

Research Paper

Estimation and prediction for the XLindley distribution under progressively Type-II censoring

ROSHANAK ZAMAN*, ADELEH FALLAH
DEPARTMENT OF STATISTICS, PAYAME NOOR UNIVERSITY, TEHRAN, IRAN

Received: May 09, 2026/ Revised: August 24, 2026/ Accepted: September 23, 2026

Abstract: In this paper, we consider the problems of estimating unknown parameters as well as predicting the failure times of the removed units in multiple stages of the progressively censored sample coming from an XLindley distribution. We obtain maximum likelihood, moment-based, bootstrap, and Bayes estimates under the squared error loss function. Since the integrals associated with the Bayesian estimators do not admit closed-form solutions, Lindley's approximation, the Markov chain Monte Carlo method, and Tierney and Kadane's approximation are used to approximate these integrals. We also compute the confidence intervals based on the asymptotic distribution of the maximum likelihood estimate, the confidence intervals based on the asymptotic distribution of the logarithm of the maximum likelihood estimate, the exact confidence intervals, the confidence intervals based on the likelihood ratio statistic, the bootstrap confidence intervals, and the Bayesian confidence intervals. Different point and interval predictors are derived based on classical and Bayesian approaches. A real example is provided to further explain the methods introduced. A Monte Carlo simulation study is conducted to evaluate and compare the performance of different estimation and prediction methods.

Keywords: Bootstrap; Lindley approximation; Markov chain monte carlo; Progressively Type-II censoring; Tierney and Kadane's approximation.

Mathematics Subject Classification (2010): 62N01, 62N05, 62F10.

1 Introduction

In many studies related to lifetime and reliability, we encounter situations where test units are removed from the study before failure occurs. Such data are called censored

*Corresponding author: r.zaman@pnu.ac.ir

data. In this paper, we consider progressively Type-II censoring. Suppose n units are placed on a lifetime test. Before the test begins, m (where $m < n$) and $R_1, \dots, R_m$ are fixed in advance such that $\sum_{i=1}^m R_i = n - m$. Upon the first failure, R_1 of the remaining $n - 1$ surviving units are randomly removed from the test. Similarly, at the time of the second observed failure, R_2 of the remaining units are randomly removed. This process continues until after the m -th failure, all remaining surviving units, i.e., $R_m = n - m - R_1 - \dots - R_{m-1}$, are removed. The resulting ordered failure times are denoted by $X_{1:m:n} < X_{2:m:n} < \dots < X_{m:m:n}$, and are called a progressively Type-II censored sample. For simplicity, we denote this sample by $X_1 < X_2 < \dots < X_m$ with censoring scheme $(R_1, \dots, R_m)$. In the special case where $R_1 = \dots = R_{m-1} = 0$ and $R_m = n - m$, this scheme reduces to conventional Type-II censoring. If $R_1 = \dots = R_m = 0$ and $n = m$, we obtain a complete sample. For further details, see Balakrishnan and Aggarwala (2000), Balakrishnan and Cramer (2014).

The XLindley distribution was introduced by Chouia and Zeghdoudi (2021). The XLindley has the cumulative distribution function (cdf)

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

and the probability density function (pdf)

$$f(x; \theta) = \frac{\theta^2}{(1 + \theta)^2} (2 + \theta + x) e^{-\theta x}, \quad x \geq 0, \theta > 0. \quad (2)$$

Chouia and Zeghdoudi (2021) studied the statistical properties like stochastic ordering, quantile function, the maximum likelihood estimator (MLE) and method of moments. Alotaibi et al. (2022) address the estimation issues of the XLindley distribution using an adaptive Type-II progressive hybrid censoring scheme. The MLE and Bayesian approaches are used to estimate the unknown parameter, reliability, and hazard rate functions. The approximate confidence intervals (ACIs) and the highest posterior density (HPD) credible intervals are also computed. Constant-stress accelerated life tests combined with progressive censoring were later addressed by Alotaibi et al. (2023), and a partially accelerated variant was studied by Nassar et al. (2023), who derived frequentist and Bayesian estimators for the acceleration factor together with the corresponding reliability functions. Estimation from record-breaking data has also received attention: Zanjiran and MirMostafaei (2024) constructed point and interval estimators from lower record values and inter-record times, showing that incorporating the timing of records improves estimation efficiency. Elbatal et al. (2025) proposed a joint progressively Type-II censoring framework based on the XLindley distribution to evaluate the reliability of two production lines operating individually and jointly. They derived maximum likelihood and Bayesian estimators for the model parameters and reliability functions, along with corresponding asymptotic, bootstrap, and HPD interval estimates. In a related direction, Thongjub and Panichkitkosolkul (2026) compared several confidence-interval procedures and reported that likelihood-based and Wald-type intervals tend to perform consistently across sample sizes. Elshahhat and Alotaibi (2026) developed a comprehensive reliability analysis under a generalized progressive mixed censoring scheme. Probability-based estimation procedures for model

parameters and key reliability features were derived. A Bayesian inference model using Markov chain Monte Carlo (MCMC) techniques and HPD credible intervals was presented.

Although the XLindley distribution has gained attention in recent studies - for instance, Alotaibi et al. (2024) explored it under adaptive hybrid censoring and Elshahhat and Alotaibi (2026) investigated generalized schemes- these works primarily focus on large-sample asymptotic properties or MCMC-based Bayesian inference. A notable gap remains regarding exact inference in small-sample settings under standard progressive Type-II censoring. The present study addresses this by deriving an exact pivotal quantity, enabling the construction of confidence intervals that maintain nominal coverage even when sample sizes are limited. Furthermore, we extend the inferential framework to include Bayesian prediction of failure times for removed units, a practical aspect often overlooked in standard estimation papers.

The aim of the present paper is two fold. First we consider the problem of estimation of parameters under classical and Bayesian approaches. Maximum likelihood, parametric bootstrap, moment-based and Bayesian estimations for the parameter are calculated based on the symmetric squared error loss (SEL) function. Since the integrals associated with the Bayesian estimators do not admit closed-form solutions, MCMC method, Tierney and Kadane's approximation (TK) and Lindley approximation method are used to approximate these integrals. We also compute ACIs, exact confidence intervals (ECIs), confidence intervals using the likelihood ratio statistic, bootstrap confidence intervals, and Bayesian confidence intervals under the respective approaches. Second, we consider the problem of prediction of the future unobserved failure times using maximum likelihood predictor (MLP), best unbiased predictor (BUP), conditional median predictor (CMP) and Bayesian predictor (BP). We also developed the prediction intervals for the future unobserved failure times using pivotal quantity method, highest conditional density (HCD) method and Bayesian method.

This paper is organized as follows. In Section 2 we provide the MLE and moments based estimate, parametric bootstrap estimate (PBE) and the Bayes estimates of the XLindley model parameter. Different confidence intervals are presented in Section 3. In Section 4, we address prediction of future failures times under classical and Bayesian approaches. In Section 5, various prediction intervals are computed. Analysis of a real data set is given in Section 6 for illustrative purposes. We also compare different proposed methods using Monte Carlo simulation in Section 7. Finally, discussion and conclusion are given in Section 8.

2 Estimation of the model parameter

In this section, we consider the estimation of the unknown parameter θ .

2.1 Maximum likelihood estimation

Let $X_1, \dots, X_m$ be a progressively Type-II censored sample from a distribution with pdf $f(\cdot)$ and cdf $F(\cdot)$, where the censoring scheme is $(R_1, \dots, R_m)$. The likelihood

function is then given by

$$L(\mathbf{x}, \theta) = C \prod_{i=1}^m f(x_i, \theta) [1 - F(x_i, \theta)]^{R_i}, \quad (3)$$

where $C = n(n - R_1 - 1)(n - R_1 - R_2 - 2) \dots (n - \sum_{i=1}^{m-1} (R_i + 1))$, where $\mathbf{x} = (x_1, \dots, x_m)$ is the observed progressively Type-II censored samples (Balakrishnan and Cramer, 2014). Using (2), (1) and (3) the likelihood function becomes

$$L(\mathbf{x}, \theta) = C \frac{\theta^{2m}}{(1 + \theta)^{2m}} e^{-\theta [\sum_{i=1}^m x_i(1 + R_i)]} \prod_{i=1}^m (2 + \theta + x_i) \prod_{i=1}^m \frac{[\theta x_i + (1 + \theta)^2]^{R_i}}{(1 + \theta)^{2R_i}}. \quad (4)$$

By taking the logarithm of the likelihood function and differentiating with respect to θ , the following equation is obtained:

$$\begin{aligned} l(x; \theta) &= \log L(X, \theta) = \log C + 2m \log \theta - 2n \log(1 + \theta) - \theta \sum_{i=1}^m x_i(1 + R_i) \\ &\quad + \sum_{i=1}^m \log(2 + \theta + x_i) + \sum_{i=1}^m R_i \log [\theta x_i + (1 + \theta)^2], \\ \frac{\partial l(x; \theta)}{\partial \theta} &= \frac{2m + 2\theta(m - n)}{\theta(1 + \theta)} + \sum_{i=1}^m \frac{1}{2 + \theta + x_i} - \sum_{i=1}^m x_i(1 + R_i) \\ &\quad + \sum_{i=1}^m R_i \left(\frac{x_i + 2(\theta + 1)}{\theta x_i + (1 + \theta)^2} \right). \end{aligned} \quad (5)$$

By solving (5), using numerical methods such as Newton-Raphson and fixed-point iteration, we can calculate the MLE of the parameter θ , i.e. $\hat{\theta}_M$. In the fixed-point iteration method, $g(\theta)$ is defined as

$$g(\theta) = 2m \left[\frac{2n}{1 + \theta} + \sum_{i=1}^m x_i(1 + R_i) - \sum_{i=1}^m \frac{1}{2 + \theta + x_i} - \sum_{i=1}^m R_i \left(\frac{2(1 + \theta) + x_i}{(1 + \theta)^2 + \theta x_i} \right) \right]^{-1}.$$

Next, we show the uniqueness and existence of the ML estimate of θ . To this end, let $\nu_1(\theta) = g(\theta)$ and $\nu_2(\theta) = \theta$. It can be easily verified that $\nu_1(\theta)$ is an increasing function with

$$\lim_{\theta \rightarrow 0} \nu_1(\theta) = \frac{2m}{2m + \sum_{i=1}^m x_i - \sum_{i=1}^m \frac{1}{2 + x_i}}, \quad \lim_{\theta \rightarrow \infty} \nu_1(\theta) = \frac{2m}{\sum_{i=1}^m x_i(1 + R_i)}.$$

So $\nu_1(\theta)$ starts from a positive real value at 0 and increases to $\frac{2m}{\sum_{i=1}^m x_i(1 + R_i)}$, which is a finite value. For large θ , $\nu_1(\theta)$ is a finite value whereas $\nu_2(\theta) \rightarrow \infty$ as $\theta \rightarrow +\infty$. This implies that there exists one real positive root, say $\hat{\theta}$, such that $g(\hat{\theta}) = \hat{\theta}$. The second derivative of the log-likelihood function is given by

$$\begin{aligned} \frac{\partial^2 \log L(\theta)}{\partial \theta^2} &= \frac{-2m}{\theta^2} + \frac{2n}{(1 + \theta)^2} - \sum_{i=1}^m \frac{1}{(2 + \theta + x_i)^2} \\ &\quad - \sum_{i=1}^m R_i \left(\frac{x_i^2 + 2x_i(2 + \theta) + 2(1 + \theta)^2}{[\theta x_i + (1 + \theta)^2]^2} \right). \end{aligned}$$

2.2 Parametric bootstrap estimation

The bootstrap is a simple and powerful method for approximating the distribution of statistics, particularly useful for estimating standard errors and confidence intervals. This method was introduced by Efron (1979). Parametric bootstrap essentially means resampling with replacement. In this method, first, numerous samples (resamples) of the same size as the original dataset are drawn with replacement. Then, the desired statistic is calculated for each resample. Finally, the distribution of these calculated statistics is analyzed. The following algorithm is employed to obtain the PBEs of θ of the XLindley distribution:

Algorithm 2.1. *Bootstrap Algorithm:*

- i. The MLE of parameter θ , $\hat{\theta}_M$ is calculated.
- ii. Using $\hat{\theta}_M$ and the original censoring scheme $R_1, \dots, R_m$, generate a progressively Type-II censored bootstrap sample $x_1^*, \dots, x_m^*$ from the fitted XLindley distribution with parameter $\hat{\theta}_M$.
- iii. Compute the MLE of θ based on the generated bootstrap sample denoted by $\hat{\theta}^*$.
- iv. Repeat step ii, B times to obtain $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$.
- v. Using the bootstrap estimates $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$, the PBE of the parameter θ , $\hat{\theta}_{BE}$, is $\hat{\theta}_{PBE}^* = \frac{1}{B} \sum_{i=1}^B \hat{\theta}_i^*$.

