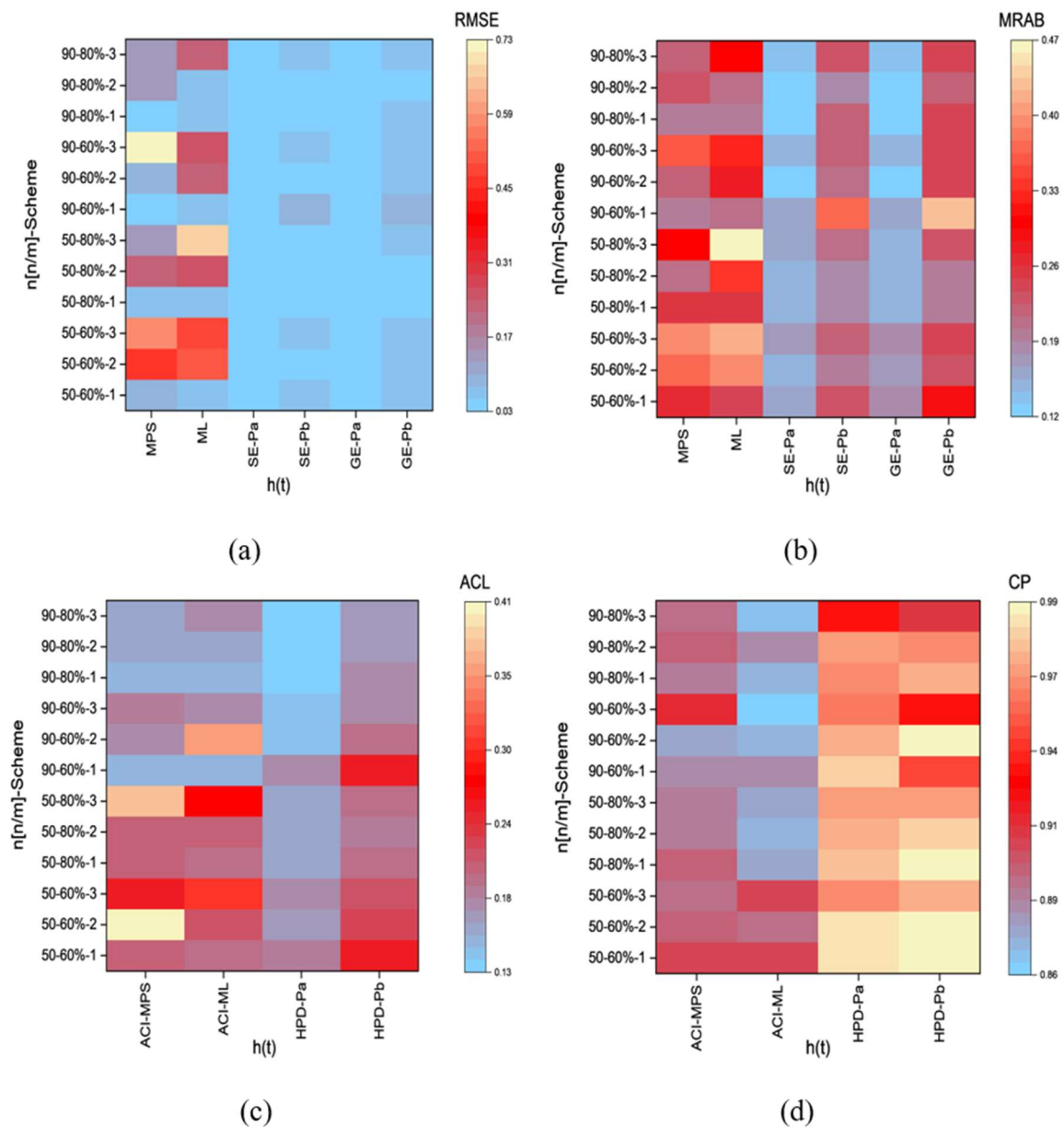


Figure 8. Heat maps of the associated estimates of $h(t)$.

Table 4. Point estimation of λ by estimator under different conditions (AE: first column; RMSE: second column; MRAB: third column).

(n, m)	Scheme	MPS		ML				BE					
								SE			GE		
								$Prior - a$			$Prior - a$		
$Prior - b$			$Prior - b$										
-2													
(50,30)	1	0.4897	0.1395	0.2052	0.5111	0.1412	0.2035	0.5258	0.0906	0.1399	0.5317	0.0888	0.1380
								0.3905	0.1297	0.2240	0.4513	0.0967	0.1557
	2	0.4967	0.1229	0.1846	0.5043	0.1284	0.1899	0.5174	0.0936	0.1453	0.5278	0.0974	0.1509
								0.3826	0.1401	0.2393	0.4481	0.1038	0.1634
	3	0.5404	0.1267	0.1785	0.507	0.1161	0.1729	0.5175	0.1000	0.1534	0.5254	0.1016	0.1582
								0.371	0.1553	0.2665	0.4373	0.1195	0.1916
(50,40)	1	0.4954	0.1023	0.1525	0.5095	0.1188	0.1716	0.5176	0.0857	0.1300	0.5261	0.0868	0.1334
								0.4168	0.1205	0.1953	0.4306	0.1102	0.1761
	2	0.5091	0.1025	0.1503	0.5157	0.1131	0.1682	0.508	0.0937	0.1403	0.515	0.0952	0.1451
								0.411	0.1278	0.2056	0.4206	0.1215	0.1903
	3	0.5257	0.1015	0.149	0.5133	0.1068	0.156	0.5074	0.0958	0.1434	0.5178	0.0927	0.1398
								0.3913	0.1401	0.2235	0.4191	0.1301	0.2073
(90,54)	1	0.4983	0.111	0.1566	0.5058	0.1219	0.1667	0.5118	0.0717	0.1108	0.5157	0.0728	0.1129
								0.4104	0.1099	0.1881	0.4255	0.0981	0.1634
	2	0.5062	0.0973	0.1435	0.5135	0.107	0.1505	0.5157	0.0769	0.1160	0.5229	0.0716	0.1108
								0.4077	0.1157	0.1934	0.4249	0.1054	0.1700
	3	0.5211	0.0912	0.1324	0.5115	0.0955	0.1362	0.5084	0.0842	0.1261	0.5156	0.0832	0.1269
								0.4096	0.1337	0.2147	0.4003	0.1376	0.2222
(90,72)	1	0.5036	0.0817	0.1205	0.5049	0.0865	0.1257	0.511	0.0661	0.0981	0.5185	0.0669	0.099
								0.4386	0.0989	0.1588	0.4112	0.1125	0.1874
	2	0.5069	0.0798	0.118	0.5102	0.0869	0.1252	0.5083	0.0731	0.1067	0.5186	0.0716	0.1046
								0.4306	0.1135	0.183	0.3984	0.1277	0.2115
	3	0.5239	0.0749	0.1106	0.5106	0.0847	0.1182	0.5115	0.0768	0.1113	0.5152	0.0722	0.1041
								0.4228	0.1311	0.2126	0.3852	0.1432	0.2415

Table 5. Point estimation of θ by estimator under different conditions (AE: first column; RMSE: second column; MRAB: third column).

