

Inference for Block Type-I Hybrid Censored Weibull Populations

Farha Sultana^a, Çağatay Çetinkaya^b, Mohammad Z. Raqab^{c,d} and Debasis Kundu^e

^aDepartment of Science and Mathematics, Indian Institute of Information Technology, Guwahati-781015, India; ^bDepartment of Actuarial Sciences, Kırıkkale University, 71450, Turkey; ^cDepartment of Statistics and Operations Research, Kuwait University, Al-Shadadiyya, Kuwait; ^dDepartment of Mathematics, The University of Jordan, Amman 11942, Jordan; ^eDepartment of Mathematics and Statistics, Indian Institute of Technology Kanpur, 208016, India

ARTICLE HISTORY

Compiled October 9, 2025

ABSTRACT

Experimenter assigns a prespecified number of units using different test facilities with the help of block censoring, which helps to improve the efficiency of test. In this paper, estimation of unknown model parameters in the context of life testing across various test facilities using block Type-I hybrid censoring has been considered. We specifically focus on products whose lifetimes follow the Weibull distribution, characterized by identical shape parameters but distinct scale parameters. The existence and uniqueness of maximum likelihood estimators for model parameters are established. Additionally, we evaluate the Bayes estimators of the unknown model parameters through a Metropolis-Hastings sampling technique within a hierarchical framework. We assess order-restricted maximum likelihood estimators for both the classical and Bayesian setups. Confidence intervals, including asymptotic and highest posterior density intervals, are also derived for all the parameters. Further, we compute the optimal hybrid censoring time points based on various optimality criteria. Finally, we conduct a comprehensive simulation study to compare the performance of the proposed methodologies, followed by the presentation of a real-world data example to show its real life application.

KEYWORDS

Block censoring; Type-I hybrid censoring; Maximum likelihood estimation; Order restricted inference; Optimal censoring time point; Weibull distribution

1. Introduction

In most reliability and life testing experiments, complete observations may not be observed. Some limitations on time, cost, labor, systemic or etc. cause censored samples. In these cases, censoring schemes are considered according to the types of missing values. The main censoring schemes are classified into two types namely Type-I censoring and Type-II censoring. Epstein (1954) introduced the Type-I hybrid censoring scheme which is a mixture of Type-I and Type II censoring schemes. In this censoring scheme (CS), the experiment starts first with n independent and identical units. In the scenario where both the experimental time T and r , the predetermined number of items to be failed, are fixed. The termination of the experiment occurs when the first of two events

happens. The experiment is terminated at the random time $T^* = \min\{X_{r:n}, T\}$, where $X_{r:n}$ is the life time of the r^{th} item. Certainly, certain life-testing experiments may lead to extended failure times among units, leading to unexpectedly prolonged test duration. Consequently, the classical hybrid Type-I experimenting time could still be too lengthy in such scenarios. To address this issue, the block censoring scheme (BCS) was introduced by Ahmadi, Doostparast, and Ahmadi (2018), offering a more efficient and improved alternative for reducing test duration in life-testing experiments.

The BCS procedure starts with n identical units stochastically divided into m sets where the i th set ($i = 1, 2, \dots, m$) contains n_i units such that $\sum_{i=1}^m n_i = n$. In this scheme m different facilities may be used to run the test simultaneously. Then, the test is terminated at $T_i^* = \min\{X_{is_i}, T_i\}$ for ($i = 1, 2, \dots, m$). If s_i denotes the number of units failed in the i^{th} facility for ($i = 1, 2, \dots, m$), then $\sum_{i=1}^m s_i = s$ denotes the total number of units failed in all the facilities. Then, entire system is terminated at $\max\{T_i^*\}$ for $i = 1, 2, \dots, m$. The BCS enables experimenters to simultaneously run a Type-I hybrid censoring scheme across all sets. This scheme allows for the removal of live units during testing, which can then be utilized for other purposes.