2.3 Moment-based estimation

Suppose $X_1, \dots, X_m$ is a progressively Type-II censored sample with distribution function $F(x; \theta)$ in (1). The random variables $Y_{i:m:n}$, ($Y_{1:m:n} < Y_{2:m:n} < \dots < Y_{m:m:n}$) with

$$Y_{i:m:n} = -\log[1 - F(X_i; \theta)] = -\log \left[\left(1 + \frac{\theta X_i}{(1 + \theta)^2} \right) e^{-\theta X_i} \right], \quad i = 1, \dots, m,$$

are a progressively Type-II censored sample from a standard exponential distribution. This follows from the probability integral transform: since $F(X_i; \theta)$ is uniformly distributed on $(0, 1)$ for every θ , the quantity $Y_{i:m:n} = -\log[1 - F(X_i; \theta)]$ has the same distributional structure as a progressively Type-II censored order statistic from the standard exponential distribution with the same censoring scheme $(R_1, \dots, R_m)$. By Theorem 2.1 of Balakrishnan and Aggarwala (2000), the spacings between successive progressively censored order statistics from the exponential distribution, appropriately normalized, are independent and identically distributed as standard exponential random variables. Under the following transformations, it is shown that $Z_1, \dots, Z_m$ are independent and identically distributed random variables with the standard exponential distribution.

$$\begin{cases} Z_1 = n \cdot Y_{1:m:n}, \\ Z_2 = (n - R_1 - 1) \cdot (Y_{2:m:n} - Y_{1:m:n}), \\ \vdots \\ Z_m = (n - R_1 - \dots - R_{m-1} - m + 1) \cdot (Y_{m:m:n} - Y_{m-1:m:n}). \end{cases}$$

Consider the transformed variables

$$V = 2Z_1 = 2nY_{1:m:n}, \quad \text{and} \quad U = 2 \sum_{i=2}^m Z_i = 2 \left[\sum_{i=1}^m (R_i + 1)Y_{i:m:n} - nY_{1:m:n} \right].$$

These quantities are independently distributed, with V having a chi-squared distribution with 2 degrees of freedom and U having a chi-squared distribution with $2m - 2$ degrees of freedom. To obtain the moment-based estimator (MBE) of θ , denoted as by $\hat{\theta}_{MBE}$, let us define $Q(\theta) = U + V = 2 \sum_{i=1}^m (R_i + 1)Y_{i:m:n}$. It can be shown that

$$Q(\theta) = 2 \sum_{i=1}^m (R_i + 1)Y_{i:m:n} = 2 \sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right]. \quad (6)$$

Since Z_1 is independent of $Z_2, \dots, Z_m$ (each being one of the m independent standard exponential spacings defined above), $V = 2Z_1$ and $U = 2 \sum_{i=2}^m Z_i$ are independent. Moreover, since $2Z_1$ is twice a standard exponential variable, $V \sim \chi^2(2)$; and since U is twice the sum of $(m-1)$ independent standard exponential variables, $U \sim \chi^2(2(m-1))$. By the additivity property of independent chi-squared random variables, $Q(\theta) = U + V$ therefore has a chi-squared distribution with $2 + (2m - 2) = 2m$ degrees of freedom, for every value of θ . Because this distribution does not depend on θ , $Q(\theta)$ is a pivotal quantity for θ , a property that will be used in Section 2.2 to construct an ECI.

Using Markov's inequality, for any $\epsilon > 0$ and $m \rightarrow \infty$, we have

$$\begin{aligned} Pr\left(\left|\frac{Q(\theta)}{2m} - 1\right| > \epsilon\right) &< \frac{E\left(\frac{Q(\theta)}{2m} - 1\right)^2}{\epsilon^2} \\ &= \frac{Var\left(\frac{Q(\theta)}{2m}\right) + E^2\left(\frac{Q(\theta)}{2m} - 1\right)}{\epsilon^2} = \frac{1}{\epsilon^2 m} \rightarrow 0. \end{aligned}$$

Therefore, when $m \rightarrow \infty$, the probability converges to zero. In other words,

$$\frac{Q(\theta)}{2m} = \frac{2 \sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right]}{2m} \xrightarrow{p} 1.$$

Then, we can obtain the MBE estimate of θ as the unique solution of the equation $Q(\theta) = 2m$, or the equation

$$\sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right] - m = 0.$$

Note that

$$\begin{aligned} \frac{\partial Q(\theta)}{\partial \theta} &= \frac{\partial}{\partial \theta} \left\{ 2 \sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right] \right\} \\ &= 2 \sum_{i=1}^m (R_i + 1) X_i \left(\frac{(\theta - 1)}{(\theta + 1)[(1 + \theta)^2 + \theta X_i]} + 1 \right) \geq 0, \end{aligned}$$

$\lim_{\theta \rightarrow 0} Q(\theta) = 0$, and $\lim_{\theta \rightarrow +\infty} Q(\theta) = +\infty$. Hence, $Q(\theta)$ is a continuous and strictly increasing function on the interval $(0, \infty)$. Therefore, the existence and uniqueness of the MBE of θ , which is a solution to $Q(\theta) = 2m$, is guaranteed.

2.4 Bayesian estimation

In this section, we obtain the Bayes estimator (BE) of the parameter θ using the Bayesian approach. We assume that θ follows a Gamma prior distribution with hyperparameters α and β , denoted as $\theta \sim \text{Gamma}(\alpha, \beta)$. The prior pdf is given by

$$\pi(\theta) = \frac{\beta^\alpha}{\Gamma(\alpha)} \theta^{\alpha-1} e^{-\beta\theta}, \quad \theta > 0, \quad \alpha, \beta > 0. \quad (7)$$

Combining the prior distribution with the likelihood function $L(\theta|\mathbf{x})$, the posterior pdf $\pi(\theta|\mathbf{x})$ is obtained via Bayes' theorem

$$\pi(\theta|\mathbf{x}) = \frac{L(\theta|\mathbf{x})\pi(\theta)}{\int_0^\infty L(\theta|\mathbf{x})\pi(\theta) d\theta}.$$

Based on the likelihood function specified in (4) and the prior distribution given in (7), the posterior pdf for θ is derived as follows.

$$\pi(\theta|\mathbf{x}) = \frac{\theta^{2m+\alpha-1}}{(1+\theta)^{2m} A(\mathbf{x})} e^{-\theta[\sum_{i=1}^m x_i(1+R_i)+\beta]} \prod_{i=1}^m (2+\theta+x_i) \prod_{i=1}^m \left(\frac{(1+\theta)^2 + \theta x_i}{(1+\theta)^2} \right)^{R_i}, \quad (8)$$

where $A(\mathbf{x})$ is given by

$$A(\mathbf{x}) = \int_0^\infty \frac{\theta^{2m+\alpha-1}}{(1+\theta)^{2m}} e^{-\theta[\sum_{i=1}^m x_i(1+R_i)+\beta]} \prod_{i=1}^m (2+\theta+x_i) \prod_{i=1}^m \left(\frac{(1+\theta)^2 + \theta x_i}{(1+\theta)^2} \right)^{R_i} d\theta.$$

To derive the BE, the SEL function is employed that is defined as $L(\theta, \hat{\theta}) = (\theta - \hat{\theta})^2$. Under the SEL, the BE is the mean of the posterior distribution. Therefore, the BE of θ , denoted by $\hat{\theta}_{BE}$, is calculated as

$$\hat{\theta}_{BE} = E[\theta|\mathbf{x}] = \int_0^\infty \theta \pi(\theta|\mathbf{x}) d\theta.$$

Considering the posterior pdf given in (8), the BE of θ under the SEL is then obtained as follows.

$$\hat{\theta}_{BE} = \int_0^\infty \frac{\theta^{2m+\alpha}}{(1+\theta)^{2m} A(\mathbf{x})} e^{-\theta[\sum_{i=1}^m x_i(1+R_i)+\beta]} \prod_{i=1}^m (2+\theta+x_i) \prod_{i=1}^m \left(\frac{(1+\theta)^2 + \theta x_i}{(1+\theta)^2} \right)^{R_i} d\theta. \quad (9)$$

The integral in (9) does not have a closed-form solution. Therefore, numerical methods are required to compute the Bayes estimates. Two approaches are considered here: the MCMC technique, TK Approximation and Lindley's approximation.

2.4.1 Bayesian estimation via the Metropolis-Hastings algorithm

For the MCMC implementation, the Metropolis-Hastings algorithm is used. A Normal proposal distribution is assumed to draw samples from the posterior pdf (Hastings, 1970; Metropolis et al., 1953; Robert and Casella, 2004).

Algorithm 2.2. *The Metropolis-Hastings algorithm:*

- i. An initial value θ_0 is chosen for the parameter θ . The MLE is a natural choice for starting the chain.
- ii. Set the iteration counter $q = 1$.
- iii. A candidate value θ^* is drawn from the proposal distribution $N(\theta_{q-1}, S_\theta^2)$, where S_θ^2 is taken as the inverse of the Fisher information evaluated at the current estimate.
- iv. The acceptance probability is computed as

$$P = \min \left\{ \frac{\pi(\theta^*|\mathbf{x})f(\theta_{q-1}|\theta^*)}{\pi(\theta_{q-1}|\mathbf{x})f(\theta^*|\theta_{q-1})}, 1 \right\},$$

where $f(x|b)$ is the pdf of the normal distribution with parameters b and S_θ^2 , otherwise set $\theta_q = \theta_{q-1}$.

v. Set $q = q + 1$.

vi. Repeat Steps iii and iv, N times to arrive at the MCMC sample $\theta_1, \dots, \theta_N$.

vii. Based on an MCMC sample $\theta_{M+1}, \dots, \theta_N$, the Bayes estimate of θ can be obtained as $\hat{\theta}_{BES} = \frac{1}{N-M} \sum_{q=M+1}^N \theta_q$, where M denotes the number of burn-in iterations.

2.4.2 Tierney and Kadane's Approximation

For a one-parameter model, Tierney and Kadane (1986) proposed a technique to approximate Bayes estimates. Let

$$\psi_0(\theta) = \frac{1}{m} \log \pi(\theta|\mathbf{x}), \quad \psi^*(\theta) = \psi_0(\theta) + \frac{1}{m} \log \xi(\theta).$$

Then, according to the TK method, the approximated Bayes estimate of θ , denoted as by $\hat{\theta}_{BETK}$, is

$$\hat{\theta}_{BETK} = \sqrt{\frac{\sigma^*}{\sigma_0}} \exp \left\{ m [\psi^*(\theta^*) - \psi_0(\theta_0)] \right\},$$

where θ^* and θ_0 maximize $\psi^*(\theta)$ and $\psi_0(\theta)$, respectively, and σ^* and σ_0 are minus the inverse of the second derivatives of $\psi^*(\theta)$ and $\psi_0(\theta)$ at the points θ^* and θ_0 , respectively. We have

$$\begin{aligned} \psi_0(\theta) &= \frac{1}{m} \left[(2m + \alpha - 1) \log \theta - 2n \log(1 + \theta) - \theta \left[\sum_{i=1}^m x_i(1 + R_i) + \beta \right] + \log C \right. \\ &\quad \left. - \log A(\mathbf{x}) + \sum_{i=1}^m \log(2 + \theta + x_i) + \sum_{i=1}^m R_i \log [\theta x_i + (1 + \theta)^2] \right]. \end{aligned}$$

Therefore, θ_0 can be derived from the following equation

$$\begin{aligned} \frac{\partial \psi_0(\theta)}{\partial \theta} &= \frac{1}{m} \left[\frac{2m + \alpha - 1}{\theta} - \frac{2n}{1 + \theta} - \sum_{i=1}^m x_i(1 + R_i) - \beta + \sum_{i=1}^m \frac{1}{2 + \theta + x_i} \right. \\ &\quad \left. + \sum_{i=1}^m R_i \left(\frac{x_i + 2(\theta + 1)}{\theta x_i + (1 + \theta)^2} \right) \right]. \end{aligned}$$

The second order derivative of $\psi_0(\theta)$ at θ_0 , namely

$$\begin{aligned} \frac{\partial^2 \psi_0(\theta)}{\partial \theta^2} &= \frac{1}{m} \left[\frac{-(2m + \alpha - 1)}{\theta^2} + \frac{2n}{(1 + \theta)^2} - \sum_{i=1}^m \frac{1}{(2 + \theta + x_i)^2} \right. \\ &\quad \left. - \sum_{i=1}^m R_i \left(\frac{x_i^2 + 2x_i(2 + \theta) + 2(1 + \theta)^2}{[\theta x_i + (\theta + 1)^2]^2} \right) \right], \end{aligned}$$

where $\sigma_0 = [-\frac{\partial^2 \psi_0(\theta)}{\partial \theta^2}]_{\theta=\theta_0}^{-1}$. We derive the approximate Bayes estimate of θ under the SEL function. Let $\xi(\theta) = \theta$ and θ^* be the maximum point of the following quantity

$$\begin{aligned} \psi^*(\theta) &= \frac{1}{m} \left[(2m + \alpha) \log \theta - 2n \log(1 + \theta) - \theta \left[\sum_{i=1}^m x_i(1 + R_i) + \beta \right] + \log C - \log A(\mathbf{x}) \right. \\ &\quad \left. + \sum_{i=1}^m \log(2 + \theta + x_i) + \sum_{i=1}^m R_i \log [\theta x_i + (1 + \theta)^2] \right]. \end{aligned}$$

Therefore, θ^* can be derived from the following equation

$$\begin{aligned} \frac{\partial \psi^*(\theta)}{\partial \theta} &= \frac{1}{m} \left[\frac{2m + \alpha}{\theta} - \frac{2n}{1 + \theta} - \sum_{i=1}^m x_i(1 + R_i) - \beta + \sum_{i=1}^m \frac{1}{2 + \theta + x_i} \right. \\ &\quad \left. + \sum_{i=1}^m R_i \left(\frac{x_i + 2(\theta + 1)}{\theta x_i + (1 + \theta)^2} \right) \right]. \end{aligned}$$

The second order derivative of $\psi^*(\theta)$ at θ^* , namely

$$\begin{aligned} \frac{\partial^2 \psi^*(\theta)}{\partial \theta^2} &= \frac{1}{m} \left[\frac{-(2m + \alpha)}{\theta^2} + \frac{2n}{(1 + \theta)^2} - \sum_{i=1}^m \frac{1}{(2 + \theta + x_i)^2} \right. \\ &\quad \left. - \sum_{i=1}^m R_i \left(\frac{x_i^2 + 2x_i(2 + \theta) + 2(1 + \theta)^2}{[\theta x_i + (\theta + 1)^2]^2} \right) \right], \end{aligned}$$

where $\sigma^* = [-\frac{\partial^2 \psi^*(\theta)}{\partial \theta^2}]_{\theta=\theta^*}^{-1}$.