(n, m)	Scheme	MPS		ML				BE						
								SE			GE			
								<i>Prior – a</i>			<i>Prior – a</i>			
								<i>Prior – b</i>			<i>Prior – b</i>			
														-2
(50,30)	1	0.6686	0.5743	0.3249	0.7282	0.3277	0.2828	0.7389	0.1897	0.1652	0.7831	0.108	0.1166	
								0.3881	0.3965	0.4825	0.5906	0.2024	0.2125	
	2	0.6529	0.5109	0.3051	0.7201	0.5525	0.3064	0.7448	0.1509	0.1402	0.7792	0.1305	0.1306	
								0.5263	0.2745	0.2983	0.6020	0.2021	0.1976	
	3	0.7168	0.2925	0.2532	0.6427	0.6425	0.3192	0.7142	0.161	0.1319	0.7689	0.1620	0.1503	
								0.4916	0.3373	0.3447	0.5740	0.2588	0.2368	
(50,40)	1	0.6688	0.2753	0.2539	0.7385	0.4545	0.2849	0.7636	0.1484	0.1337	0.7840	0.1244	0.126	
								0.4136	0.3722	0.4489	0.5818	0.2229	0.2243	
	2	0.7019	0.5158	0.2736	0.7586	0.5247	0.307	0.7419	0.1820	0.1581	0.7640	0.1600	0.1465	
								0.4847	0.3278	0.3537	0.5564	0.2576	0.2582	
	3	0.7139	0.4281	0.2498	0.7301	0.6584	0.3084	0.7320	0.2008	0.1717	0.7731	0.1638	0.1503	
								0.4697	0.3551	0.3737	0.5324	0.3006	0.2906	
(90,54)	1	0.6973	0.257	0.2188	0.7721	0.2877	0.2455	0.7454	0.1392	0.1282	0.7553	0.126	0.1121	
								0.5003	0.3001	0.3329	0.4933	0.2947	0.3423	
	2	0.7081	0.3537	0.228	0.7554	0.4585	0.2615	0.7646	0.1363	0.1173	0.7868	0.1052	0.1076	
								0.4593	0.3279	0.3876	0.5536	0.2381	0.2622	
	3	0.7152	0.2474	0.2033	0.7242	0.6738	0.2884	0.7426	0.1831	0.1551	0.7734	0.1594	0.1448	
								0.4140	0.4003	0.4497	0.4875	0.3416	0.3507	
(90,72)	1	0.7284	0.2088	0.1898	0.7674	0.2309	0.2009	0.7629	0.1267	0.1141	0.7891	0.1242	0.1168	
								0.5251	0.2668	0.2999	0.5261	0.2624	0.2993	
	2	0.7159	0.4108	0.2002	0.7461	0.2387	0.2108	0.7584	0.1549	0.1343	0.7915	0.1274	0.1191	
								0.4086	0.3846	0.4566	0.505	0.2954	0.3278	
	3	0.7502	0.5628	0.2173	0.7647	0.5311	0.236	0.7626	0.1778	0.1529	0.7863	0.1524	0.1393	
								0.361	0.4457	0.5231	0.4594	0.3608	0.3915	

Table 6. Point estimation of $S(t)$ by estimator under different conditions (AE: first column; RMSE: second column; MRAB: third column).

(n, m)	Scheme	MPS		ML				BE					
								SE			GE		
								$Prior - a$			$Prior - a$		
$Prior - b$			$Prior - b$										
-2													
(50,30)	1	0.9226	0.1517	0.0392	0.8771	0.0541	0.0327	0.9253	0.0174	0.0147	0.9271	0.017	0.0149
								0.9051	0.0262	0.0228	0.9046	0.0272	0.0236
	2	0.9145	0.0598	0.0336	0.9228	0.0879	0.0364	0.9274	0.0164	0.0144	0.9283	0.0174	0.0151
								0.9088	0.023	0.0197	0.9095	0.0227	0.019
	3	0.9184	0.0589	0.0367	0.9448	0.1018	0.041	0.9269	0.0188	0.0162	0.9272	0.0191	0.0163
								0.9068	0.0262	0.0223	0.9084	0.0252	0.0214
(50,40)	1	0.9164	0.0412	0.0303	0.9331	0.0532	0.0343	0.9302	0.0159	0.0140	0.9292	0.0159	0.0141
								0.9087	0.0217	0.0190	0.9097	0.0213	0.0182
	2	0.9197	0.0774	0.0349	0.9204	0.1022	0.0397	0.9288	0.0166	0.0144	0.9284	0.0165	0.0145
								0.9081	0.0232	0.0197	0.9087	0.0218	0.019
	3	0.9165	0.0763	0.0357	0.9257	0.1217	0.0412	0.9278	0.0183	0.0158	0.9294	0.0175	0.0154
								0.9068	0.0253	0.0214	0.9061	0.0255	0.0223
(90,54)	1	0.9213	0.042	0.0245	0.9425	0.049	0.0296	0.9226	0.0177	0.0142	0.9234	0.017	0.0135
								0.8886	0.0406	0.0392	0.8889	0.0403	0.0389
	2	0.9223	0.079	0.0276	0.9306	0.0806	0.0335	0.9292	0.0144	0.0127	0.9294	0.0144	0.0127
								0.9038	0.0253	0.023	0.9041	0.0250	0.0229
	3	0.9179	0.049	0.0285	0.9457	0.1544	0.0418	0.9289	0.0041	0.0044	0.9293	0.0180	0.0156
								0.9028	0.0279	0.0248	0.9030	0.0274	0.0246
(90,72)	1	0.9244	0.0334	0.0237	0.922	0.0397	0.026	0.9308	0.0175	0.0152	0.9309	0.0155	0.0136
								0.9018	0.0265	0.0252	0.9022	0.0261	0.0248
	2	0.9196	0.0963	0.0285	0.9333	0.0422	0.0283	0.9308	0.0158	0.014	0.9317	0.0150	0.0134
								0.9042	0.0244	0.0229	0.9037	0.0249	0.0231
	3	0.924	0.1124	0.0304	0.938	0.0862	0.0312	0.9307	0.0179	0.0156	0.9315	0.0179	0.0158
								0.9006	0.0289	0.0272	0.9012	0.0286	0.0264

Table 7. Point estimation of $h(t)$ by estimator under different conditions (AE: first column; RMSE: second column; MRAB: third column).