On the other hand, Ahmadi et al. (2018) proposed the BSC as an extension of the first-failure censoring schemes (FFCS). The FFCS is introduced by Sarhan and Greenberg (1962) and Johnson (1964). In FFCS, there are m facilities and $N = m \times n$ test units. These N units are randomly allocated to m sets each containing n units and then each set is tested on one of the m facilities. When the first failure on each facility is observed, the life-test is terminated. Wu, Tsai, and Ouyang (2001) proposed the limited failure censoring scheme (LFCS) as an extension of FFCS. In this CS, m life-tests are performed sequentially in m facilities and each test is terminated when the s_i^{th} failure is observed ($i = 1, \dots, m$). Then, the system termination time of LFCS is given as $\sum_{i=1}^m X_{s_i:n_i}^{(i)}$ where $X_{s_i:n_i}^{(i)}$ denotes the s_i^{th} order statistics in each test set with sizes n_i . Ahmadi et al. (2018) proposed Type-II BSC to provide another extension of FFCS. In this model, m sets where the i^{th} set ($i = 1, \dots, m$) contains n_i units and there exists m facilities to run the experiment for the m sets. By performing m sets simultaneously, s_i failures are observed in the each facility and then totally S failures in all of the facilities ($\sum_{i=1}^m s_i = S$) are obtained. This paper is an extension of both FFCS and BSC. Both FFCS and BSC can be obtained as special cases of the present CS. This model contains m sets where each of them contains n_i units. In this model, we propose to use predetermined time point, T_i , for each set. Thus, each set terminate at the time $T_i^* = \min\{X_{is_i}, T_i\}$. The flexibility of this hybrid censoring model is quite useful when the failure times of the units are quite high. We call this CS as Type-I hybrid BSC. The expected experimental time and the number of failure number in this model is approximately same with the classical Type-I hybrid censoring scheme when the all components are tested completely jointly. However, the BSC allows the experimenters a simultaneously experimental process for different facilities. Further, by using this BCS scheme, we can utilize m different facilities effectively.

The BCS scheme becomes popular in the literature because of its usefulness over normal CS. In previous studies, Zhu (2020) addressed the BCS method in the context of the traditional Type-II censoring scheme on Weibull populations. Wang, Wu, Lin, and Tripathi (2023) further explored the concept of block censoring for progressive censored competing risk data, focusing on inverted exponentiated exponential distributions. Kumari, Tripathi, Sinha, and Wang (2023) studied block progressive censoring for reliability estimation of bathtub-shaped distribution. Singh, Tripathi, Lodhi, and Wang (2024) studied the block progressive type-II censoring scheme under the unit

inverse Weibull distribution and Singh, Tripathi, Wang, and Wu (2024) handled the block adaptive Type-II progressive hybrid censoring scheme for the Weibull distribution. In addition to these previous studies, our aim is to focus on implementing the BCS scheme in each of the m sets, employing a Type-I hybrid censoring approach that has been used quite extensively in reliability acceptance test in MIL-STD-781-C (US Department of Defense, 1977). The block Type-I hybrid censoring scheme proposes an expansion of the Type-II censoring scheme handled by the various authors since this hybrid censoring scheme is a mixture of Type-I and Type-II censoring ideas. In this censoring scheme, each set in the system terminated at $T_i^* = \min\{X_{is_i}, T_i\}$ and the experiment ends at $\max\{T_i^*\}$ for $i = 1, 2, \dots, m$ sets. Another aim of this study is to analyze and determine the optimal censoring scheme for given parameters. The optimal censoring plan has never been handled before within block censoring studies. This paper aims to help to experimenters to guide to construct optimal censoring plan to save cost, time, labor force, and many more. This paper aims to address that issue.

The paper is outlined as follows. In Section 2, we provide a detailed description of the model and the process of maximum likelihood estimation (MLE). This section also includes discussions on inferences under order restrictions between scale parameters λ_i , $i = 1, 2, \dots, m$, along with the discussion of the existence and uniqueness of the MLEs. We present the corresponding approximate confidence intervals (CIs) for MLEs using the asymptotic normality of the MLEs. Section 3 is devoted to obtaining Bayesian estimations and their corresponding highest posterior density (HPD) credible intervals, achieved through using gamma-Dirichlet priors and a Monte Carlo Markov Chain (MCMC) with the Metropolis-Hastings (M-H) algorithm. The Bayesian estimation under order restriction is further discussed. Additionally, in Section 4, we establish an optimal censoring plan by considering the expected duration of the experiment and an information measure as proposed by Gupta and Kundu (2006). To validate the theoretical findings, we conduct simulation studies in Section 5. Furthermore, the proposed methods are illustrated through the analysis of a real-life data set on the breaking strengths of jute fibers in Section 6. Finally, we conclude with a summary of our findings in Section 7.

2. Model Description and MLEs

Suppose we have n identical units for the life testing experiment. These units are divided into m sets, where the i^{th} set ($i = 1, 2, \dots, m$) includes n_i units, and the total number of units across all sets is n , as seen in Figure 1.