2.4.3 Bayesian estimation using Lindley's approximation

Lindley proposed a method for approximating the ratio of two integrals, which has been employed here to obtain approximate Bayes estimates Lindley (1980). For a one-parameter case, Lindley's approximation for a function of the parameter, say $U(\theta)$, is expressed as

$$E(U(\theta)|\mathbf{x}) = U(\theta) + \frac{1}{2} [U_2 + 2U_1\rho_1] [\sigma(\theta)]^2 + \frac{1}{2} \ell_3 U_1 [\sigma(\theta)]^4, \quad (10)$$

where $U(\theta) = \theta$, $U_j = \frac{\partial^j U(\theta)}{\partial \theta^j}$, for $j = 1, 2$, $\ell_3 = \frac{\partial^3 \log L(\mathbf{x}; \theta)}{\partial \theta^3}$ are defined. Also, $\rho_1 = \frac{\partial}{\partial \theta} \log \pi(\theta)$ and $\sigma(\theta)$ is the inverse of the Fisher information. Based on (10),

the approximate Bayes estimate of θ using Lindley's method, denoted as by $\hat{\theta}_{BEL}$, is obtained as

$$\hat{\theta}_{BEL} = \hat{\theta} + \left(\frac{\alpha - 1}{\hat{\theta}} - \beta \right) \sigma^2(\hat{\theta}) + \frac{1}{2} \ell_3(\hat{\theta}) \sigma^4(\hat{\theta}),$$

where $\hat{\theta}$ is the MLE, $\sigma(\hat{\theta})$ is the inverse of the Fisher information evaluated at the MLE, $\rho_1 = \frac{\partial}{\partial \theta} \log \pi(\theta) = \frac{\alpha-1}{\theta} - \beta$, and

$$\begin{aligned} \ell_3 = \frac{\partial^3 \log L(\mathbf{x}; \theta)}{\partial \theta^3} &= \frac{4m}{\theta^3} - \frac{4n}{(1+\theta)^3} + 2 \sum_{i=1}^m \frac{1}{(2+\theta+x_i)^3} \\ &\quad - 6 \sum_{i=1}^m R_i \left(\frac{x_i + 2(\theta+1)}{[\theta x_i + (\theta+1)^2]^2} \right) + 2 \sum_{i=1}^m R_i \left[\frac{x_i + 2(\theta+1)}{\theta x_i + (\theta+1)^2} \right]^3. \end{aligned}$$

3 Confidence intervals based on different methods

In this section, we compute ACIs, ECIs, confidence intervals using the likelihood ratio statistic, bootstrap confidence intervals, and Bayesian confidence intervals

3.1 Approximate confidence intervals

The ACIs are obtained based on the asymptotic normality of the MLEs. Under standard regularity conditions, the MLE $\hat{\theta}_M$ is asymptotically normally distributed with mean θ and variance $\widehat{Var}(\hat{\theta}_M) = \left[-\frac{\partial^2 \log L(\theta)}{\partial \theta^2} \right]_{\theta=\hat{\theta}_M}^{-1}$. Therefore, the approximate distribution of the pivotal quantity $(\hat{\theta}_M - \theta) / \sqrt{\widehat{Var}(\hat{\theta}_M)}$ is a standard normal. Hence, an asymptotic $100(1-\gamma)\%$ confidence interval for θ can be constructed as

$$(\hat{\theta}_l^{ACI}, \hat{\theta}_u^{ACI}) = \left(\hat{\theta}_M - z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)}, \hat{\theta}_M + z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)} \right), \quad (11)$$

where $z_{1-\gamma/2}$ denotes the $(1-\gamma/2)$ -quantile of the standard normal distribution. Since θ is a positive parameter and the lower end of the confidence interval in (11) could be less than zero, a modified confidence interval of θ can be proposed as

$$(\hat{\theta}_{l+}^{ACI}, \hat{\theta}_u^{ACI}) = \left[\max \left(0, \hat{\theta}_M - z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)} \right), \hat{\theta}_M + z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)} \right].$$

Alternatively, to ensure the resulting confidence limits for θ to be positive, we construct an ACI for θ based on logarithm transformed MLE by approximating the distribution of $\frac{\log(\hat{\theta}_M) - \log(\theta)}{\sqrt{Var(\log(\hat{\theta}_M))}}$ by the standard normal distribution, where $Var(\log(\hat{\theta}_M))$ can be approximated by delta method as $\widehat{Var}(\log(\hat{\theta}_M)) = \widehat{Var}(\hat{\theta}_M) / \hat{\theta}_M^2$ (see Meeker and Escobar, 1998). Hence, an approximate $100(1-\gamma)\%$ Confidence Interval for θ based on the log transformation is obtained as

$$(\hat{\theta}_l^{ACI*}, \hat{\theta}_u^{ACI*}) = \left[\hat{\theta}_M \exp \left(-\frac{z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)}}{\hat{\theta}_M} \right), \hat{\theta}_M \exp \left(\frac{z_{1-\gamma/2} \sqrt{\widehat{Var}(\hat{\theta}_M)}}{\hat{\theta}_M} \right) \right].$$

3.2 Exact confidence interval

Here, an ECI for the parameter θ of the XLindley distribution is investigated. To obtain an ECI for θ , consider the random variable

$$Q(\theta) = 2 \sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right],$$

given in (6). Since the pivotal quantity $Q(\theta)$ has a chi-square distribution with $2m$ degrees of freedom, an exact $100(1 - \gamma)\%$ confidence interval for θ can be constructed from the relation

$$P \left(\chi_{\gamma/2, 2m}^2 < 2 \sum_{i=1}^m (R_i + 1) \left[\theta X_i + \log \left(\frac{(1 + \theta)^2}{\theta X_i + (1 + \theta)^2} \right) \right] < \chi_{(1-\gamma/2), 2m}^2 \right) = 1 - \gamma,$$

where $\chi_{\gamma/2, 2m}^2$ and $\chi_{(1-\gamma/2), 2m}^2$ denote the lower and upper $\gamma/2$ percentage points of a chi-square distribution with $2m$ degrees of freedom. Since $Q(\theta)$ is a monotonic increasing function in θ , we can compute the $100(1 - \gamma)\%$ ECI for θ as $(\hat{\theta}_l, \hat{\theta}_u)$, where $\hat{\theta}_l$ and $\hat{\theta}_u$, respectively, are the solutions of equations

$$\begin{aligned} 2 \sum_{i=1}^m (R_i + 1) \left[\hat{\theta}_l x_i + \log \left(\frac{(1 + \hat{\theta}_l)^2}{\hat{\theta}_l x_i + (1 + \hat{\theta}_l)^2} \right) \right] &= \chi_{\gamma/2, 2m}^2, \\ 2 \sum_{i=1}^m (R_i + 1) \left[\hat{\theta}_u x_i + \log \left(\frac{(1 + \hat{\theta}_u)^2}{\hat{\theta}_u x_i + (1 + \hat{\theta}_u)^2} \right) \right] &= \chi_{(1-\gamma/2), 2m}^2. \end{aligned}$$

3.3 Confidence interval using the likelihood ratio statistic

Consider the hypothesis test $H_0 : \theta = \theta_0$ against $H_a : \theta \neq \theta_0$. The likelihood ratio test (LRT) statistic for the above hypothesis test is,

$$\lambda(\mathbf{x}) = \frac{\sup_{\theta \in \Theta_0} L(\theta | \mathbf{x})}{\sup_{\theta \in \Theta} L(\theta | \mathbf{x})} = \frac{L(\theta_0 | \mathbf{x})}{L(\hat{\theta} | \mathbf{x})}, \quad (12)$$

where Θ is the parameter space, $\hat{\theta}$ is the MLE of θ and Θ_0 is the parameter space under the hypothesis H_0 . Based on the relations (4) and (12), the LRT statistic is obtained as

$$\begin{aligned} \lambda(\mathbf{x}) &= \left(\frac{\theta_0(1 + \hat{\theta})}{\hat{\theta}(1 + \theta_0)} \right)^{2m} \exp \left(-(\theta_0 - \hat{\theta}) \sum_{i=1}^m (R_i + 1) x_i \right) \prod_{i=1}^m \frac{2 + \theta_0 + x_i}{2 + \hat{\theta} + x_i} \\ &\quad \times \prod_{i=1}^m \left(\frac{(1 + \hat{\theta})^2 [(1 + \theta_0)^2 + \theta_0 x_i]}{(1 + \theta_0)^2 [(1 + \hat{\theta})^2 + \hat{\theta} x_i]} \right)^{R_i}, \end{aligned}$$

where $x_1, \dots, x_m$ are the observed progressively Type-II censored samples with the censoring scheme $R_1, \dots, R_m$. By the asymptotic property of the LRT statistic for large samples, the function $-2 \log \lambda(\mathbf{X})$ in the distribution converges to the chi-square

distribution with one degree of freedom Casella and Berger (2002). In the LRT, the hypothesis H_0 is rejected at the significance level γ whenever $-2 \log \lambda(\mathbf{x}) > \chi_{1-\gamma}^2(1)$. Therefore, based on the complement of the rejection region of the hypothesis H_0 the confidence interval $(1 - \gamma)\%$ is obtained as

$$\begin{aligned} K(\theta) &= \{\theta : -2 \log \lambda(\mathbf{x}) \leq \chi_{(1-\gamma),1}^2\} \\ &= \{\theta : -2[\log L(\theta_0|\mathbf{x}) - \log L(\hat{\theta}|\mathbf{x})] \leq \chi_{(1-\gamma),1}^2\} \\ &= \{\theta : \log L(\theta_0|\mathbf{x}) \geq \log L(\hat{\theta}|\mathbf{x}) - \frac{1}{2}\chi_{(1-\gamma),1}^2\}. \end{aligned}$$

To construct the LRT based confidence interval (LRT-CI) for θ , the lower and upper bounds, denoted by θ_L and θ_U , are obtained by numerically solving the equation $-2 \log \lambda(\mathbf{x}) = \chi_{(1-\gamma),1}^2$, with respect to θ . Given the unimodality of the likelihood function (established in Subsection 2.1), this equation has exactly two roots. The MLE, $\hat{\theta}_M$, serves as the initial reference point, ensuring that the search for θ_L is conducted to its left and for θ_U to its right. The uniqueness of the roots is guaranteed by the strict unimodality of the likelihood function, which shows that the function $-2 \log \lambda(\mathbf{x})$ is strictly increasing for $\theta < \hat{\theta}_M$ and strictly decreasing for $\theta > \hat{\theta}_M$. Consider the function as follows

$$\begin{aligned} h(\theta) &= -2 \log \lambda(\mathbf{x}) - \chi_{(1-\gamma),1}^2, \\ \lim_{\theta \rightarrow 0} h(\theta) &= +\infty, \quad \lim_{\theta \rightarrow \infty} h(\theta) = +\infty, \quad \text{and} \quad -2 \log \lambda(\hat{\theta}_M) = 0 < \chi_{(1-\gamma),1}^2. \end{aligned}$$

Therefore, $h(\theta)$ changes sign from positive to negative in the interval $(0, \hat{\theta}_M)$ and from negative to positive in the interval $(\hat{\theta}_M, \infty)$. This property guarantees the uniqueness of the roots in each of the two intervals. To find the lower and upper bounds of the interval, it is sufficient to obtain the roots of the $h(\theta) = 0$ equation in the interval $(0, \hat{\theta}_M)$ and the interval $(\hat{\theta}_M, \infty)$, respectively.

3.4 Bootstrap confidence intervals

It is evident that the confidence intervals based on the asymptotic results do not perform very well for small sample size. Therefore, we propose two confidence interval on the bootstrap methods: (i) percentile bootstrap method (Boot-p) and (ii) bootstrap-t method (Boot-t).