(n, m)	Scheme	MPS		ML				BE					
								SE		GE			
								<i>Prior – a</i>		<i>Prior – a</i>			
								<i>Prior – b</i>		<i>Prior – b</i>			
-2													
(50,30)	1	0.2011	0.0741	0.2683	0.1878	0.0592	0.242	0.1938	0.0404	0.1667	0.1971	0.0428	0.1792
								0.2215	0.0538	0.2311	0.2339	0.0656	0.2881
	2	0.2299	0.4497	0.3735	0.1661	0.4973	0.3925	0.1883	0.0362	0.1541	0.1931	0.0413	0.1682
								0.2133	0.0478	0.2005	0.2211	0.0536	0.2255
	3	0.217	0.5736	0.3982	0.1655	0.4777	0.4144	0.1897	0.0401	0.1710	0.1953	0.0436	0.1843
								0.2170	0.0527	0.2230	0.2219	0.0561	0.2406
(50,40)	1	0.2058	0.0627	0.2552	0.1984	0.0647	0.2553	0.1824	0.0338	0.1451	0.1896	0.0350	0.1476
								0.2116	0.0425	0.1810	0.2169	0.0466	0.2001
	2	0.1942	0.2237	0.2046	0.2052	0.2439	0.3358	0.1840	0.0339	0.1451	0.1898	0.0348	0.1485
								0.2123	0.0443	0.1859	0.2175	0.0459	0.2009
	3	0.2131	0.1273	0.3069	0.1894	0.675	0.4708	0.1862	0.0366	0.1564	0.1887	0.0363	0.1543
								0.2151	0.0482	0.2020	0.2231	0.0531	0.2340
(90,54)	1	0.197	0.0474	0.1958	0.1794	0.0523	0.2025	0.1967	0.0383	0.1571	0.2018	0.0404	0.1628
								0.2519	0.0781	0.3699	0.2634	0.0884	0.4317
	2	0.1947	0.0904	0.2204	0.1699	0.2216	0.2725	0.1815	0.0296	0.1294	0.1877	0.0305	0.1268
								0.2205	0.0474	0.2114	0.2282	0.0537	0.2460
	3	0.1898	0.7276	0.3584	0.1554	0.2483	0.326	0.1836	0.0341	0.148	0.1873	0.0352	0.1485
								0.221	0.0501	0.2189	0.2265	0.0532	0.2416
(90,72)	1	0.1911	0.0475	0.1937	0.1813	0.0516	0.1997	0.1798	0.0280	0.1233	0.1831	0.0290	0.1262
								0.2228	0.0471	0.2169	0.2292	0.0528	0.2486
	2	0.184	0.1361	0.2317	0.1677	0.0541	0.2117	0.1791	0.0290	0.1271	0.1816	0.0276	0.1203
								0.2171	0.0419	0.1889	0.2238	0.0478	0.2201
	3	0.1874	0.1211	0.2138	0.1594	0.2271	0.3026	0.1799	0.0326	0.1391	0.1818	0.0326	0.1396
								0.2235	0.0496	0.2268	0.2279	0.0533	0.2454

Table 8. Interval estimation of λ by estimator under different conditions.

(n, m)	Scheme	ACI		HPD					
		MPS		ML		<i>Prior – a</i>		<i>Prior – b</i>	
		ACL	CP	ACL	CP	ACL	CP	ACL	CP
(50,30)	1	0.5959	0.929	0.5407	0.919	0.3655	0.965	0.4281	0.951
	2	0.5383	0.932	0.4926	0.924	0.3515	0.949	0.422	0.932
	3	0.4397	0.923	0.4459	0.951	0.3365	0.910	0.3939	0.853
(50,40)	1	0.4299	0.945	0.4102	0.915	0.319	0.956	0.3898	0.918
	2	0.3964	0.926	0.3765	0.903	0.3102	0.929	0.3732	0.864
	3	0.3581	0.928	0.3554	0.912	0.3021	0.905	0.3628	0.834
(90,54)	1	0.4106	0.919	0.3758	0.906	0.3131	0.969	0.3428	0.894
	2	0.3530	0.918	0.3558	0.880	0.2922	0.963	0.3662	0.932
	3	0.3037	0.918	0.3012	0.900	0.2713	0.912	0.3400	0.804
(90,72)	1	0.2872	0.913	0.2770	0.886	0.2583	0.947	0.3430	0.900
	2	0.2674	0.919	0.2657	0.906	0.2523	0.952	0.3364	0.876
	3	0.2542	0.888	0.2642	0.902	0.2366	0.925	0.3188	0.763

Table 9. Interval estimation of θ by estimator under different conditions.

(n, m)	Scheme	ACI				HPD			
		MPS		ML		<i>Prior – a</i>		<i>Prior – b</i>	
		ACL	CP	ACL	CP	ACL	CP	ACL	CP
(50,30)	1	1.3061	0.908	1.0387	0.891	0.5877	0.986	0.8699	0.965
	2	1.1711	0.921	0.9627	0.888	0.5510	0.972	0.8449	0.955
	3	0.8076	0.904	0.9926	0.926	0.5016	0.950	0.7761	0.875
(50,40)	1	0.9013	0.921	0.7698	0.873	0.5367	0.980	0.8515	0.939
	2	0.7790	0.902	0.6721	0.881	0.5254	0.960	0.8133	0.898
	3	0.6540	0.895	0.671	0.888	0.5017	0.944	0.7724	0.868
(90,54)	1	0.7660	0.842	0.5899	0.800	0.6550	0.978	0.8402	0.887
	2	0.6332	0.836	0.6084	0.805	0.5377	0.981	0.8682	0.951
	3	0.5169	0.867	0.5271	0.826	0.4844	0.931	0.7569	0.830

Continued on next page

(n, m)	Scheme	ACI				HPD			
		MPS		ML		<i>Prior – a</i>		<i>Prior – b</i>	
		ACL	CP	ACL	CP	ACL	CP	ACL	CP
(90,72)	1	0.4541	0.767	0.4141	0.748	0.4985	0.938	0.8488	0.925
	2	0.4432	0.788	0.4027	0.765	0.4908	0.945	0.8223	0.898
	3	0.3628	0.769	0.4004	0.772	0.4378	0.897	0.7363	0.802

Table 10. Interval estimation of $S(t)$ by estimator under different conditions.