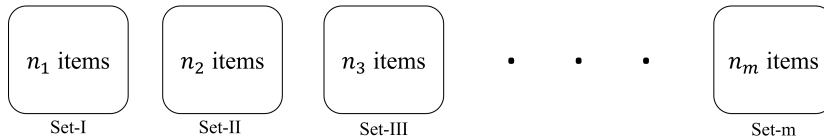


Figure 1. Schematic presentation of blocks.

Then, we consider the following censoring scheme:

- All items exist and perform life-testing experiments simultaneously in m sets.
- The sets in the experiment will be removed from the test at the random time $T_i^* = \min\{X_{is_i}, T_i\}$, where $s_i (< n_i)$ and T_i pre-determined number of failures

and experimental time, respectively.

- The test is terminated at $\max\{T_i^*\}$, $i = 1, 2, \dots, m$.
- Note that the BCS means here running Type-I hybrid censoring in all m sets simultaneously.
- Let X_{ij} denotes the j^{th} observed test unit from i^{th} set. Therefore, the data looks like

$$\begin{bmatrix} \text{Set} - 1 \\ \text{Set} - 2 \\ \vdots \\ \text{Set} - m \end{bmatrix} = \begin{bmatrix} X_{11} & X_{12} & X_{13} & \cdots & U_1 \\ X_{21} & X_{22} & X_{23} & \cdots & U_2 \\ \vdots & \vdots & \vdots & \cdots & \vdots \\ X_{m1} & X_{m2} & X_{m3} & \cdots & U_m \end{bmatrix}$$

where $U_i = \min(X_{is_i}, T_i)$.

When $T_i \rightarrow \infty$ and s_i for $i = 1, 2, \dots, m$, then BCS is particular case of FFCS introduced by Sarhan and Greenberg (1962) and Johnson (1964). Assume that $X_{ij} \sim \text{Weibull}(\alpha, \lambda_i)$, $i = 1, 2, \dots, m$, $j = 1, 2, \dots, n_i$, where α is the common shape parameter and λ_i different scale parameters. In particular for two sets ($i = 2$), $X_{11}, X_{12}, \dots, X_{1n_1} \sim \text{Weibull}(\alpha, \lambda_1)$ and $X_{21}, X_{22}, \dots, X_{2n_2} \sim \text{Weibull}(\alpha, \lambda_2)$. The following probability density (PDF) and cumulative distribution functions (CDF), are respectively, given by

$$g(x; \alpha, \lambda_i) = \alpha \lambda_i x^{\alpha-1} e^{-\lambda_i x^\alpha}, \quad G(x; \alpha, \lambda_i) = 1 - e^{-\lambda_i x^\alpha}. \quad (1)$$

Thus, the corresponding likelihood function of the Set-I can be obtained as

$$L(\alpha, \lambda_1, X_1) = \frac{n_1}{(n_1 - k_1)!} \prod_{j=1}^{k_1} f(x_{1j}) [1 - F(U_1)]^{n_1 - k_1}, \quad (2)$$

where

$$k_1 = \begin{cases} s_1, & \text{if } X_{1s_1} < T_1, \\ d_1, & \text{if } X_{1s_1} > T_1 \end{cases}, \quad U_1 = \begin{cases} x_{1s_1}, & \text{if } X_{1s_1} < T_1, \\ T_1, & \text{if } X_{1s_1} > T_1 \end{cases},$$

with d_1 being the number of failures before time point T_1 occurred, such as $x_{1d_1} < T_1 < x_{1d_1+1}$. A case with $k = 0$ would suggest no failures during the entire experiment, which is unrealistic in life-testing practice and falls outside the scope of our study. Thus, our MLE expressions are valid only for cases where $k \geq 1$, as intended in reliability analysis. For the Set-1, by changing PDF and CDF (1) in (2) we obtain the following likelihood function

$$\begin{aligned} L(\alpha, \lambda_1, X_1) &\propto \prod_{j=1}^{k_1} \alpha \lambda_1 x_{1j}^{\alpha-1} \exp \left\{ -\lambda_1 x_{1j}^\alpha - \lambda_1 (n_1 - k_1) U_1^\alpha \right\} \\ &\propto \alpha^{k_1} \lambda_1^{k_1} \exp \left\{ (\alpha - 1) \sum_{j=1}^{k_1} \ln(x_{1j}) - \lambda_1 \left(\sum_{j=1}^{k_1} x_{1j}^\alpha + (n_1 - k_1) U_1^\alpha \right) \right\}. \end{aligned} \quad (3)$$