3.4.1 Boot-p method

A widely used bootstrap method is the parametric Boot-p. A typical algorithm for the confidence interval of works as follows:

Algorithm 3.1. *Boot-p method algorithm:*

- i. Repeat steps i-iv in Algorithm 2.1.
- ii. Arrange $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$ in ascending orders and obtain $\hat{\theta}_{(1)}^*, \dots, \hat{\theta}_{(B)}^*$.
- iii. Pair of two-sided $100(1 - \gamma)\%$ Boot-p confidence interval for θ is given by $[\hat{\theta}_{B(\gamma/2)}^*, \hat{\theta}_{B(1-\gamma/2)}^*]$, where $\hat{\theta}_{B(\gamma/2)}^*$ is the percentile of $(\gamma/2)$ in the bootstrap sample.

3.4.2 Bootstrap-t method

Algorithm 3.2. *Boot-t method algorithm:*

- i. Repeat steps i-iv in Algorithm 2.1.
- ii. The variance of $\widehat{V}(\widehat{\theta}^*)$ is calculated using the relation $\widehat{V}(\widehat{\theta}^*) = \frac{\widehat{\theta}^{*2}}{\mathbf{V}}$, where $\mathbf{V}$ denotes the inverse of the Fisher information.
- iii. The T^* statistic is defined as $T^* = \frac{(\widehat{\theta}^* - \widehat{\theta})}{\sqrt{\widehat{V}(\widehat{\theta}^*)}}$.
- iv. Repeat steps i and ii, B times and obtain $T_1^*, \dots, T_B^*$.
- v. Subsequently, $T_1^*, \dots, T_B^*$ are arranged in ascending order and represented by $T_{(1)}^*, \dots, T_{(B)}^*$. Hence, the two-sided $100(1 - \gamma)\%$ Boot-t confidence interval for θ is constructed as $\left(\widehat{\theta} + T_{B(\gamma/2)}^* \sqrt{\frac{\widehat{\theta}^{*2}}{\mathbf{V}}}, \widehat{\theta} + T_{B(1-\gamma/2)}^* \sqrt{\frac{\widehat{\theta}^{*2}}{\mathbf{V}}} \right)$.

3.5 Bayesian confidence intervals

Based on the above simulated values of θ , $\{\theta_q; q = M + 1, \dots, N\}$, and using the method proposed by Chen and Shao (1999), we obtain the HPD credible interval for θ . We assume that $\theta_{[M+1]} < \dots < \theta_{[N-M]}$ is the ordered MCMC sample of $\{\theta_q; q = M + 1, \dots, N\}$. A $100(1 - \gamma)\%$ the HPD credible interval for θ is

$$\left(\theta_{[(\gamma/2)(N-M)]}, \theta_{[(1-\gamma/2)(N-M)]} \right),$$

where $\theta_{[(\gamma/2)(N-M)]}$ and $\theta_{[(1-\gamma/2)(N-M)]}$ are the $[(\gamma/2)(N - M)]$ -th smallest integer and the $[(1 - \gamma/2)(N - M)]$ -th smallest integer of $\{\theta_q; q = 1, 2, \dots, N\}$, respectively.

4 Prediction of future failures times

In this section, we address prediction of future failures times under classical and Bayesian approaches. Our interest is to obtain the prediction of the life-lengths $Z = Y_{jk}$, $j = 1, 2, \dots, m$ and $k = 1, 2, \dots, R_j$ of all censored units in all m stages of censoring based on observed data $\mathbf{x} = (x_1, \dots, x_m)$. Different methods such as the MLP, the BUP, the CMP and the BP are considered.

4.1 Maximum likelihood predictor

The likelihood approach was first proposed by Kaminsky and Rhodin (1985) to predict the future order statistics and estimate the parameters involved in the model. Based on the censored sample $\mathbf{x}$, the predictive likelihood function (PLF) of $Z = Y_{jk}$ and θ is given by

$$\begin{aligned} L(z, \theta | \mathbf{x}) &\propto \frac{[F(z, \theta) - F(x_j, \theta)]^{k-1} (1 - F(z, \theta))^{R_j - k} f(z, \theta)}{(1 - F(x_j, \theta))^{R_j}} \\ &\times \prod_{i=1, i \neq j}^m f(x_i, \theta) \prod_{i=1, i \neq j}^m [1 - F(x_i, \theta)]^{R_i}, \quad z > x_j. \end{aligned} \quad (13)$$

Considering the relations (2), (1) and (13) the logarithm of the PLF Z and θ as

$$\begin{aligned} \log L(z, \theta | \mathbf{x}) \propto & 2(m+1) \log \theta - 2n \log(1+\theta) - \theta z(R_j - k + 1) + \theta x_j R_j + \log(2 + \theta + z) \\ & - \theta \left[\sum_{i=1, i \neq j}^m x_i (1 + R_i) \right] + (R_j - k) \log[\theta z + (1 + \theta)^2] \\ & - R_j \log[\theta x_j + (1 + \theta)^2] + \sum_{i=1, i \neq j}^m \log(2 + \theta + x_i) + \sum_{i=1, i \neq j}^m R_i \log[\theta x_i + (1 + \theta)^2] \\ & + (k-1) \log \left[e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2) \right]. \end{aligned} \quad (14)$$

Maximizing (14) with respect to θ and z , we could find the MLP of Z and the predictive MLE (PMLE) of θ . It follows from (14) that the predictive likelihood equations (PLEs) for Z and θ are given by

$$\begin{aligned} \frac{\partial \log L(z, \theta | \mathbf{x})}{\partial \theta} = & \frac{2(m+1)}{\theta} - \frac{2n}{1+\theta} - z(R_j - k + 1) + x_j R_j - \sum_{i=1, i \neq j}^m x_i (1 + R_i) \\ & + \frac{(R_j - k)[z + 2(1 + \theta)]}{\theta z + (1 + \theta)^2} + \frac{1}{2 + \theta + z} - \frac{R_j[x_j + 2(1 + \theta)]}{\theta x_j + (1 + \theta)^2} \\ & + \sum_{i=1, i \neq j}^m \frac{1}{2 + \theta + x_i} + \sum_{i=1, i \neq j}^m \frac{R_i[x_i + 2(1 + \theta)]}{\theta x_i + (1 + \theta)^2} \\ & + (k-1) \left[\frac{e^{-\theta z} ([z + 2(1 + \theta)] - z[\theta z + (1 + \theta)^2])}{e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2)} \right] \\ & + (k-1) \left[\frac{e^{-\theta x_j} (x_j[\theta x_j + (1 + \theta)^2] - [x_j + 2(1 + \theta)])}{e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2)} \right], \end{aligned} \quad (15)$$

$$\begin{aligned} \frac{\partial \log L(z, \theta | \mathbf{x})}{\partial \theta} = & -\theta(R_j - k + 1) + \frac{\theta(R_j - k)}{\theta z + (1 + \theta)^2} + \frac{1}{2 + \theta + z} \\ & + (k-1) \left[\frac{e^{-\theta z} [\theta - z\theta^2 - \theta(\theta + 1)^2]}{e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2)} \right]. \end{aligned} \quad (16)$$

By solving (15) and (16) with respect to Z and θ simultaneously, the MLP of Z and PMLE of θ can be obtained. Numerical methods can be applied to obtain the solutions of these PLEs and obtain the MLP of Z and PMLE of θ .

4.2 Best unbiased predictor

According to the Markovian property of record statistics, the conditional pdf Z given $\mathbf{x} = (x_1, \dots, x_m)$ is equal to the pdf Z given x_j (see Balakrishnan and Aggarwala, 2000). The conditional density function Z given x_j for $z \geq x_j$ is as

$$f_Z(z | x_j; \theta) = k \binom{R_j}{k} \frac{[F(z, \theta) - F(x_j, \theta)]^{k-1} (1 - F(z, \theta))^{R_j-k} f(z, \theta)}{(1 - F(x_j, \theta))^{R_j}}, \quad (17)$$

According to the relations (2), (1) and (17), the conditional distribution of Z given the condition of x_j is calculated as

$$\begin{aligned} f_Z(z|x_j; \theta) &= k \binom{R_j}{k} \exp[-\theta z(R_j - k + 1) + \theta R_j x_j] \\ &\quad \times \frac{\theta^2 (2 + \theta + z) [\theta z + (1 + \theta)^2]^{R_j - k}}{[\theta x_j + (1 + \theta)^2]^{R_j}} \\ &\quad \times [e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2)]^{k-1}, \quad z > x_j. \end{aligned}$$

The BUP for Z is obtained by solving the following equation

$$\begin{aligned} \hat{Z}_{BUP} &= E(Z|x_j) = \int_{x_j}^{\infty} z f_Z(z|x_j; \theta) dz = \frac{k \binom{R_j}{k} \theta^2 e^{\theta R_j x_j}}{[\theta x_j + (1 + \theta)^2]^{R_j}} \\ &\quad \times \int_{x_j}^{\infty} z \exp[-\theta z(R_j - k + 1)] (2 + \theta + z) [\theta z + (1 + \theta)^2]^{R_j - k} \\ &\quad \times [e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) - e^{-\theta z} (\theta z + (1 + \theta)^2)]^{k-1} dz. \end{aligned}$$

If the parameters θ is unknown, the BUP of Z can be approximated by replacing θ by its MLE.

4.3 Conditional median predictor

As another possible predictor, it can be considered CMP (Raqab and Nagaraja, 1995). Note that a predictor $\hat{Z}$ is called the CMP of Z , if it is the median of the conditional pdf of Z given the x_j observation. That is,

$$Pr_{\theta}(Z \leq \hat{Z}|X_j = x_j) = Pr_{\theta}(Z \geq \hat{Z}|X_j = x_j).$$

Hence, the CMP of Z is obtained by solving the equation

$$\int_{x_j}^{\hat{Z}} f_Z(z|x_j; \theta) dz = \frac{1}{2}.$$

We have

$$\begin{aligned} \frac{1}{2} &= \int_{x_j}^{\hat{Z}} f_Z(z|x_j; \theta) dz \\ &= \int_{x_j}^{\infty} k \binom{R_j}{k} \frac{[F(z, \theta) - F(x_j, \theta)]^{k-1} (1 - F(z, \theta))^{R_j - k} f(z, \theta)}{(1 - F(x_j, \theta))^{R_j}} dz \\ &= \int_{x_j}^{\hat{Z}} k \binom{R_j}{k} \left[\frac{1 - F(z, \theta)}{1 - F(x_j, \theta)} \right]^{R_j - k} \left[1 - \frac{1 - F(z, \theta)}{1 - F(x_j, \theta)} \right]^{k-1} \frac{f(z, \theta)}{1 - F(x_j, \theta)} dz. \end{aligned}$$

Using the transformation $u = (1 - F(z, \theta))/(1 - F(x_j, \theta))$, we have

$$\frac{1}{2} = \int_0^{\frac{1 - F(\hat{Z}, \theta)}{1 - F(x_j, \theta)}} \frac{u^{R_j - k} (1 - u)^{k-1}}{\text{Beta}(R_j - k + 1, k)} du. \quad (18)$$

Therefore, by (18), the CMP of Z can be obtained by solving the equation

$$\frac{1 - F(\widehat{Z}, \theta)}{1 - F(x_j, \theta)} = \frac{e^{-\theta \widehat{Z}} (\theta \widehat{Z} + (1 + \theta)^2)}{e^{-\theta x_j} (\theta x_j + (1 + \theta)^2)} = MB(R_j - k + 1, k), \quad (19)$$

where $MB(R_j - k + 1, k)$ is the median of the beta distribution with parameters $R_j - k + 1$ and k . Now, from (19), the for Z , $\widehat{Z}^{CMP}$, can be obtained by solving the nonlinear equation

$$\left[e^{-\theta \widehat{Z}^{CMP}} (\theta \widehat{Z}^{CMP} + (1 + \theta)^2) \right] - MB(R_j - k + 1, k) \left[e^{-\theta x_j} (\theta x_j + (1 + \theta)^2) \right] = 0.$$

If the parameters θ is unknown, the CMP of Z can be approximated by replacing θ by its MLE.

4.4 Bayesian predictor

The Bayesian predictive pdf Z given x_j is

$$f_1(z|x_j; \theta) = \int_0^\infty f_Z(z|x_j; \theta) \pi(\theta|\mathbf{x}) d\theta, \quad z \geq x_j.$$

The BP Z under the SEL function is given by

$$Z_{BP} = \int_{x_j}^\infty z f_1(z|x_j; \theta) dz. \quad (20)$$

Due to the complexity of the BP pdf $f_1(z|x_j; \theta)$, BPs are not obtained in a simple and explicit way. Therefore, based on the examples generated by the Metropolis-Hastings algorithm $\{\theta_q; q = M + 1, \dots, N\}$, the BP pdf $f_1(z|x_j; \theta)$, is estimated as

$$\widehat{f}_1(z|x_j; \theta) = \frac{1}{N - M} \sum_{q=M+1}^N f_Z(z|x_j; \theta_q). \quad (21)$$

According to the relations (20) and (21), the approximate BP for Z under the SEL function is obtained as

$$\begin{aligned} \widehat{Z}_{BP} &= \int_{x_j}^\infty z \widehat{f}_1(z|x_j; \theta) dz = \frac{1}{N - M} \sum_{q=M+1}^N \int_{x_j}^\infty z f_Z(z|x_j; \theta_q) dz \\ &= \frac{1}{N - M} \sum_{q=M+1}^N I_1(x_j; \theta_q), \end{aligned}$$

where $I_1(x_j; \theta) = \int_{x_j}^\infty z f_Z(z|x_j; \theta) dz$.