(n, m)	Scheme	ACI				HPD			
		MPS		ML		<i>Prior – a</i>		<i>Prior – b</i>	
		ACL	CP	ACL	CP	ACL	CP	ACL	CP
(50,30)	1	0.1615	0.899	0.1714	0.924	0.0932	0.995	0.1262	0.997
	2	0.1653	0.895	0.1713	0.926	0.0865	0.991	0.1128	0.993
	3	0.1778	0.935	0.1751	0.911	0.0864	0.971	0.1089	0.958
(50,40)	1	0.1483	0.899	0.1572	0.892	0.0805	0.987	0.1059	0.991
	2	0.15	0.902	0.165	0.907	0.08	0.972	0.1014	0.957
	3	0.1517	0.907	0.1626	0.908	0.0813	0.949	0.1019	0.932
(90,54)	1	0.1196	0.892	0.1237	0.898	0.0946	0.983	0.1283	0.926
	2	0.1207	0.887	0.1277	0.874	0.0764	0.981	0.1075	0.973
	3	0.1256	0.925	0.1344	0.875	0.0754	0.947	0.0968	0.859
(90,72)	1	0.1087	0.897	0.1157	0.894	0.0706	0.958	0.0997	0.946
	2	0.1117	0.913	0.1153	0.9	0.0699	0.957	0.092	0.919
	3	0.1113	0.898	0.1187	0.886	0.069	0.914	0.0905	0.837

Table 11. Interval estimation of $h(t)$ by estimator under different conditions.

(n, m)	Scheme	ACI				HPD			
		MPS		ML		<i>Prior – a</i>		<i>Prior – b</i>	
		ACL	CP	ACL	CP	ACL	CP	ACL	CP
(50,30)	1	0.204	0.907	0.2012	0.910	0.1885	0.986	0.2549	0.993
	2	0.4103	0.897	0.2187	0.896	0.1740	0.986	0.2233	0.993
	3	0.2534	0.895	0.2982	0.907	0.1767	0.964	0.216	0.975
(50,40)	1	0.2046	0.898	0.1978	0.878	0.1559	0.979	0.2021	0.991
	2	0.2115	0.891	0.2056	0.873	0.1559	0.975	0.1932	0.982
	3	0.3762	0.892	0.2784	0.878	0.1603	0.967	0.1964	0.971
(90,54)	1	0.1526	0.887	0.1474	0.887	0.1791	0.981	0.2538	0.948
	2	0.1762	0.877	0.3585	0.873	0.143	0.973	0.2021	0.991
	3	0.1885	0.917	0.1827	0.862	0.1446	0.961	0.1834	0.932
(90,72)	1	0.1486	0.889	0.1461	0.873	0.1288	0.965	0.183	0.972
	2	0.1637	0.897	0.1573	0.885	0.1279	0.97	0.167	0.963
	3	0.1599	0.894	0.1826	0.87	0.1293	0.932	0.1692	0.911

From the heat maps in Figures 5–8 and Tables 4–11, the following conclusions can be drawn:

(1) All estimates of λ , θ , $S(t)$ and $h(t)$ are good estimators because of low RMSE, MRAB, and ACL values and high CP values. Through changing the color of the heat-maps to from down to up, we can find that, in most cases, with the increase of n and m , the estimation performance of all obtained estimators will improve, corresponding to lower RMSE, MRAB, ACL, and CP values.

(2) Among all estimates, Bayesian estimates based on SE loss and GE loss are more accurate than MPS and ML estimates because of lower RMSE, MRAB, and ACL values and higher CP values.

(3) Different prior parameters will affect the effectiveness of Bayesian estimation. In Bayesian estimation, the Bayesian estimator based on SE loss and the Bayesian estimator based on GE loss have better estimation performance under the gamma prior function with $Prior - a$ as parameter than $Prior - b$ because the variance of $Prior - a$ is smaller.

(4) Different censored schemes may affect the estimated effectiveness to some extent. In most cases, estimates based on scheme-1 work better. In the face of progressively type-II censored samples of the NGL distribution, the Bayesian estimation under the gamma prior function with $Prior - a$ can be used to estimate the corresponding unknown parameters and related functions.

(5) The Bayesian estimation under the gamma prior function with $Prior - a$ is overestimated while with $Prior - b$ is underestimated.

(6) The point (or interval) estimation of the MPS estimates and ML estimates of $S(t)$ is significantly weaker than the result obtained by Bayes.

(7) In summary, with a more complete amount of data from progressively type-II censored, Bayesian estimation via M-H algorithm can obtain better estimates when estimating unknown parameters of the NGL distribution.

6. Practical application

This section demonstrates the flexibility of the proposed distribution and the usefulness of the various estimation methods through a practical application. The dataset, which was previously used by Bekker et al. [33] and later applied by Habib et al. [34], is the annual survival time of 46 patients who received chemotherapy and radiation therapy. The specific data is shown in Table 12.

Table 12. The annual survival time of 46 patients.

0.047	0.115	0.121	0.132	0.164	0.197	0.203	0.26	0.282	0.296	0.334	0.395
0.458	0.466	0.501	0.507	0.529	0.534	0.54	0.57	0.641	0.644	0.696	0.841
0.863	1.099	1.219	1.271	1.326	1.447	1.485	1.553	1.581	1.589	2.178	2.343
2.416	2.444	2.825	2.83	3.578	3.658	3.743	3.978	4.003	4.033		

In order to test whether the NGL distribution is suitable for the dataset, ML estimation is first used to estimate λ and θ , and the corresponding values with standard errors (St) are obtained as 0.90788 (0.33141) and 0.20264 (0.41243). Then the Kolmogorov-Smirnov (KS) value and P-value can be obtained as 0.10158 (0.69157). Compared with the truncated Nadarajah-Haghighi Raykeigh distribution, which presents a KS value with a P-value of 0.1080 (0.6307) for this dataset [34], it can be considered that the NGL distribution has a better fitting effect and is suitable for this dataset. As illustrated in Figure 9, the NGL distribution to the fitting effect of the dataset is demonstrated. This includes the fitted CDF, the probability-probability (PP), the scaled total time on test (TTT) transform [35], and the counter of the log-likelihood function. Figure 9 indicates that the NGL

distribution is very close to the real data distribution and the existence and uniqueness of the obtained MLE estimates $\hat{\lambda}$ and $\hat{\theta}$.