For the Set- i , $i = 1, 2, \dots, m$ the likelihood function looks like

$$L(\alpha, \lambda_i, X_i) \propto \alpha^{k_i} \lambda_i^{k_i} \exp \left\{ (\alpha - 1) \sum_{j=1}^{k_i} \ln(x_{ij}) - \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha \right) \right\} \quad (4)$$

where

$$k_i = \begin{cases} s_i, & \text{if } X_{is_i} < T_i \\ d_i, & \text{if } X_{is_i} > T_i \end{cases} \quad \text{and} \quad U_i = \begin{cases} x_{is_i}, & \text{if } X_{is_i} < T_i, \\ T_i, & \text{if } X_{is_i} > T_i \end{cases}.$$

Let $\mathbf{X} = (X_1, X_2, \dots, X_m)$ and $\boldsymbol{\lambda} = (\lambda_1, \lambda_2, \dots, \lambda_m)$ be the available BCS data and corresponding parameter vector. The likelihood function for the BCS can be written by generalizing (4) for m sets as given in the following:

$$\begin{aligned} L(\alpha, \boldsymbol{\lambda}, \mathbf{X}) &\propto \prod_{i=1}^m L(\alpha, \lambda_i, X_i) \propto \alpha^{\sum_{i=1}^m k_i} \exp \left\{ \sum_{i=1}^m k_i \ln(\lambda_i) \right\} \\ &\times \exp \left\{ (\alpha - 1) \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) - \sum_{i=1}^m \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha \right) \right\}. \end{aligned} \quad (5)$$

Then, the log-likelihood function of the observed sample is obtained as

$$\begin{aligned} \ell(\alpha, \boldsymbol{\lambda}, \mathbf{X}) &\propto \sum_{i=1}^m k_i \ln(\alpha) - \sum_{i=1}^m \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha \right) \\ &+ \sum_{i=1}^m k_i \ln(\lambda_i) + (\alpha - 1) \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) \end{aligned} \quad (6)$$

The MLEs of the parameters can be obtained by differentiating partially the log-likelihood function $\ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})$ with respect to α and λ_i , $i = 1, 2, \dots, m$, as given in the following:

$$\frac{\partial \ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})}{\partial \alpha} = \frac{\sum_{i=1}^m k_i}{\alpha} + \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) - \sum_{i=1}^m \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha \ln x_{ij} + (n_i - k_i) U_i^\alpha \ln U_i \right),$$

and

$$\frac{\partial \ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})}{\partial \lambda_i} = \frac{k_i}{\lambda_i} - \sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha.$$

Theorem 2.1. Let X_{ij} be the j^{th} observed test unit from i^{th} set for $i = 1, 2, \dots, m$ and $j = 1, 2, \dots, n_i$. If it is assumed that X_i 's, are i.i.d. with the Weibull(α, λ_i) under the BCS, the MLEs of λ_i , denoted by $\hat{\lambda}_i$, exist and given by,

$$\hat{\lambda}_i = \frac{k_i}{\sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha}.$$

Theorem 2.2. Let X_{ij} be the j^{th} observed test unit from i^{th} set for $i = 1, 2, \dots, m$ and $j = 1, 2, \dots, n_i$. If it is assumed that X_i 's, are i.i.d. by the Weibull(α, λ_i) under the BSC, the MLEs of α , denoted by $\hat{\alpha}$ exist and can be obtained with the solution of the following non-linear equation

$$\Omega(\alpha) = \frac{\sum_{i=1}^m k_i}{\alpha} + \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) - \sum_{i=1}^m \hat{\lambda}_i \left(\sum_{j=1}^{k_i} x_{ij}^{\alpha} \ln x_{ij} + (n_i - k_i) U_i^{\alpha} \ln U_i \right) = 0. \quad (7)$$

The non-linear equation (7) requires an iterative method, such as the Newton-Raphson method, since an analytical solution is not possible. The MLE of the α is an approximation for solving this equation. The Appendix section contains the proofs for the existence and uniqueness of $\hat{\alpha}$ and $\hat{\lambda}_i$. Recognizing that there are situations in which the assumption of equal shape parameters is not appropriate is crucial. One might inquire about scenarios involving different shape parameters. In such instances, the log-likelihood function, as presented in equation (6), is observed as:

$$\begin{aligned} \ell(\alpha, \boldsymbol{\lambda}, \mathbf{X}) &\propto \sum_{i=1}^m k_i \ln(\alpha_i) - \sum_{i=1}^m \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^{\alpha_i} + (n_i - k_i) U_i^{\alpha_i} \right) \\ &+ (\alpha_i - 1) \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) + \sum_{i=1}^m k_i \ln(\lambda_i). \end{aligned} \quad (8)$$

However, in this case, the iterative methods outlined in Theorem (2.2) must be repeated m times, leading to prolonged computational duration for the analysis. As a result, dealing with unequal shape parameters present significant computational challenges, particularly when the number of m is large. Hence, we propose considering an equal parameter case within the scope of this study.