5 Interval prediction

In this section, we consider two approaches to obtain the prediction intervals for $Z = Y_{jk}$, $j = 1, \dots, m$ and $k = 1, \dots, R_j$ based on the observed progressively Type-II censoring sample $\mathbf{x} = (x_1, \dots, x_m)$.

5.1 Pivotal approach

Considering the pdf of Z under the condition of x_j of the relation (17) or subsection 4.3, it is clear that the random variable

$$U = \frac{1 - F(Z, \theta)}{1 - F(x_j, \theta)} = \frac{e^{-\theta Z} (\theta Z + (1 + \theta)^2)}{e^{-\theta x_j} (\theta x_j + (1 + \theta)^2)} | X_j = x_j,$$

has a beta distribution with parameters $R_j - k + 1$ and k (denoted by Beta $(R_j - k + 1, k)$). Therefore, U is a pivotal quantity which can be used to construct prediction interval of Z . Using the above pivotal quantity, a prediction interval of $100(1 - \gamma)\%$ for Z is obtained from the relation, $\Pr(B_{\gamma/2} < U < B_{1-\gamma/2} | x_j) = 1 - \gamma$, where $B_{\gamma/2}$ is the 100γ -th percentile of Beta $(R_j - k + 1, k)$ distribution. Therefore, a $100(1 - \gamma)\%$ predictive interval for Z is (L_1^*, U_1^*) , where L_1^* and U_1^* satisfy, respectively, the following non-linear equations

$$\begin{aligned} e^{-\theta L_1^*} [\theta L_1^* + (1 + \theta)^2] - B_{\gamma/2} (e^{-\theta x_j} [\theta x_j + (1 + \theta)^2]) &= 0, \\ e^{-\theta U_1^*} [\theta U_1^* + (1 + \theta)^2] - B_{1-\gamma/2} (e^{-\theta x_j} [\theta x_j + (1 + \theta)^2]) &= 0. \end{aligned}$$

5.2 Highest conditional density approach

Here, the HCD method is used for constructing the prediction interval of Z . The conditional distribution of U given $X_j = x_j$ follows a Beta $(R_j - k + 1, k)$ distribution with pdf

$$f(u | x_j) = \frac{u^{R_j - k} (1 - u)^{k-1}}{\text{Beta}(R_j - k + 1, k)} du, \quad 0 < u < 1,$$

which is a unimodal function of u for $1 < k < R_j$. A $100(1 - \gamma)\%$ HCD predictive interval for Z is (L_2^*, U_2^*) , where L_2^* and U_2^* are respectively the solutions of

$$\begin{aligned} e^{-\theta L_2^*} [\theta L_2^* + (1 + \theta)^2] - d_1 (e^{-\theta x_j} [\theta x_j + (1 + \theta)^2]) &= 0, \\ e^{-\theta U_2^*} [\theta U_2^* + (1 + \theta)^2] - d_2 (e^{-\theta x_j} [\theta x_j + (1 + \theta)^2]) &= 0, \end{aligned}$$

with d_1 and d_2 satisfy

$$\int_{d_1}^{d_2} f(u | x_j) du = 1 - \gamma, \quad f(d_1 | x_j) = f(d_2 | x_j). \quad (22)$$

(22) can be simplified as

$$B_{d_2}(R_j - k + 1, k) - B_{d_1}(R_j - k + 1, k) = 1 - \gamma, \quad \left(\frac{1 - d_2}{1 - d_1}\right)^{k-1}, \quad \left(\frac{d_1}{d_2}\right)^{R_j - k},$$

where $B_t(a, b) = \int_0^t \frac{x^{a-1} (1-x)^{b-1}}{\text{Beta}(a, b)} dx$ is the incomplete beta function.

For $k = R_j$, the pdf $f(u | x_j)$ is decreasing function with $f(0 | x_j) = R_j$ and $f(1 | x_j) = 0$. Hence, the values of d_1 and d_2 must be calculated as

$$\int_0^{d_2} f(u | x_j) du = 1 - \gamma \implies d_2 = 1 - \gamma^{1/R_j}.$$

For $k = 1$, the pdf $f(u|x_j)$ is increasing function with $f(0|x_j) = 0$ and $f(1|x_j) = R_j$. Hence, the values of d_1 and d_2 must be calculated as

$$\int_{d_1}^1 f(u|x_j) du = 1 - \gamma \implies d_1 = \gamma^{1/R_j}.$$

Finally, for $k = R_j = 1$, $f(u|x_j)$ is pdf of $U(0, 1)$ distribution. Consequently, the end points of the prediction interval can be chosen simply with equal probabilities. That is, $d_1 = \gamma/2$ and $d_2 = 1 - \gamma/2$.

5.3 Bayesian prediction interval

A $100(1 - \gamma)\%$ Bayesian prediction interval (BPI) for Z is given by (L_3^*, U_3^*) , where L_3^* and U_3^* can be obtained by solving the following two nonlinear equations simultaneously

$$\begin{aligned} \Pr(Z > L_3^*|x_j) &= \int_{L_3^*}^{\infty} f_1(z|x_j; \theta) dz = \int_{L_3^*}^{\infty} \hat{f}_1(z|x_j; \theta) dz = 1 - \frac{\gamma}{2}, \\ \Pr(Z > U_3^*|x_j) &= \int_{U_3^*}^{\infty} f_1(z|x_j; \theta) dz = \int_{U_3^*}^{\infty} \hat{f}_1(z|x_j; \theta) dz = \frac{\gamma}{2}. \end{aligned}$$

By using $\hat{f}_1(z|x_j; \theta)$ defined in (21) and using the MCMC sample $\{\theta_q; q = M + 1, \dots, N\}$, we can compute the lower and upper bounds L_3^* and U_3^* from the relations

$$\begin{aligned} 1 - \frac{\gamma}{2} &= \frac{1}{N - M} \sum_{q=M+1}^N \int_{L_3^*}^{\infty} f_Z(z|x_j; \theta_q) dz, \\ \frac{\gamma}{2} &= \frac{1}{N - M} \sum_{q=M+1}^N \int_{U_3^*}^{\infty} f_Z(z|x_j; \theta_q) dz. \end{aligned}$$

6 Illustrative example: Gastric cancer survival data

Here we consider a real dataset and discuss the various estimation techniques described in this paper. This dataset is taken from Stablein et al. (1981) and shows the survival time (in years) of a group of patients treated with chemotherapy plus radiation therapy for gastric cancer.

We fitted nine distributions to the Gastric Cancer Survival Data: Exponential, Weibull, Lindley, XLindley, power Lindley (PL), power XLindley (PXL), generalized Lindley (GL), extended Lindley (EL), and exponentiated XLindley (EXL) distributions. The XLindley, Lindley and exponential distributions have one parameter but other distributions have two parameters. Each distribution was fitted by the method of maximum likelihood.

Table 1 lists the parameter estimates the log-likelihood values, the values of the Akaike information criterion (AIC), the values of the Bayesian information criterion (BIC), the p-values based on the Kolmogorov-Smirnov (K-S) statistic. We see that the XLindley distribution has the smallest log-likelihood value, the smallest AIC value, the smallest BIC value and the largest p-value based on the K-S test. After the XLindley

Table 1: Parameter estimates, -log-likelihood values and goodness of fit measures.

Distribution	$\hat{\lambda}$	$\hat{\theta}$	$-\log L$	AIC	BIC	K-S	p-value
Exponential	—	0.372	79.45	160.90	163.06	0.174	0.052
Weibull	2.814	2.986	68.23	140.46	144.78	0.085	0.623
Lindley	—	0.481	72.15	146.30	148.46	0.112	0.312
XLindley	—	0.623	66.89	137.78	139.94	0.067	0.812
PL	0.412	1.234	67.12	138.24	142.56	0.071	0.754
PXL	0.598	1.112	66.98	137.96	142.28	0.069	0.789
GL	0.501	1.045	69.87	143.78	148.06	0.098	0.451
EL	0.489	0.987	70.12	144.24	148.56	0.104	0.398
EXL	0.612	0.956	67.23	138.46	142.78	0.070	0.772

distribution, the power XLindley, power Lindley and exponentiated XLindley distributions have the smallest log-likelihood value, the smallest AIC value, the smallest BIC value, and the largest p-value based on the K-S test. The exponential, Lindley, extended Lindley, generalized Lindley and Weibull distributions have the largest log-likelihood value, the largest AIC value, the largest BIC value, the smallest p-value based on the K-S test, respectively. Figure 1 displays the empirical histograms of the data sets alongside the fitted pdfs of the considered models for Data. Moreover, Figure 2 exhibits the probability-probability (P-P) plots for Data. From Figures 1 and 2, we may also conclude visually that the XLindley, power XLindley, power Lindley and exponentiated XLindley distributions offer the best fits for Data among the distributions considered here.

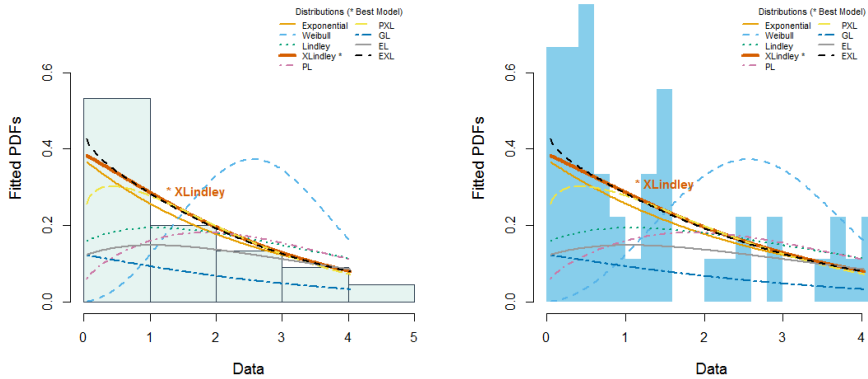


Figure 1: Empirical histogram of data and the fitted pdfs of the considered model.

From this data, we generate two progressively censored samples using the following schemes:

scheme	m	R
1	40	$(0, 0, 0, 0, 0, 0, 0, 0, 0, \dots, 0^{*40}, 5)$
2	28	$(0, 0, 0, 0, 0, 0, 5, 0, 0, 0, 0, 0, 3, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 4, 0, 0, 5)$

Scheme 1 shows the Type-II censoring, in which all censoring occurs only at the end of

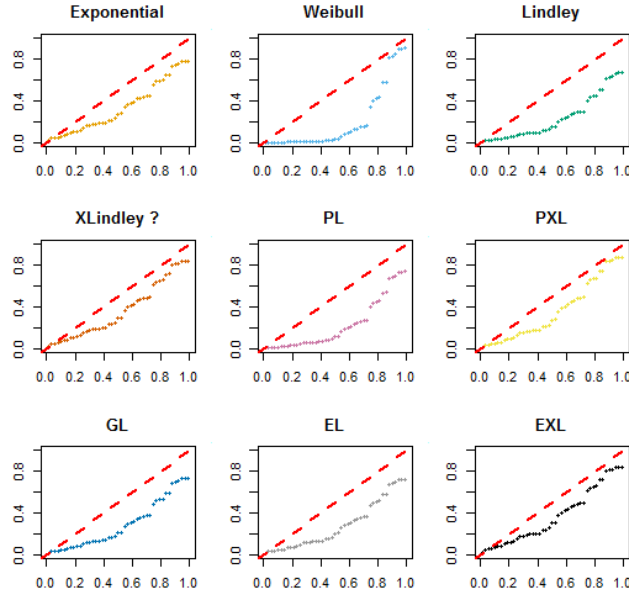


Figure 2: Probability plots for the fits of the Exponential, Weibull, Lindley, XLindley, PL, PXL, GL, EL, and EXL distributions.

the experiment. This scheme is often designed to save costs by allowing the experiment to be terminated early after a fixed number of failures. To achieve the Type-II censored sample, suppose that in this study, as soon as the 40th survival time was observed, the experiment was stopped and the survival times of the remaining 5 patients were not observed. Therefore, we are faced with a Type-II censored sample with $n = 45$ and $m = 40$. Based on scheme 2, the generated progressively censored samples are as follows:

scheme	Censored data						
2	0.047	0.115	0.121	0.132	0.164	0.197	0.203
	0.260	0.282	0.296	0.334	0.395	0.458	0.466
	0.501	0.507	0.529	0.540	0.641	0.644	0.841
	0.863	1.099	1.271	1.447	1.485	1.553	1.589

Scheme 2 involves progressively removals at various intermediate failure times, reflecting situations where units are withdrawn at multiple stages (e.g., due to scheduled inspections, budget constraints, or ethical considerations in clinical trials). The specific pattern used here is chosen to demonstrate the flexibility of the proposed methods under a more complex, non-monotone censoring plan.

We computed point and interval estimates of θ from the XLindley distribution using the various estimation techniques described in Sections 2 and 3. Bayes estimates are obtained using Lindley's approximation, Metropolis-Hastings procedures and TK based on the SEL function. The results are presented in Tables 2 and 3. For computing the Bayes estimates and corresponding confidence intervals, we have used the approximate non-informative prior with $\alpha = \beta = 0.0001$. For MCMC method using Metropolis-Hastings algorithm, we sample $N = 50000$ values and discard the initial $M = 5000$

as burn-in sample and calculate the Bayes estimates and corresponding confidence intervals based on the remaining $N - M = 45000$ samples.