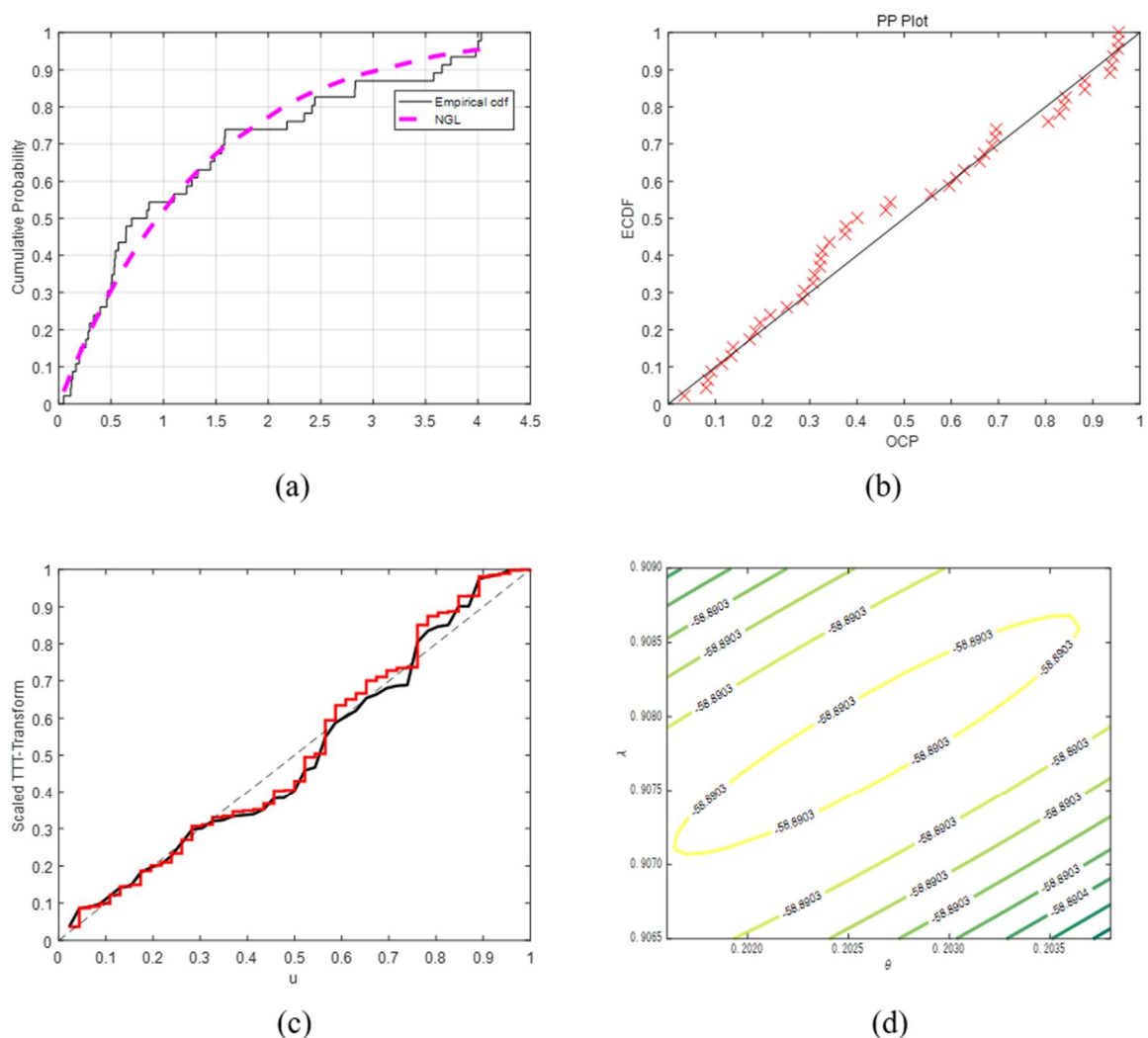


Figure 9. (a) Fitted CDF of NGL, (b) PP, (c) scaled TTT-Transform, (d) counter of log-likelihood function from the annual survival time of 46 patients.

Let $m=20$, and three types of progressively type-II censored samples are obtained from this dataset, as shown in Table 13. The point estimates of λ and θ are obtained by the above estimation method (MPS, ML, Bayesian estimation) and are shown in Table 14. At the same time, standard errors (St.es) are used to judge the accuracy of the estimated results. Here, the SE.es of λ and θ obtained by MPS and ML estimation are the square root of the diagonal elements of $I^{-1}(\theta)$, respectively.

Table 13. Three progressively type-II censored samples from the annual survival time.

Sample	Scheme										
S_1	(0*19,26)	0.047	0.115	0.121	0.132	0.164	0.197	0.203	0.26	0.282	0.296
		0.334	0.395	0.458	0.466	0.501	0.507	0.529	0.534	0.54	0.57
S_2	(13,0*18,13)	0.047	0.501	0.507	0.529	0.534	0.54	0.57	0.641	0.644	0.696
		0.841	0.863	1.099	1.219	1.271	1.326	1.447	1.485	1.553	4.033
S_3	(26,0*19)	0.047	1.271	1.326	1.447	1.485	1.553	1.581	1.589	2.178	2.343
		2.416	2.444	2.825	2.83	3.578	3.658	3.743	3.978	4.003	4.033

Table 14. Point estimates (St.es) of λ, θ from the annual survival time.

S_i	Parameter	MPS	MLE	SE	GE $\gamma = -2$
S_1	λ	0.63616(0.11222)	0.74652(0.15355)	0.2378(0.0579)	0.2539(0.0557)
	θ	0.999(0.17427)	0.999(0.4182)	0.7465(0.2911)	0.8367(0.2369)
S_2	λ	0.21672(0.038498)	0.1919(0.033919)	0.0663(0.0120)	0.0670(0.0112)
	θ	0.91092(0.14331)	0.95217(0.15405)	0.7755(0.1317)	0.7933(0.1235)
S_3	λ	0.01693(0.00463)	0.01611(0.00455)	0.0523(0.0216)	0.0935(0.0141)
	θ	0.001(0.42653)	0.001(0.44385)	0.2585(0.3328)	0.8966(0.0917)

Through the M-H algorithm, 5000 MCMC samples are generated, and the first 1000 samples are omitted. The sample sequence $\{\lambda_{(i)}\}, \{\theta_{(i)}\}$ formed can approximately obey the posterior distribution, and the variance of the posterior distribution can be estimated by the sample variance, that is:

$$St_{\lambda} = \sqrt{\frac{1}{3999} \sum_{i=1}^{4000} (\lambda_{(i)} - \hat{\lambda})^2}, St_{\theta} = \sqrt{\frac{1}{3999} \sum_{i=1}^{4000} (\theta_{(i)} - \hat{\theta})^2}.$$

It can be seen from Table 14 that under three different censored schemes:

- All estimation methods have a better fitting effect when fitting the dataset.
- The results estimated by ML and MPS are similar, and the results estimated by Bayes under two different losses are similar.
- The MPS and ML estimation fit the scale parameter λ preferentially when fitting, which makes the scale parameter more accurate but tends to lead to overflow of the shape parameter α . Bayesian estimation weighs the two parameters more than the other two estimation methods and is less prone to overflow of the range of estimates.
- In general, the results of Bayesian estimates based on GE loss functions are more consistent with the dataset than those of other methods.