2.1. Order-Restricted MLE

If α is fixed, the existence of $\hat{\lambda}_i$ specified in Theorem 2.1 may not be guaranteed for $i \in \tilde{I}^c$, where $I = \{1, 2, \dots, m\}$ and $\tilde{I} = \{i \in I \mid n_i \geq 1\}$. Considering order-restricted inference significantly reduce the probability of non-existence of unrestricted MLEs. Assuming the presence of at least one failure in every set, we impose the natural constraint on the scale parameters as follows:

$$\lambda_1 < \lambda_2 < \dots < \lambda_m.$$

Thus, the order-restricted MLEs of the unknown parameters can be obtained by maximizing (9) with respect to $\alpha, \lambda_1, \dots, \lambda_m$ based on the order restriction of $\lambda_1 < \lambda_2 < \dots < \lambda_m$.

Lemma 2.3. In the case of fixed shape parameter α , if $\ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})$ is maximized under the restriction on the scale parameters $\lambda_1 < \lambda_2 < \dots < \lambda_m$, the MLE of λ_u , denoted by $\tilde{\lambda}_u^{(\alpha)}$, for $k = 1, 2, \dots, m$ is given by

$$\tilde{\lambda}_u^{(\alpha)} = \left[\min_{s \leq u} \max_{t \geq u} \frac{\sum_{r=s}^t \Delta_r(\alpha)}{\sum_{r=s}^t k_r} \right]^{-1}, \quad s, t \in \tilde{I} = I,$$

where $\Delta_r(\alpha) = \sum_{j=1}^{k_r} x_{rj}^\alpha + (n_r - k_r)U_r^\alpha$.

Proof. See Appendix 7. □

Then, the profile log-likelihood function of α , $\omega(\alpha)$ is needed to maximize for obtaining the MLE of α , where

$$\omega(\alpha) = \frac{\sum_{i=1}^m k_i}{\alpha} + \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) - \sum_{i=1}^m \tilde{\lambda}_i^{(\alpha)} \left(\sum_{j=1}^{k_i} x_{ij}^\alpha \ln x_{ij} + (n_i - k_i)U_i^\alpha \ln U_i \right). \quad (9)$$

Proceeding similar to the classical MLE, the non-linear equation $\omega(\alpha)$ cannot be solved analytically and an iterative method such as the [Newton-Raphson](#) method is needed to solve such equations. The approximate solution of this non-linear equation be the MLE of the α parameter.

2.2. Asymptotic Confidence Interval

In this section, the asymptotic confidence intervals (ACIs) of the unknown parameters are constructed by using the observed Fisher information matrix. Suppose the parameter vector $\eta = (\alpha, \lambda_1, \dots, \lambda_m)^T$. The observed Fisher information matrix is given by

$$F = (f_{ij}) = \left(-\frac{\partial^2 l}{\partial \eta_i \partial \eta_j} \right) - \begin{bmatrix} \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \alpha^2} & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \alpha \partial \lambda_1} & \dots & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \alpha \partial \lambda_m} \\ \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_1 \partial \alpha} & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_1^2} & \dots & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_1 \partial \lambda_m} \\ \dots & \dots & \dots & \dots \\ \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_m \partial \alpha} & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_m \partial \lambda_1} & \dots & \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_m^2} \end{bmatrix},$$

where

$$\begin{aligned} \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \alpha^2} &= \frac{\sum_{i=1}^m k_i}{\alpha^2} + \sum_{i=1}^m \lambda_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha \ln^2 x_{ij} + (n_i - k_i)U_i^\alpha \ln^2 U_i \right), \\ \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \alpha \partial \lambda_i} &= \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_i \partial \alpha} = \sum_{j=1}^{k_i} x_{ij}^\alpha \ln x_{ij} + (n_i - k_i)U_i^\alpha \ln U_i, \\ \frac{\partial^2 \ell(\eta, \mathbf{X})}{\partial \lambda_i^2} &= \frac{k_i}{\lambda_i^2}. \end{aligned}$$