According to the Geweke test, the p-values under scheme 1 and scheme 2 are calculated to be 0.667 and 0.383, respectively. Hence, for both schemes, the mean of the parameter θ at the beginning of the chain (after the initial deletion) is not significantly different from its mean at the end of the chain and the chain converges well to the posterior distribution. The acceptance rate of the algorithm is about 70%, which indicates that the chain moves efficiently. According to the Heidelberger-Welch test, the p-values under scheme 1 and scheme 2 are calculated to be 0.213 and 0.343, respectively. Hence, for both schemes, there is no statistically significant evidence of chain instability and even considering the entire chain after the initial deletion, the data show stationary behavior and there is no need to delete any further part of the chain beginning. The estimated autocorrelation of the generated chain is low for both schemes (0.12 for Scheme 1 and 0.18 for Scheme 2), indicating that the Metropolis-Hastings algorithm, with the proposed normal proposal distribution and appropriate variance, has performed an efficient exploration of the posterior space with low dependence between successive draws. The Gelman-Rubin test value is obtained for scheme 1 equal to 1.001 and for scheme 2 equal to 1.002. These values are significantly lower than the conventional threshold of 1.1, which clearly indicates that the between-chain variance is negligible compared to the within-chain variance. As a result, the complete convergence of the MCMC chains to the posterior distribution is confirmed for both censoring schemes. Also, Figures 3 and 4 present the histograms of the sequences generated by the Metropolis-Hastings algorithm for Schemes 1 and 2. These histograms indicate that the choice of the normal distribution as an approximation to both the posterior pdf and the proposal distribution is entirely appropriate. Moreover, to assess the convergence of the generated chains, trace plots can be employed. These plots show that the generated values of θ from the Metropolis-Hastings algorithm are scattered around the mean for each scheme. In the above example, to verify the existence and uniqueness of the MLE of the parameter θ , the log-likelihood function $\log L(\theta; \mathbf{x})$ is plotted. Figure 5 shows that the log-likelihood function is unimodal over the interval $(0, \infty)$. Consequently, the MLE exists and is unique.

In the prediction context, we apply the methods described in Sections 4 and 5 to obtain the point and interval predictors of $Z = Y_{j:k}$, $j = 1, \dots, m$ and $k = 1, \dots, R_j$. The results are presented in Tables 4 and 5.

Table 2: Point estimates of θ for the gastric cancer survival data.

Method	MLE	MBE	PBE	BES	BEL	BETK
Scheme 1	0.8847	0.8785	0.8942	0.8868	0.8873	0.8864
Scheme 2	1.0653	1.0660	1.0934	1.0683	1.0773	1.0682

7 Monte Carlo simulation study

In this section, we compare the performances of the different methods of estimation and prediction based on Monte Carlo simulations. For each choice of n and m , four

Table 3: Interval estimates of θ for the gastric cancer survival data.

	Method			Method	
Scheme 1	ACI	(0.6612, 1.1083)	ACI*	(0.6872, 1.1391)	
	ECI	(0.6790, 1.1050)	LRT-CI	(0.6899, 1.1181)	
	Boot-p	(0.6986, 1.1424)	Boot-t	(0.6978, 1.1363)	
	HPD	(0.6798, 1.1128)			
Scheme 2	ACI	(0.6965, 1.4342)	ACI*	(0.7536, 1.5061)	
	ECI	(0.7794, 1.4016)	LRT-CI	(0.7870, 1.4128)	
	Boot-p	(0.8108, 1.4817)	Boot-t	(0.8162, 1.4593)	
	HPD	(0.7796, 1.4066)			

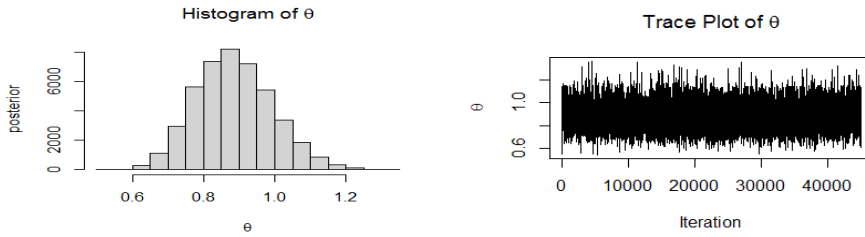


Figure 3: Histogram and trace plots of XLindley parameter based on Scheme 1.

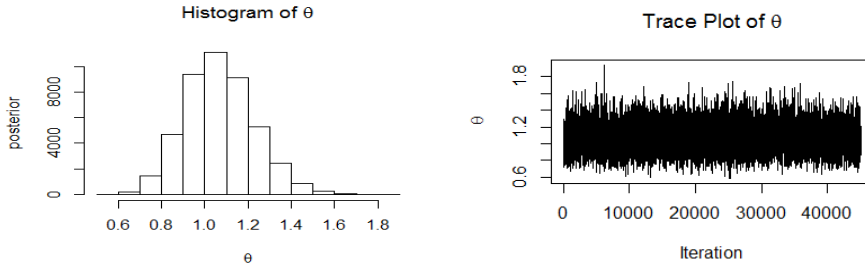


Figure 4: Histogram and trace plots of XLindley parameter based on Scheme 2.

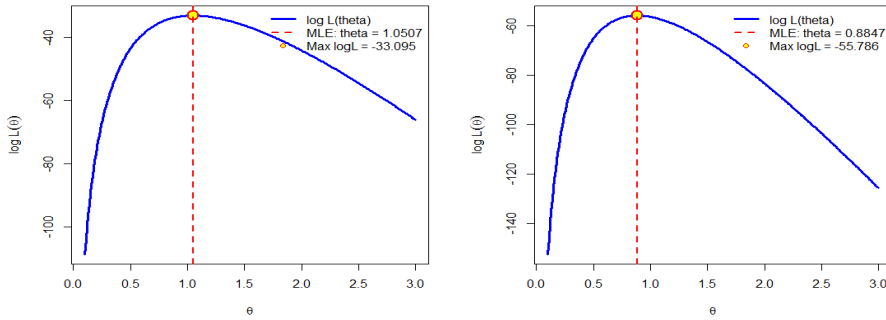
Figure 5: Plot of the log-likelihood function $\log L(\theta; x)$ versus θ for Scheme 1 and Scheme 2.

Table 4: Point prediction of the future lifetimes based on the gastric cancer survival data.

Scheme		PMLE	MLP	BUP	CMP	Bayesian
1	$Y_{40:1}$	0.9004	3.5782	3.8428	3.7614	3.8467
	$Y_{40:2}$	0.9003	3.8695	4.1705	4.0757	4.1798
	$Y_{40:3}$	0.9001	4.2422	4.6044	4.4890	4.6215
	$Y_{40:4}$	0.8999	4.7668	5.2484	5.0931	5.2750
	$Y_{40:5}$	0.8998	5.6500	6.5175	6.2272	6.5607
2	$Y_{7:1}$	1.0836	0.2030	0.4411	0.3695	0.4491
	$Y_{7:2}$	1.0950	0.4650	0.7343	0.6516	0.7481
	$Y_{7:3}$	1.0941	0.7974	1.1184	1.0202	1.1486
	$Y_{7:4}$	1.0929	1.2582	1.6826	1.5519	1.7255
	$Y_{7:5}$	1.0921	2.0304	2.7788	2.5362	2.8417
	$Y_{13:1}$	1.0954	0.4580	0.8465	0.7305	0.8574
	$Y_{13:2}$	1.0939	0.9258	1.4161	1.2646	1.4378
	$Y_{13:3}$	1.0926	1.7136	2.5196	2.2576	2.5796
	$Y_{25:1}$	1.0940	1.4470	1.7277	1.6427	1.7356
	$Y_{25:2}$	1.0933	1.7646	2.0979	1.9942	2.1163
	$Y_{25:3}$	1.0926	2.2078	2.6454	2.5081	2.6726
	$Y_{25:4}$	1.0917	2.9564	3.7179	3.4685	3.7702
	$Y_{28:1}$	1.0940	1.5890	1.8129	1.7449	1.8184
	$Y_{28:2}$	1.0934	1.8339	2.0904	2.0107	2.1030
	$Y_{28:3}$	1.0929	2.1472	2.4569	2.3609	2.4841
	$Y_{28:4}$	1.0923	2.5851	2.9999	2.8705	3.0355
	$Y_{28:5}$	1.0917	3.3262	4.0656	3.8235	4.1305

Table 5: Interval prediction of the future lifetimes based on the gastric cancer survival data.

Scheme		Pivotal	HCD	BPI
1	$Y_{40:1}$	(3.584, 4.547)	(3.578, 4.366)	(3.584, 4.582)
	$Y_{40:2}$	(3.649, 5.224)	(3.613, 5.029)	(3.649, 5.294)
	$Y_{40:3}$	(3.788, 6.068)	(3.409, 5.743)	(3.784, 6.191)
	$Y_{40:4}$	(4.018, 7.360)	(4.106, 8.243)	(4.010, 7.561)
	$Y_{40:5}$	(4.433, 10.26)	(4.624, ∞)	(4.414, 10.58)
2	$Y_{7:1}$	(0.209, 1.071)	(0.202, 0.911)	(0.209, 1.127)
	$Y_{7:2}$	(0.268, 1.666)	(0.234, 1.495)	(0.267, 1.770)
	$Y_{7:3}$	(0.393, 2.399)	(0.489, 2.117)	(0.388, 2.580)
	$Y_{7:4}$	(0.600, 3.509)	(0.679, 4.261)	(0.586, 3.774)
	$Y_{7:5}$	(0.970, 5.977)	(1.139, ∞)	(0.941, 6.381)
	$Y_{13:1}$	(0.468, 1.870)	(0.458, 1.611)	(0.467, 1.950)
	$Y_{13:2}$	(0.575, 3.112)	(0.575, 3.112)	(0.572, 3.278)
	$Y_{13:3}$	(0.864, 5.670)	(0.996, ∞)	(0.852, 6.002)
	$Y_{25:1}$	(1.454, 2.414)	(1.447, 2.284)	(1.454, 2.529)
	$Y_{25:2}$	(1.690, 4.380)	(1.497, 3.082)	(1.525, 3.373)
	$Y_{25:3}$	(1.690, 4.380)	(1.190, 3.990)	(1.683, 4.586)
	$Y_{25:4}$	(2.016, 6.848)	(2.164, ∞)	(2.001, 7.245)
	$Y_{28:1}$	(1.594, 2.410)	(1.589, 2.257)	(1.594, 2.455)
	$Y_{28:2}$	(1.650, 2.980)	(1.618, 2.816)	(1.649, 3.074)
	$Y_{28:3}$	(1.650, 2.980)	(1.445, 3.416)	(1.762, 3.840)
	$Y_{28:4}$	(1.962, 4.773)	(2.037, 5.512)	(1.951, 5.019)
	$Y_{28:5}$	(2.313, 7.295)	(2.475, ∞)	(2.290, 7.595)

progressively Type-II right censored schemes are considered, which are listed below.

Scheme 1: $R_1 = \dots = R_{m-1} = 0$ and $R_m = n - m$.

Table 6: Simulated biases and MSEs of the point estimation procedures of $\theta = 0.5$.