7. Conclusions

A Lindley distribution with multiple models can be applied to various fields, such as product life research and species presence distribution. This paper focused on the examination of the NGL distribution's inherent properties and significant statistical characteristics. It was determined that the proposed distribution possesses favorable properties and that such properties exhibit relatively straightforward numerical expressions. This finding is instrumental in facilitating subsequent research endeavors. Subsequently, this paper actively explored the application of maximum product spacing

estimation, maximum likelihood estimation, and Bayesian estimation to the NGL distribution. The development of a comprehensive simulation plan, incorporating two distinct failure rates, three disparate censored schemes, and four evaluation criteria, was undertaken to investigate the impact of various estimation methods on the proposed distribution point estimation and interval estimation. Finally, it was concluded that the Bayesian estimator under SE loss and GE loss has relatively good estimation effect, and this estimation method can be integrated into the NGL distribution for further exploration in future research. In addition, in order to ensure that the NGL distribution has practical application significance, a real dataset, containing annual survival rates, was used to carry out research. The analysis showed that the proposed distribution has a good fitting effect on the original dataset, and we found that the proposed distribution has a good fit to the censored data of the dataset. The above estimation methods were used to finally conclude that the distribution also has a good fitting effect on the censored data of the real dataset in the Bayesian estimation. In future studies, we will continue to increase the integration of the proposed distribution with real industry scenarios.

Author contributions

Jiajie Shi: Writing-original draft, Method, Software, Writing-review & editing; Haiping Ren: Methods, Writing-original draft, Writing-review & editing. All authors have read and approved the final version of the manuscript for publication.

Use of AI tools declaration

The authors declare they have not used Artificial Intelligence (AI) tools in the creation of this article.

Acknowledgments

This research was funded by National Natural Science Foundation of China, grant number 71661012.

Conflict of interest

The authors declare no conflict of interest.

References

1. D. V. Lindley, Fiducial distributions and Bayes' theorem, *J. R. Stat. Soc. B.*, **1** (1958), 102–107.
2. C. S. Rajitha, A. Akhilnath, Generalization of the Lindley distribution with application to COVID-19 data, *Int. J. Data Sci. Anal.*, 2022. <https://doi.org/10.1007/s41060-022-00369-2>
3. Y. E. Fatehi, D. S. Chhaya, Extended odd Weibull-Lindley distribution, *Aip. Adv.*, **14** (2024), 035317. <https://doi.org/10.1063/5.0192518>
4. S. K. Ashour, M. A. Eltehiwy, Exponentiated power Lindley distribution, *J. Adv. Res.*, **6** (2015), 895–905. <https://doi.org/10.1016/j.jare.2014.08.005>

5. M. Alizadeh, S. F. Bagheri, E. Samani, S. Ghobadi, S. Nadarajah, Exponentiated power Lindley power series class of distributions: Theory and applications, *Commun. Stat. simul. C.*, **47** (2018), 2499–2531. <https://doi.org/10.1080/03610918.2017.1350270>
6. L. Schelin, S. S. Luna, Spatial prediction in the presence of left-censoring, *Comput. Stat. Data Anal.*, **74** (2014), 125–141. <https://doi.org/10.1016/j.csda.2014.01.004>
7. M. Du, J. Sun, Statistical analysis of interval-censored failure time data, *J. Appl. Probab. Stat. (Chinese)*, **37** (2021), 627–654. <https://doi.org/10.3969/j.issn.1001-4268.2021.06.006>
8. S. W. Lagakos, General right censoring and its impact on the analysis of survival data, *Biometrics*, **35** (1979), 139–156. <https://doi.org/10.2307/2529941>
9. O. E. Abo-Kasem, R. Ahmed, A. Ibrahim, Maximum product spacings for the Kumaraswamy distributions based on progressive type-II censoring schemes, *Sci. Rep.*, **13** (2023), 12063. <https://doi.org/10.1038/s41598-023-38594-9>
10. N. Balakrishnan, N. Kannan, C. T. Lin, S. Wu, Inference for the extreme value distribution under progressive type-II censoring, *J. Stat. Comput. Simul.*, **74** (2004), 25–45. <https://doi.org/10.1080/0094965031000105881>
11. J. I. Seo, S. B. Kang, Y. Kim, Robust Bayesian estimation of a bathtub-shaped distribution under progressive type-II censoring, *Commun. Stat. simul. C.*, **46** (2017), 1008–1023. <https://doi.org/10.1080/03610918.2014.988256>
12. R. Alshenawy, A. Al-Alwan, E. M. Almetwally, A. Z. Afify, H. M. Almongy, Progressive type-II censoring schemes of extended Odd Weibull exponential distribution with applications in medicine and engineering, *Mathematics.*, **8** (2020), 1679. <https://doi.org/10.3390/math8101679>
13. R. Shanker, S. N. Sharma, R. Shanker, A two-parameter Lindley distribution for modeling waiting and survival times data, *Appl. Math.*, **4** (2013), 363–368. <https://doi.org/10.4236/am.2013.42056>
14. C. W. Wu, A. Darmawan, N. Y. Wu, A double sampling plan for truncated life tests under two-parameter Lindley distribution, *Ann. Oper. Res.*, **340** (2024), 619–641. <https://doi.org/10.1007/s10479-024-05955-0>
15. C. Tanış, A new Lindley distribution: applications to COVID-19 patients data, *Soft. Comput.*, **28** (2024), 2863–2874. <https://doi.org/10.1007/s00500-023-09339-7>
16. H. K. T. Ng, L. Luo, Y. Hu, F. Duan, Parameter estimation of three-parameter Weibull distribution based on progressively type-II censored samples, *J. Stat. Comput. Simul.*, **82** (2012), 1661–1678. <https://doi.org/10.1080/00949655.2011.591797>
17. R. Kumar Singh, S. Kumar Singh, U. Singh, Maximum product spacings method for the estimation of parameters of generalized inverted exponential distribution under progressive type II censoring, *J. Stat. Manag. Syst.*, **19** (2016), 219–245. <https://doi.org/10.1080/09720510.2015.1023553>
18. N. Alsadat, D. A. Ramadan, E. M. Almetwally, A. H. Tolba, Estimation of some lifetime parameter of the unit half logistic-geometry distribution under progressively type-II censored data, *J. Radiat. Res. Appl. Sci.*, **16** (2023), 100674. <https://doi.org/10.1016/j.jrras.2023.100674>
19. R. Alotaibi, H. Rezk, A. Elshahhat, Analysis of generalized type-II progressively hybrid Lindley-exponential data and its modeling in physics, engineering, and management, *Aip. Adv.*, **14** (2024), 045213. <https://doi.org/10.1063/5.0197082>
20. H. S. Mohammed, A. Elshahhat, R. Alotaibi, Inferences of inverted Gompertz parameters from an adaptive type-II progressive hybrid censoring, *Phys. Scripta*, **98** (2023), 105222. <https://doi.org/10.1088/1402-4896/acf5ad>