The asymptotic distribution of $\hat{\eta} = (\alpha, \lambda_1, \dots, \lambda_m)^T$ is given by $\hat{\eta} - \eta \sim N_{(m+1)}(0, F^{-1})$. Therefore, $100(1 - \delta)\%$ ($0 < \delta < 1$) ACI of $\hat{\eta}_i$ is given by

$$\hat{\eta}_i \pm z_{\frac{\delta}{2}} \sqrt{v_{ii}}$$

where v_{ii} is the $(i, i)^{th}$ element of the matrix F^{-1} and $z_{\frac{\delta}{2}}$ is 100th percentile of standard normal distribution $N(0, 1)$. However, the ACIs cannot always guarantee that the lower bound of the interval is positive. Therefore, we can propose the delta method introduced by Greene (2012). Due to the asymptotic normality property of MLEs, $\hat{\eta}_i \sim N(\eta, v_{ii})$. Using a logarithmic transformation and delta method on η_i , the asymptotic normal distribution of $\ln(\eta_i)$ is obtained as $\ln(\hat{\eta})_i \sim N(\ln(\eta), v_{ii}(\hat{\eta})^{-2})$. Thus, the $100(1 - \delta)\%$ ACI for $\ln(\eta_i)$ is obtained as

$$\log \hat{\eta} \pm Z_{\frac{\delta}{2}} \sqrt{\frac{v_{ii}}{\hat{\eta}^2}}$$

and the revised $100(1 - \delta)\%$ ACI of $\hat{\eta}_i$ by delta method can be given as

$$\left(\hat{\eta} \exp \left\{ -Z_{\frac{\delta}{2}} \sqrt{\frac{v_{ii}}{\hat{\eta}^2}} \right\}, \hat{\eta} \exp \left\{ Z_{\frac{\delta}{2}} \sqrt{\frac{v_{ii}}{\hat{\eta}^2}} \right\} \right)$$

3. Bayesian Inference

In this section, we provide the Bayesian estimation for the unknown parameters $\eta = (\alpha, \lambda_1, \dots, \lambda_m)^T$. Due to its positive aspect in small sample cases, the Bayesian method is a robust alternative to the MLEs for estimation. A crucial step in the Bayesian inference process involves establishing prior distributions for the unknown parameters. Arnold and Press (1983) suggests that no specific prior can be claimed to be superior. Instead, they propose focusing on a flexible family of prior distributions and selecting the one that best suits the problem at hand. In our research, all parameters are confined within the range $[0, \infty)$. In many studies, the prior distributions for unknown parameters are typically assumed to be independent. However, it is worth exploring the possibility of interrelated priors for model parameters, especially considering the constraints of the experiments and the mutual influence observed in practical applications. Therefore, we have assumed a dependent prior for η using the concept of Peña and Gupta (1990). In this purpose, we take $\lambda_{1m} = \lambda_1 + \lambda_2 + \dots + \lambda_m$ and it is assumed that λ_{1m} has a Gamma prior $GA(a_0, b_0)$ with the pdf

$$\pi_0(\lambda_{1m} | a_0, b_0) = \frac{b_0^{a_0}}{\Gamma(a_0)} \lambda_{1m}^{a_0-1} e^{-b_0 \lambda_{1m}}, \quad (10)$$

where $a_0 > 0, b_0 > 0, \lambda_{1m} > 0$. For given λ_{1m} , $\left(\frac{\lambda_1}{\lambda_{1m}}, \frac{\lambda_2}{\lambda_{1m}}, \dots, \frac{\lambda_m}{\lambda_{1m}} | \lambda_{1m}, c_1, c_2, \dots, c_m \right)$ follows a Dirichlet prior, say $\pi_1(\cdot | c_1, c_2, \dots, c_m)$ and that is given by

$$\pi_1 \left(\frac{\lambda_1}{\lambda_{1m}}, \frac{\lambda_2}{\lambda_{1m}}, \dots, \frac{\lambda_m}{\lambda_{1m}} | \lambda_{1m}, c_1, c_2, \dots, c_m \right) = \frac{\Gamma(c_{1m})}{\Gamma(c_1) \Gamma(c_2) \dots \Gamma(c_m)} \prod_{i=1}^m \left(\frac{\lambda_i}{\lambda_{1m}} \right)^{c_i}, \quad (11)$$

where $\lambda_i > 0$ for $i = 1, 2, \dots, m$ and $c_i > 0$ are the parameters of the Dirichlet distribution (known as the gamma-Dirichlet distribution, which is an extension to higher dimension of a gamma-beta distribution) with $c_{1m} = c_1 + c_2 + \dots + c_m$. Then, the joint prior of $(\lambda_1, \lambda_2, \dots, \lambda_m)$, namely $\pi(\lambda_1, \lambda_2, \dots, \lambda_m | a_0, b_0, c_1, c_2, \dots, c_m)$, can

be expressed as

$$\pi_1(\lambda_1, \lambda_2, \dots, \lambda_m) = \frac{\Gamma(c_{1m})}{\Gamma(a_0)} \prod_{i=1}^m \frac{b_0^{a_0}}{\Gamma(c_i)} \lambda_i^{c_i-1} e^{-b_0 \lambda_i} (\lambda_{1m} b_0)^{a_0-c_{1m}}, \quad (12)$$