				Lindleys				MCMC		TK		
n	m	Scheme		MLE	MBE	PBE	Prior1	Prior2	Prior1	Prior2	Prior1	Prior2
10	4	1	Bias	0.161	0.164	0.303	0.473	0.044	0.169	0.094	0.172	0.091
			MSE	0.134	0.135	0.267	0.195	0.037	0.144	0.050	0.146	0.049
		2	Bias	0.020	0.015	0.070	0.010	0.017	0.015	0.026	0.014	0.024
			MSE	0.029	0.029	0.050	0.042	0.014	0.030	0.018	0.030	0.018
		3	Bias	0.145	0.148	0.290	0.382	0.065	0.154	0.076	0.157	0.076
			MSE	0.138	0.139	0.291	0.179	0.023	0.149	0.053	0.151	0.052
20	10	4	Bias	0.280	0.285	0.448	0.502	0.040	0.295	0.175	0.296	0.180
			MSE	0.148	0.151	0.310	0.184	0.015	0.163	0.055	0.163	0.057
		1	Bias	0.034	0.036	0.068	0.043	0.030	0.036	0.029	0.036	0.029
			MSE	0.017	0.017	0.023	0.020	0.011	0.017	0.013	0.017	0.013
		2	Bias	0.028	0.030	0.063	0.036	0.025	0.030	0.023	0.030	0.023
			MSE	0.016	0.016	0.022	0.018	0.011	0.017	0.012	0.017	0.012
20	14	3	Bias	0.025	0.027	0.059	0.033	0.022	0.027	0.021	0.027	0.021
			MSE	0.017	0.018	0.024	0.021	0.011	0.018	0.013	0.018	0.013
		4	Bias	0.031	0.033	0.067	0.040	0.028	0.034	0.026	0.034	0.026
			MSE	0.019	0.019	0.026	0.022	0.013	0.020	0.014	0.020	0.014
		1	Bias	0.019	0.021	0.042	0.022	0.018	0.021	0.017	0.021	0.017
			MSE	0.012	0.012	0.015	0.013	0.010	0.012	0.010	0.012	0.010
40	20	2	Bias	0.022	0.024	0.047	0.026	0.021	0.024	0.020	0.024	0.020
			MSE	0.011	0.011	0.014	0.011	0.009	0.011	0.009	0.011	0.009
		3	Bias	0.025	0.027	0.049	0.029	0.023	0.027	0.023	0.027	0.023
			MSE	0.013	0.013	0.017	0.014	0.011	0.014	0.011	0.014	0.011
		4	Bias	0.035	0.037	0.060	0.038	0.032	0.037	0.032	0.037	0.032
			MSE	0.013	0.014	0.017	0.014	0.011	0.014	0.011	0.014	0.011
40	28	1	Bias	0.017	0.018	0.033	0.021	0.018	0.018	0.016	0.018	0.016
			MSE	0.007	0.007	0.008	0.007	0.006	0.007	0.006	0.007	0.006
		2	Bias	0.019	0.021	0.036	0.023	0.020	0.020	0.018	0.021	0.018
			MSE	0.008	0.009	0.010	0.009	0.007	0.009	0.007	0.009	0.007
		3	Bias	0.012	0.013	0.028	0.015	0.013	0.013	0.011	0.013	0.011
			MSE	0.007	0.007	0.009	0.008	0.006	0.007	0.006	0.007	0.006
60	20	4	Bias	0.018	0.019	0.035	0.022	0.019	0.019	0.017	0.019	0.017
			MSE	0.008	0.008	0.009	0.008	0.007	0.008	0.007	0.008	0.007
		1	Bias	0.010	0.011	0.021	0.012	0.011	0.011	0.010	0.011	0.010
			MSE	0.005	0.005	0.006	0.005	0.005	0.005	0.005	0.005	0.005
		2	Bias	0.011	0.012	0.022	0.012	0.011	0.012	0.011	0.012	0.011
			MSE	0.005	0.005	0.006	0.005	0.005	0.005	0.005	0.005	0.005
60	28	3	Bias	0.010	0.011	0.021	0.012	0.011	0.011	0.010	0.011	0.010
			MSE	0.005	0.005	0.006	0.005	0.005	0.005	0.005	0.005	0.005
		4	Bias	0.012	0.013	0.023	0.014	0.012	0.013	0.012	0.013	0.012
			MSE	0.005	0.005	0.006	0.006	0.005	0.005	0.005	0.005	0.005
		1	Bias	0.015	0.016	0.031	0.023	0.020	0.016	0.015	0.016	0.014
			MSE	0.007	0.007	0.008	0.008	0.006	0.007	0.006	0.007	0.006
60	20	2	Bias	0.016	0.017	0.032	0.024	0.020	0.017	0.015	0.017	0.015
			MSE	0.008	0.008	0.009	0.009	0.007	0.008	0.007	0.008	0.007
		3	Bias	0.017	0.018	0.033	0.025	0.021	0.018	0.016	0.018	0.016
			MSE	0.007	0.007	0.008	0.008	0.006	0.007	0.006	0.007	0.006
		4	Bias	0.015	0.016	0.031	0.023	0.020	0.016	0.014	0.016	0.015
			MSE	0.007	0.007	0.008	0.008	0.006	0.007	0.006	0.007	0.006

Scheme 2: $R_1 = n - m$ and $R_2 = \dots = R_m = 0$.

Scheme 3: $R_1 = R_m = \frac{n-m}{2}$ and $R_2 = \dots = R_{m-1} = 0$

Scheme 4: $R_{\frac{m}{2}} = n - m$ and $R_i = 0$ for $i \neq \frac{m}{2}$.

For given m, n and $(R_1, \dots, R_m)$, we have generated progressively censored sample of

Table 7: Simulated average interval lengths (AILs) and coverage probabilities (CPs) of 95% confidence intervals of $\theta = 0.5$.

n	m	schme		ACI	ACT*	LRT-CIT	ECI	Boot-p	Boot-t	HPD1	HPD2
10	4	1	AIL	0.728	0.782	0.659	0.645	0.831	0.832	0.648	0.573
			CP	0.976	0.981	0.951	0.954	0.916	0.910	0.958	0.969
		2	AIL	0.797	0.862	0.735	0.719	0.929	0.929	0.724	0.618
			CP	0.960	0.920	0.890	0.910	0.864	0.869	0.915	0.944
		3	AIL	0.728	0.785	0.665	0.651	0.836	0.838	0.653	0.577
			CP	0.976	0.976	0.964	0.969	0.932	0.931	0.959	0.973
	20	4	AIL	0.751	0.809	0.690	0.675	0.869	0.874	0.679	0.592
			CP	0.969	0.955	0.933	0.935	0.901	0.904	0.936	0.965
		1	AIL	0.549	0.573	0.477	0.471	0.545	0.547	0.471	0.444
			CP	0.984	0.976	0.940	0.943	0.919	0.921	0.941	0.957
		2	AIL	0.547	0.572	0.499	0.494	0.567	0.569	0.494	0.459
			CP	0.958	0.965	0.938	0.933	0.921	0.922	0.934	0.951
20	10	3	AIL	0.538	0.562	0.477	0.471	0.543	0.543	0.472	0.443
			CP	0.969	0.974	0.944	0.946	0.928	0.927	0.943	0.955
		4	AIL	0.542	0.566	0.485	0.479	0.553	0.550	0.480	0.449
			CP	0.974	0.981	0.955	0.956	0.929	0.929	0.954	0.966
	14	1	AIL	0.428	0.440	0.401	0.397	0.438	0.438	0.398	0.380
			CP	0.966	0.957	0.949	0.943	0.938	0.938	0.946	0.954
		2	AIL	0.427	0.439	0.413	0.410	0.451	0.451	0.409	0.391
			CP	0.966	0.965	0.957	0.952	0.935	0.937	0.955	0.968
		3	AIL	0.431	0.443	0.409	0.405	0.446	0.446	0.406	0.387
			CP	0.966	0.959	0.945	0.950	0.921	0.926	0.951	0.959
40	20	4	AIL	0.426	0.438	0.406	0.403	0.444	0.444	0.403	0.385
			CP	0.967	0.964	0.950	0.952	0.926	0.923	0.951	0.957
		1	AIL	0.370	0.378	0.323	0.322	0.345	0.344	0.322	0.313
			CP	0.979	0.969	0.947	0.949	0.920	0.920	0.945	0.951
		2	AIL	0.369	0.376	0.341	0.339	0.361	0.362	0.338	0.328
			CP	0.971	0.974	0.951	0.953	0.940	0.939	0.951	0.955
	28	3	AIL	0.369	0.376	0.328	0.326	0.349	0.349	0.326	0.317
			CP	0.976	0.975	0.953	0.957	0.933	0.936	0.952	0.958
		4	AIL	0.370	0.378	0.332	0.331	0.353	0.354	0.330	0.320
			CP	0.977	0.967	0.952	0.951	0.932	0.929	0.951	0.952
		1	AIL	0.294	0.298	0.275	0.274	0.287	0.288	0.274	0.268
			CP	0.967	0.967	0.950	0.953	0.946	0.948	0.954	0.957
60	20	2	AIL	0.294	0.298	0.285	0.284	0.296	0.296	0.283	0.277
			CP	0.951	0.949	0.944	0.947	0.939	0.938	0.945	0.949
		3	AIL	0.294	0.298	0.279	0.278	0.291	0.291	0.277	0.272
			CP	0.975	0.971	0.962	0.963	0.948	0.944	0.960	0.964
		4	AIL	0.294	0.298	0.280	0.279	0.291	0.292	0.278	0.272
			CP	0.965	0.970	0.959	0.958	0.952	0.951	0.953	0.961
	28	1	AIL	0.399	0.409	0.312	0.310	0.333	0.332	0.311	0.302
			CP	0.982	0.991	0.938	0.943	0.924	0.929	0.939	0.945
		2	AIL	0.403	0.414	0.340	0.338	0.361	0.363	0.337	0.327
			CP	0.981	0.988	0.958	0.957	0.934	0.936	0.956	0.961
		3	AIL	0.401	0.411	0.320	0.318	0.340	0.342	0.318	0.309
			CP	0.987	0.994	0.952	0.952	0.944	0.948	0.951	0.956
80	20	4	AIL	0.409	0.420	0.330	0.329	0.351	0.357	0.329	0.318
			CP	0.978	0.989	0.931	0.927	0.919	0.915	0.933	0.935

size m from XLindley distribution with $\theta = 0.5$. For the Bayesian estimators, two prior distributions are used.

Prior 1: We considered priors very close to the non-informative prior for θ , $\alpha =$

$\beta = 0.0001$.

Prior 2: Informative prior distributions for the parameter such that the prior mean of the parameter is equal to the corresponding true value and the prior variance is equal to 0.125. In this case, for the combination of $\theta = 0.5$ parameter, $\alpha = 2$ and $\beta = 4$ values are obtained.

The Geweke test Geweke (1992), the Raftery and Lewis's diagnostic (Raftery and Lewis, 1992, 1996), the Heidelberger and Welch's convergence diagnostic Heidelberger and Welch (1983), and the Gelman-Rubin test Gelman et al. (1992) are used to check the convergence of the generated MCMC. To check the convergence, we consider a simulation case, say ($n = 40, m = 28, \theta = 0.5$, under censoring scheme 1) and iterate the Metropolis-Hastings algorithm 1000 times. In this iteration, the average acceptance rate of the Metropolis-Hastings algorithm is 68.2%, and the average effective sample size is 31,200 out of the 45,000 retained samples. The average Gelman-Rubin statistic is calculated to be 1.002, indicating that the inter-chain variance is negligible compared to the intra-chain variance and that the chains converge to a common stationary distribution regardless of the initial values. In addition, the average p -value of the Geweke test is 0.49 and the average Heidelberger-Welch p -value is 0.28, both confirming no significant evidence of non-convergence, while the average lag-1 autocorrelation of the retained draws is 0.15, indicating low dependence between successive samples and efficient mixing of the chain. These quantitative results provide strong empirical evidence that the MCMC sampler performs uniformly well on repeatedly simulated datasets. Figure 6 shows the MCMC (the figure is for $n = 40, m = 28, \theta = 0.5, \alpha = \beta = 0.0001$ and $\alpha = 2$ and $\beta = 4$, under censoring scheme 1), from which the convergence of the Metropolis-Hastings algorithm may be confirmed.

The MLEs, PBEs, MBEs, BESs, BETKs and BELs are computed using the methods described in Section 2. The approximate Bayes estimates under the SEL functions are computed for θ using Lindley's approximation, MCMC and TK. We compare the performances of the MLEs, PBEs, MBEs, BESs, BETKs and BELs in terms of biases, and mean square errors (MSEs), which can be estimated as

$$\widehat{Bias} = \frac{1}{1000} \sum_{i=1}^{1000} (\hat{\theta}_i - \theta), \quad \text{and} \quad \widehat{MSE} = \frac{1}{1000} \sum_{i=1}^{1000} (\hat{\theta}_i - \theta)^2,$$

respectively, where $\hat{\theta}_i$ is the estimate of θ obtained in the i -th simulation for $i = 1, \dots, 1000$. The simulation results are reported in Table 6.

The performance of different confidence intervals including ACI, ACI^* , ECI, Boot-p, Boot-t and HPD is evaluated based on their average interval lengths (AILs) and simulated coverage probabilities (CPs). If $(L_i; U_i)$ is the 95% confidence interval of θ obtained in the i -th simulation, then the average interval lengths (AILs) and the coverage probabilities (CPs) are computed as

$$\widehat{AILs} = \frac{1}{1000} \sum_{i=1}^{1000} (U_i - L_i), \quad \text{and} \quad \widehat{CPs} = \frac{1}{1000} \sum_{i=1}^{1000} I(L_i < \theta < U_i),$$

where $I(\cdot)$ denotes the indicator function. The simulation results are reported in Table 7.

On the prediction front, the performances of the point predictors including MLP, BUP, CMP and BPs are compared based on their biases and mean squared prediction errors (MSPEs). For interval prediction, we compute the 95% interval predictions for Z based on the pivotal quantity method, the HCD Method and the Bayesian interval predictor. The predictive intervals are also compared by means of their estimated ALs and CPs. The results are reported in Tables 8 and 9.

Table 8: Simulated biases and MSPEs of different point predictors.