21. S. Dey, M. Nassar, D. Kumar, Alpha power transformed inverse Lindley distribution: A distribution with an upside-down bathtub-shaped hazard function, *J. Comput. Appl. Math.*, **348** (2019), 130–145. <https://doi.org/10.1016/j.cam.2018.03.037>
22. W. H. Greene, *Econometric analysis 4th edition*, New Jersey: Prentice-Hall, 2000.
23. H. S. Mohammed, M. Nassar, R. Alotaibi, A. Elshahhat, Analysis of adaptive progressive type-II hybrid censored dagum data with Applications, *Symmetry*, **14** (2022), 2146. <https://doi.org/10.3390/sym14102146>
24. K. Ghosh, S. Jammalamadaka, A general estimation method using spacings, *J. Stat. Plan. Infer.*, **93** (2001), 71–82. [https://doi.org/10.1016/S0378-3758\(00\)00160-9](https://doi.org/10.1016/S0378-3758(00)00160-9)
25. A. A. Mohamed, R. M. Refaey, G. R. AL-Dayian, Bayesian and E-Bayesian estimation for odd generalized exponential inverted Weibull distribution, *J. Bus. Environ. Sci.*, **3** (2024), 275–301.
26. G. R. AL-Dayian, A. EL-Helbawy, F. Abd EL-Maksoud, Bayesian and maximum likelihood estimation for mixture models of the new Topp-Leone-G Family, *Environ. Sci.*, **4** (2025), 210–253.
27. Y. Y. Abdelall, A. S. Hassan, E. M. Almetwally, A new extention of the odd inverse Weibull-G family of distributions: Bayesian and non-Bayesian estimation with engineering applications, *Comput. J. Math. Stat. Sci.*, **3** (2024), 359–388. <https://doi.org/10.21608/cjmss.2024.285399.1050>
28. A. Tolba, Bayesian and non-Bayesian estimation methods for simulating the parameter of the Akshaya distribution, *Comput. J. Math. Stat. Sci.*, **1** (2022), 13–25. <https://doi.org/10.21608/cjmss.2022.270897>
29. M. Nagy, Expected Bayesian estimation based on generalized progressive hybrid censored data for Burr-XII distribution with applications, *Aip. Adv.*, **14** (2024), 015357. <https://doi.org/10.1063/5.0184910>
30. R. Calabria, G. Pulcini, An engineering approach to Bayes estimation for the Weibull distribution, *Microelectron. Reliab.*, **34** (1994), 789–802. [https://doi.org/10.1016/0026-2714\(94\)90004-3](https://doi.org/10.1016/0026-2714(94)90004-3)
31. M. H. Chen, Q. M. Shao, Monte Carlo estimation of Bayesian credible and HPD intervals, *J. Comput. Graph. Stat.*, **8** (1999), 69–92. <https://doi.org/10.1080/10618600.1999.10474802>
32. H. Ren, Q. Gong, X. Hu, Estimation of entropy for generalized Rayleigh distribution under progressively Type-II censored samples, *Axioms*, **12** (2023), 776. <https://doi.org/10.3390/axioms12080776>
33. A. Bekker, J. J. J. Roux, P. J. Mosteit, A generalization of the compound Rayleigh distribution: using a Bayesian method on cancer survival times, *Commun. Stat. Theor. M.*, **29** (2000), 1419–1433. <https://doi.org/10.1080/03610920008832554>
34. K. H. Habib, M.A. Khaleel, H. Al-Mofleh, P. E. Oguntunde, S. J. Adeyeye, Parameters estimation for the $[0, 1]$ truncated Nadarajah Haghighi Rayleigh distribution, *Sci. Afr.*, **23** (2024), e02105. <https://doi.org/10.1016/j.sciaf.2024.e02105>
35. B. Klefsjö, TTT-plotting-a tool for both theoretical and practical problems, *J. Stat. Plan. Infer.*, **29** (1991), 99–110. [https://doi.org/10.1016/0378-3758\(92\)90125-C](https://doi.org/10.1016/0378-3758(92)90125-C)

Appendix

A. The parameters, the reliability function, and hazard function of the NGL distribution are estimated using MPS estimation

```
clear
tic
a0=0.5; b0=0.75; m=30;n=50;num=1000;t=0.5;S0=0.9248;h0=0.1842;ca=0;cb=0;cS=0;ch=0;
Ra=zeros(1,m-1),n-m];
Rb=[floor((n-m)/2),zeros(1,m-2),(n-m-floor((n-m)/2))];
Rc=[n-m,zeros(1,m-1)];
a = zeros(num,1);
b = zeros(num,1);
ACI1 = zeros(num, 2);
ACI2 = zeros(num, 2);
ACI3 = zeros(num, 2);
ACI4 = zeros(num, 2);
for i=1:num
G=rand(1,m);
for j=1:m
H(j)=G(j).^(1/(j+sum(Ra(m-j+1:m)))));
end
for j=1:m
Z(j)=1-prod(H(m-j+1:m));
end
x=(-b0.*lambertw(-1,exp(-1./b0).*(Z-1)./b0)-1)./(a0.*b0);
[a(i,1),b(i,1),ACI1(i,:),ACI2(i,:),S(i,1),h(i,1),ACI3(i,:),ACI4(i,:)] =MPS(t,Ra,n,x,a0,b0);
if a0>=ACI1(i,1)&&a0<=ACI1(i,2)
ca=ca+1;
end
if b0>=ACI2(i,1)&&b0<=ACI2(i,2)
cb=cb+1;
end
if S0>=ACI3(i,1)&&S0<=ACI3(i,2)
cS=cS+1;
end
if h0>=ACI4(i,1)&&h0<=ACI4(i,2)
ch=ch+1;
end
end
end
aMLi=mean(a)
mse_a=sqrt(mean((a-a0).^2))
MRE_a=mean(abs((a-a0)./a0))
ACL_a=mean(ACI1(:,2)-ACI1(:,1))
ca=ca/num
```

```

bMLi=mean(b)
mse_b=sqrt(mean((b-b0).^2))
MRE_b=mean(abs((b-b0)./b0))
ACL_b=mean(ACI2(:,2)-ACI2(:,1))
cb=cb/num
SMLi=mean(S)
mse_S=sqrt(mean((S-S0).^2))
MRE_S=mean(abs((S-S0)./S0))
ACL_S=mean(ACI3(:,2)-ACI3(:,1))
cS=cS/num
hMLi=mean(h)
mse_h=sqrt(mean((h-h0).^2))
MRE_h=mean(abs((h-h0)./h0))
ACL_h=mean(ACI4(:,2)-ACI4(:,1))
ch=ch/num
toc