In general, this is a dependent prior but if $a_0 = c_{1m}$, then λ_i 's are independent gamma priors with parameters b_0 and c_i ($i = 1, 2, \dots, m$). On the other hand, the gamma priors are most considerable of their flexible structure and conjugate type. Therefore, we propose a gamma prior for α with density functions given as in the following

$$\alpha \sim \mathbf{GA}(a_1, b_1), \quad \pi(\alpha) \propto \alpha^{a_1-1} \exp\{-\alpha b_1\},$$

where a_1 and b_1 , are the hyper-parameters and they are assumed as non-negative and known. Suppose there is no any prior information about the parameters. In that case, that is in a non-informative case, quite small non-negative values such as 0.0001 can be used according to the suggestion of Congdon (2001). These values of the hyper-parameters are almost like Jeffrey's priors, but they are proper, inversely. Different selections for the hyper-parameters can be evaluated in the more informative cases. Among them, considering means of the samples can be used. Since a_0/b_0 and λ_i/λ_{1m} values provide the means of the gamma and Dirichlet distributions respectively, these hyper-parameters values can be determined as providing them equal to the MLE value. Given $(\lambda_1, \lambda_2, \dots, \lambda_m)$ and α directly calculate the joint posterior distribution $\pi(\lambda_1, \lambda_2, \dots, \lambda_m, \alpha | data)$ as

$$\begin{aligned} \pi_\eta(\alpha, \boldsymbol{\lambda}, \mathbf{X}) &\propto \left(\prod_{i=1}^m L(\alpha, \lambda_i, X_i) \right) \pi_1(\lambda_0, \lambda_1, \dots, \lambda_m) \pi(\alpha) \\ &\propto \alpha^{a_1 + \sum_{i=1}^m k_i - 1} \left(\prod_{i=1}^m \lambda_i^{k_i + c_i - 1} \right) \exp \left\{ - \sum_{i=1}^m \lambda_i \Delta_i(\alpha) - b_0 \lambda_{1m} \right\} \lambda_{1m}^{a_0 - c_{1m}} \\ &\times \exp \left\{ -\alpha \left(b_1 - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) \right) \right\}, \end{aligned} \quad (13)$$

where $\Delta_i(\alpha) = \left(\sum_{j=1}^{k_i} x_{ij}^\alpha + (n_i - k_i) U_i^\alpha \right)$. Under the squared error loss function (SELF), the Bayes estimate of η , denoted by $\hat{\eta}^B$ can be obtained as the mean of the posterior function as given in the following

$$E(\eta | \mathbf{X}) = \int_0^\infty \int_0^\infty \dots \int_0^\infty \eta \pi_\eta(\alpha, \boldsymbol{\lambda}, \mathbf{X}) d\alpha d\lambda_1 d\lambda_2 \dots d\lambda_m. \quad (14)$$

Here, an expression for the posterior means of η is not seen to be possible to obtain, explicitly. For this reason, the posterior conditional distributions of parameters are calculated using the Markov chain Monte Carlo method (MCMC) with Gibbs sampling

as expressed in the following

$$\begin{aligned}\pi(\lambda_i|\alpha, \lambda_j, \mathbf{X}) &\propto \lambda_i^{k_i+c_i-1} \bar{\lambda}_{1m}^{a_0-c_{1m}} \exp \left\{ -\lambda_i \left(\Delta_i(\alpha) + b_0 \right) \right\} \\ \pi(\alpha|\lambda, \mathbf{X}) &\propto \alpha^{a_1+\sum_{i=1}^m k_i-1} \exp \left\{ -\sum_{i=1}^m \lambda_i \Delta_i(\alpha) \right\} \exp \left\{ -\alpha \left(b_1 - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}) \right) \right\},\end{aligned}\tag{15}$$

where $\bar{\lambda}_{1m} = \lambda_i + \lambda_j$, $i, j = 1, 2, \dots, m$ and $j \neq i$.