(n, m)	Scheme	j	k		Classic predictions			Bayesian predictions	
					MLP	BUP	CMP	Prior 1	Prior 2
(20, 10)	$(0^{*9}, 10)$	10	1	Bias	-0.2818	-0.0026	-0.0862	0.0200	0.0187
				MSPE	0.1567	0.0834	0.0880	0.0845	0.0836
		2		Bias	-0.3007	0.0093	-0.0810	0.0553	0.0521
				MSPE	0.2766	0.2086	0.2105	0.2180	0.2103
		3		Bias	-0.3521	-0.0099	-0.1043	0.0598	0.0552
				MSPE	0.4670	0.3812	0.3828	0.3944	0.3813
		4		Bias	-0.4108	-0.0307	-0.1303	0.0691	0.0571
				MSPE	0.7152	0.5968	0.6040	0.6237	0.6036
		5		Bias	-0.4213	0.0049	-0.1020	0.1261	0.1216
				MSPE	0.8904	0.7904	0.7993	0.8358	0.7924
(40, 28)	$(0^{*27}, 12)$	10	1	Bias	-0.2177	0.0028	-0.0640	0.0087	0.0087
				MSPE	0.0904	0.0448	0.0481	0.0451	0.0450
		2		Bias	-0.2148	0.0182	-0.0538	0.0301	0.0304
				MSPE	0.1433	0.1016	0.1027	0.1021	0.1018
		3		Bias	-0.2299	0.0167	-0.0587	0.0362	0.0351
				MSPE	0.2265	0.1805	0.1814	0.1824	0.1810
		10	1	Bias	-0.0958	-0.0013	-0.0300	0.0013	0.0011
				MSPE	0.0188	0.0098	0.0106	0.0100	0.0099
(60, 30)	$(0^{*29}, 30)$	2		Bias	-0.1001	-0.0028	-0.0330	0.0022	0.0020
				MSPE	0.0308	0.0212	0.0221	0.0213	0.0212
		3		Bias	-0.1016	-0.0010	-0.0318	0.0064	0.0063
				MSPE	0.0378	0.0285	0.0292	0.0287	0.0286

For point estimation, from Table 6, we observe that the MLE works better than the MBE. The PBE does not perform well and gives the largest MSE in all cases considered. Comparing the Bayesian estimators obtained from Lindley approximation methods, TK approximation and MCMC, we observe that MCMC method based on the non-informative prior (i.e., Prior 1) is better than Lindley approximation method in terms of both biases and MSEs. In addition, the Bayesian estimators obtained by TK and the MCMC perform similarly in terms of biases and MSEs. Also, the BEs based on the informative prior perform better than the BEs based on the non-informative prior. From Table 6, we see that the Bayesian point estimates under the informative prior perform better than the MLE, MBE and PBE estimates.

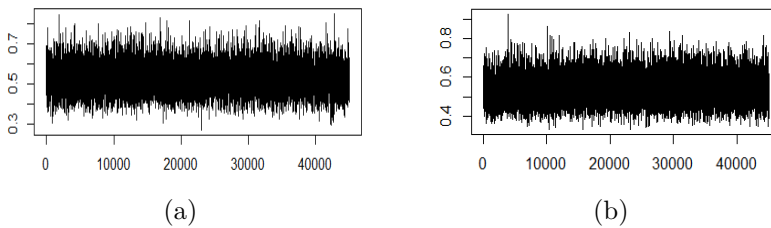
For interval estimation, from Table 7, we observe that all the interval estimation procedures considered here, as the sample size m increases, the simulated average width of the intervals decreases. The average interval lengths of the 95% confidence intervals based on the ECIs are smaller than the corresponding average interval lengths of the ACIs based on the asymptotic normality of the MLE, confidence intervals based on LRT, bootstrap confidence intervals (Boot-p and Boot-t). The bootstrap method generally performs poorly, producing the longest intervals in most cases considered. Additionally, ACIs based on the asymptotic normality of log MLE are larger than

Table 9: Simulated ALs and CPs of different 95% prediction intervals.

(n, m)	Scheme	j	k		BPI			
					Pivotal	HCD	Prior 1	Prior 2
(20, 10)	(0 [*] 9, 10)	1		AIL	0.9899	0.8137	1.1478	1.1368
				CP	0.933	0.917	0.95	0.951
		2		AIL	1.5284	1.3658	1.8413	1.8139
				CP	0.919	0.913	0.947	0.948
		3		AIL	1.9651	1.839	2.4354	2.3962
				CP	0.910	0.899	0.954	0.952
	(0 [*] 9, 10)	4		AIL	2.4327	2.3434	3.0939	3.0349
				CP	0.918	0.916	0.951	0.956
		5		AIL	2.9616	2.9245	3.8555	3.7692
				CP	0.874	0.874	0.936	0.941
(20, 10)	(0 [*] 9, 10)	1		AIL	0.8009	0.6568	0.8432	0.8413
				CP	0.944	0.953	0.950	0.949
		2		AIL	1.2124	1.0794	1.2946	1.2911
				CP	0.951	0.946	0.958	0.957
		3		AIL	1.5559	1.4483	1.6811	1.6777
				CP	0.930	0.928	0.941	0.942
	(0 [*] 9, 10)	1		AIL	0.3430	0.2811	0.3605	0.3604
				CP	0.938	0.932	0.945	0.945
		2		AIL	0.5106	0.4539	0.5455	0.5430
				CP	0.947	0.936	0.957	0.952
		3		AIL	0.6366	0.5896	0.6879	0.6870
				CP	0.927	0.926	0.940	0.942

those based on the MLE. Bayesian intervals using the informative prior are shorter than those using the non-informative prior. From Table 7, we see that Bayesian confidence intervals, HPD, under the informative prior perform better than the other confidence intervals proposed in this paper.

For point prediction, from Table 8, we observe that BUP performs better than the MLP and CMP in terms of biases and MSPEs. For fixed value of m and n , the MSPEs of all the point predictors are increasing with k as expected. Comparing the Bayesian point predictions based on different priors, the Bayesian point predictors based on informative priors (i.e., Prior 2) perform better than the Bayesian point predictors based on the non-informative prior (i.e., Prior 1), in terms of biases and MSPEs. We also observe that the MLP does not perform well because it provides the largest biases and MSPEs among all the point predictors considered here. We also observe that the BUP performs better than the Bayesian point predictors in terms of estimated MSPEs.

Figure 6: Plot of MCMC for (a) $\theta = 0.5$ and $\alpha = \beta = 0.0001$, (b) $\theta = 0.5$ and $\alpha = 2$ and $\beta = 0.0001$.

For prediction intervals, we observe from Table 9 that the simulated coverage probabilities are close to the nominal level 95% in most cases. For fixed values of n and m , the average interval lengths of different prediction intervals increase as k increases. It can be seen that Bayesian prediction intervals are wider than the prediction intervals obtained by the pivotal quantity method and the HCD method. Among the pivotal quantity method and the HCD method, the simulated average widths of prediction intervals based on HCD method are smaller.

8 Conclusions

In this paper, we have considered the problem of estimation and prediction for the XLindley distribution under progressive Type-II censoring from classical and Bayesian viewpoints. We derived the MLE, MBEs, PBEs, and Bayes estimates for the unknown parameter of an XLindley distribution. Also, confidence intervals based on the asymptotic distribution of the MLE, confidence intervals based on the asymptotic distribution of the MLE logarithm, ECIs, bootstrap confidence intervals (Boot-p and Boot-t methods), and Bayesian confidence intervals have been calculated. Based on the simulation results, the MLE method is recommended when prior information about the parameter is not available, and the Bayesian estimation approach is recommended when accurate prior information about the parameter is available. In the discussion of interval estimation, Bayesian intervals (with informative prior) and ECIs also performed more favorably. In practice, the ECI requires only univariate root-finding and does not depend on MCMC sampling or specialized Bayesian software, which makes it a readily accessible option for practitioners without programming or MCMC expertise. Various predictors, including the MLP, BUP, CMP, and the Bayesian point predictor, have been calculated for the failure time. We have also presented prediction intervals for the future failure times. Based on the simulation results, it is shown that the BUP performs well when compared to other point predictors in terms of prediction bias and MSPEs. It is also observed that the HCD prediction interval behaves well based on the average width criterion. These comparative conclusions are not tied to the particular informative prior used here; a fuller sensitivity analysis to prior misspecification, and application of the methods to further real censored data sets, are left for future work. Another area for future work could involve studying stress-strength reliability with Type-II censored data. Additionally, linear inference based on progressive censoring, Type-I right censoring, or record data could be developed for the XLindley distribution. Since the MBE (and its corresponding confidence interval) and MLE involve solving some nonlinear equations, the function `uniroot` in the statistical software R (R Core Team, 2026) was used to solve the corresponding nonlinear equations and calculate these estimates. All the computations in the paper were done using the statistical software R (R Core Team, 2026) and the packages `MHadaptive` (Chivers, 2012), `nleqslv` (Hasselman, 2018), and `rootSolve` (Soetaert, 2009).

References

- Alotaibi, R., Nassar, M. and Elshahhat, A. (2022). Computational analysis of XLindley parameters using adaptive Type-II progressive hybrid censoring with applications in chemical engineering. *Mathematics*, **10**(18):3325.
- Alotaibi, R., Nassar, M. and Elshahhat, A. (2023). Reliability estimation under normal operating conditions for progressively type-II XLindley censored data. *Axioms*, **12**(4):352.
- Alotaibi, R., Nassar, M. and Elshahhat, A. (2024). Data analysis with applications in physics and engineering using XLindley model with improved adaptive Type-II progressively censored samples. *Heliyon*, **10**(14):e33598.
- Balakrishnan, N. and Aggarwala, R. (2000). *Progressive Censoring: Theory Methods and Applications*, Boston: Birkhäuser.
- Balakrishnan, N. and Cramer, E. (2014). *Art of Progressive Censoring: Applications to Reliability and Quality*. New York: Birkhauser.
- Casella, G. and Berger, R.L. (2002). *Statistical Inference*. 2nd ed., CRC Press..
- Chen, M.H. and Shao, Q.M. (1999). Monte Carlo estimation of Bayesian credible and HPD intervals. *Journal of computational and Graphical Statistics*, **8**(1):69–92.
- Chivers, C. (2012). MHadaptive: General Markov Chain Monte Carlo for Bayesian Inference using adaptive Metropolis-Hastings sampling. R package.
- Chouia, S. and Zeghdoudi, H. (2021). The XLindley distribution: properties and application. *Journal of Statistical Theory and Applications*, **20**(2):318–327.
- Efron, B. (1979), Bootstrap methods: Another look at the jackknife, *The Annals of Statistics*, **7**(1):1–26.
- Elbatal, I., Elshahhat, A., Semary, H.E. and Nassar, M. (2025). Reliability assessment of two production lines using joint progressively type-II censored XLindley samples. *Scientific Reports*, **15**(1):26841.
- Elshahhat, A. and Alotaibi, R. (2026). The XLindley survival model under generalized progressively censored data: theory, inference, and applications. *Axioms*, **15**(1):56.
- Gelman, A. and Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, **7**(4):457–472.
- Geweke, J. (1992). Evaluating the accuracy of sampling-based approaches to the calculations of posterior moments. *Bayesian Statistics*, **4**:641–649.
- Hasselmann, B. (2018). nleqslv: solve systems of nonlinear equations, R package version 3.3.2.
- Hastings, W.K. (1970). Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, **57**:97–109.

- Heidelberger, P. and Welch, P.D. (1983). Simulation run length control in the presence of an initial transient. *Operations Research*, **31**(6):1109–1144.
- Kaminsky, K.S. and Rhodin, L.S. (1985). Maximum likelihood prediction, *Annals of the Institute of Statistical Mathematics*, **37**(Part A):507–517.
- Lindley, D.V. (1980). Approximate Bayesian methods, *Trabajos de estadística y de investigación operativa*, **31**(1):223–237.
- Meeker, W.Q. and Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, New York: Wiley.
- Metropolis, N., Rosenbluth, A.W., Rosenbluth, M.N., Teller, A.H. and Teller, E. (1953). Equation of state calculations by fast computing machines. *The Journal of Chemical Physics*, **21**(6):1087–1092.
- Nassar, M., Alotaibi, R. and Elshahhat, A. (2023). Reliability estimation of XLindley constant-stress partially accelerated life tests using progressively censored samples. *Mathematics*, **11**(6):1331.
- Raftery, A.E. and Lewis, S.M. (1992). Comment: One long run with diagnostics: Implementation strategies for Markov chain Monte Carlo. *Statistical Science*, **7**(4):493–497.
- Raftery, A.E. and Lewis, S.M. (1996). Implementing MCMC. In *Markov Chain Monte Carlo in Practice*, Gilks, W.R., Richardson, S. and Spiegelhalter, D.J., pp. 115–130, Boca Raton: Chapman and Hall/CRC.
- Raqab, M.Z. and Nagaraja, H.N. (1995). On some predictors of future order statistics, *Metron*, **53**(12):185–204.
- R Core Team (2026). *R: A Language and Environment for Statistical Computing*. Austria, Vienna: R Foundation for Statistical Computing. <https://www.R-project.org/>
- Robert, C. and Casella, G. (2004). *Monte Carlo Statistical Methods*, 2nd edition, New York: Springer.
- Soetaert, K. (2009). *rootSolve: Nonlinear root finding, equilibrium and steady-state analysis of ordinary differential equations*. R-package version 1.8.2.
- Stablein, D.M., Carter, J. and Novak, J.W. (1981). Analysis of survival data with nonproportional hazard functions. *Controlled Clinical Trials*, **2**(2):149–159.
- Thongjub, N. and Panichkitkosolkul, W. (2026). Improved inference for the XLindley distribution: A comparative analysis of confidence intervals and applications. *Statistics Optimization and Information Computing*, **15**(6):4900–4921.
- Tierney, L., and Kadane, J.B. (1986). Accurate approximations for posterior moments and marginal densities. *Journal of the American Statistical Association*, **81**(393):82–86.
- Zanjiran, F. and MirMostafaei, S.M.T.K. (2024). On estimation and prediction for the XLindley distribution based on record data. *Reliability: Theory and Applications*, **19**(2 (78)):258–272.