```

B. The parameters, the reliability function, and hazard function of the NGL distribution are estimated using ML estimation

```

clear
tic
a0=0.5; b0=0.75; m=30;n=50;num=1000;t=0.5;S0=0.9248;h0=0.1842;ca=0;cb=0;cS=0;ch=0;
Ra=zeros(1,m-1),n-m];
Rb=[floor((n-m)/2),zeros(1,m-2),(n-m-floor((n-m)/2))];
Rc=[n-m,zeros(1,m-1)];
a = zeros(num,1);
b = zeros(num,1);
S=zeros(num,1);
h=zeros(num,1);
ACI1 = zeros(num, 2);
ACI2 = zeros(num, 2);
ACI3 = zeros(num, 2);
ACI4 = zeros(num, 2);
for i=1:num
G=rand(1,m);
for j=1:m
H(j)=G(j).^(1/(j+sum(Ra(m-j+1:m)))));
end
for j=1:m
Z(j)=1-prod(H(m-j+1:m));
end
x=(-b0.*lambertw(-1,exp(-1./b0).*(Z-1)./b0)-1)./(a0.*b0);
[a(i,1),b(i,1),ACI1(i,:),ACI2(i,:),S(i,1),h(i,1),ACI3(i,:),ACI4(i,:)] =MLi(t,Ra,n,x,a0,b0);
if a0>=ACI1(i,1)&&a0<=ACI1(i,2)
ca=ca+1;

```

```

end
if b0>=ACI2(i,1)&&b0<=ACI2(i,2)
cb=cb+1;
end
if S0>=ACI3(i,1)&&S0<=ACI3(i,2)
cS=cS+1;
end
if h0>=ACI4(i,1)&&h0<=ACI4(i,2)
ch=ch+1;
end
end
end
aMLi=mean(a)
mse_a=sqrt(mean((a-a0).^2))
MRE_a=mean(abs((a-a0)./a0))
ACL_a=mean(ACI1(:,2)-ACI1(:,1))
ca=ca/num
bMLi=mean(b)
mse_b=sqrt(mean((b-b0).^2))
MRE_b=mean(abs((b-b0)./b0))
ACL_b=mean(ACI2(:,2)-ACI2(:,1))
cb=cb/num
SMLi=mean(S)
mse_S=sqrt(mean((S-S0).^2))
MRE_S=mean(abs((S-S0)./S0))
ACL_S=mean(ACI3(:,2)-ACI3(:,1))
cS=cS/num
hMLi=mean(h)
mse_h=sqrt(mean((h-h0).^2))
MRE_h=mean(abs((h-h0)./h0))
ACL_h=mean(ACI4(:,2)-ACI4(:,1))
ch=ch/num
toc

```

C. The parameters, the reliability function, and hazard function of the NGL distribution are estimated using Bayesian estimation

```

clear
a0=0.5; b0=0.75; m=30;n=50;num=1000;t=0.5;S0=0.9248;h0=0.1842;c1=0;c2=0;c3=0;c4=0;
Ra=[zeros(1,m-1),n-m];
Rb=[floor((n-m)/2),zeros(1,m-2),(n-m-floor((n-m)/2))];
Rc=[n-m,zeros(1,m-1)];
a_BMH=zeros(num,1);
b_BMH=zeros(num,1);
S_BMH=zeros(num,1);
h_BMH=zeros(num,1);
aHPD=zeros(num,1);

```

```

bHPD=zeros(num,1);
SHPD=zeros(num,1);
hHPD=zeros(num,1);
for i=1:num
G=rand(1,m);
for j=1:m
H(j)=G(j).^(1/(j+sum(Rc(m-j+1:m))));
end
for j=1:m
Z(j)=1-prod(H(m-j+1:m));
end
x=(-b0.*lambertw(-1,exp(-1./b0).*(Z-1)./b0)-1)./(a0.*b0);
[a_BMH(i,1),b_BMH(i,1),S_BMH(i,1),h_BMH(i,1),aHPD(i,1),bHPD(i,1),SHPD(i,1),hHPD(i,1),cp1
,cp2,cp3,cp4]=BMH(t,Rc,n,x,0.4,0.7);
c1=c1+cp1;c2=c2+cp2;c3=c3+cp3;c4=c4+cp4;
end
aBMH=mean(a_BMH)
mse_a=sqrt(mean((a_BMH-a0).^2))
MRE_a=mean(abs((a_BMH-a0)./a0))
ACL_a=mean(aHPD)
ca=c1/num
bBMH=mean(b_BMH)
mse_b=sqrt(mean((b_BMH-b0).^2))
MRE_b=mean(abs((b_BMH-b0)./b0))
ACL_b=mean(bHPD)
cb=c2/num
SBMH=mean(S_BMH)
mse_S=sqrt(mean((S_BMH-S0).^2))
MRE_S=mean(abs((S_BMH-S0)./S0))
ACL_S=mean(SHPD)
cS=c3/num
hBMH=mean(h_BMH)
mse_h=sqrt(mean((h_BMH-h0).^2))
MRE_h=mean(abs((h_BMH-h0)./h0))
ACL_h=mean(hHPD)
ch=c4/num
Counter diagram of the parameters of the NGL distribution
clear
syms a b;
a0=0.5; b0=0.75; m=50;n=100;
Ra=[zeros(1,m-1),n-m];
Rb=[floor((n-m)/2),zeros(1,m-2),(n-m-floor((n-m)/2))];
Rc=[n-m,zeros(1,m-1)];
G=rand(1,m);

```

```

for j=1:m
H(j)=G(j).^(1/(j+sum(Ra(m-j+1:m)))));
end
for j=1:m
Z(j)=1-prod(H(m-j+1:m));
end
P=Ra;
x=(-b0.*lambertw(-1,exp(-1./b0).*(Z-1)./b0)-1)./(a0.*b0);
lb=m.*log(a)-a.*sum(x.*(1+P))+sum(log(1-b+a.*b.*x))+sum(P.*log(1+a.*b.*x));
lc= n.*log(a)-a.*sum(x)+sum(log(1-b+a.*b.*x));
laa=matlabFunction(lb, 'Vars', {a, b});
a1 = 0.2:0.01:1;
b1 = 0.36:0.01:0.75;
[a1,b1] = meshgrid(a1,b1);
la=real(laa(b1,a1));
[M,c]=contour(a1,b1,la,'showText','on');
xlabel('θ');
ylabel('λ');
c.LineWidth=3;
colorbar.

```



AIMS Press

© 2025 the Author(s), licensee AIMS Press. This is an open access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0>)