It is evident that there is no analytical way to reduce random samples of λ_i and $\pi(\alpha|\lambda_i, \mathbf{X})$ to well-known distributions. As a result, λ_i and α samples cannot be obtained directly using conventional techniques. Additionally, it is seen that the conditional posterior density plots resemble the Gaussian distribution (see Fig. 2). In this scenario, Tierney (1994) suggests using the Metropolis-Hasting (M-H) sampling in the Gibbs algorithm using a normal proposal distribution. By creating samples using the Gibbs sampling with the M-H algorithm under the standard proposal, we perform the MCMC approach. The algorithm that is employed is this one.

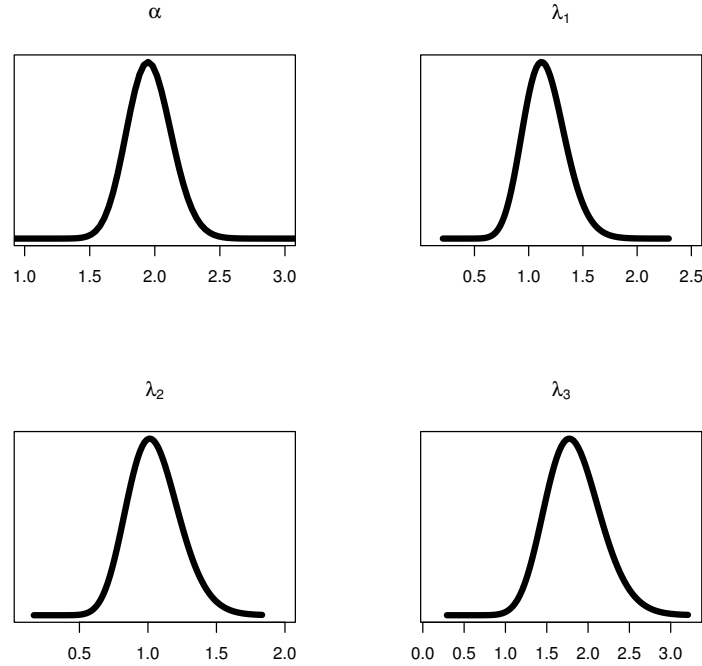


Figure 2. Posterior probability density function plots of α and λ_i values for $m = 3$.

Step 1: Determine the initial values of $\alpha^{(0)}$ and $\lambda_i^{(0)}$, for $i = 1, 2, \dots, m$.

Step 2: Set $t = 1$.

Step 3: Using the M-H algorithm, generate $\alpha^{(t)}$ from $\pi(\alpha|\lambda, \mathbf{X})$ with normal proposal $N(\alpha^{(t-1)}, \sigma_{\hat{\alpha}})$.

Step 3.1: Let $v = \alpha^{(t-1)}$ and generate w from the proposal as $w =$

$$N(\alpha^{(t-1)}, \sigma_{\hat{\alpha}}).$$

Step 3.2: Let $p(v, w) = \min \left\{ 1, \frac{\pi(w|\lambda^{(t)})N(v, \sigma_{\hat{\alpha}})}{\pi(v|\lambda^{(t)})N(w, \sigma_{\hat{\alpha}})} \right\}$.

Step 3.3: Generate U from $U(0, 1)$, then accept proposal if $U \leq p(v, w)$ and set $\alpha^{(t)} = w$ or otherwise $\alpha^{(t)} = v$.

Step 4: Generate $\lambda_i^{(t)}$ using the M-H algorithm with a normal proposal using the similar procedure given in Step 3.

Step 5: Set $t = t + 1$.

Step 6: Repeat 2-6 with M iteration.

Then, under the squared error loss function, the approximate posterior means of the parameters, denoted by $\hat{\alpha}^B$ and $\hat{\lambda}_i^B$, can be obtained as

$$\hat{\alpha}^B = \frac{1}{M-S} \sum_{t=S+1}^M \alpha^{(t)} \text{ and } \hat{\lambda}_i^B = \frac{1}{M-S} \sum_{t=S+1}^M \lambda_i^{(t)},$$

where S indicates the burn-in process. Then, following the method introduced by Chen and Shao (1999), the $100(1 - \gamma)\%$ HPD credible intervals of the Bayesian estimations for the parameters can be obtained as

$$\left(\hat{\alpha}_{[\frac{\gamma}{2}(M-S)]}^B, \hat{\alpha}_{[(1-\frac{\gamma}{2})(M-S)]}^B \right) \quad \text{and} \quad \left(\hat{\lambda}_{i[\frac{\gamma}{2}(M-S)]}^B, \hat{\lambda}_{i[(1-\frac{\gamma}{2})(M-S)]}^B \right),$$

where $[\frac{\gamma}{2}(M-S)]$ and $[(1-\frac{\gamma}{2})(M-S)]$ are the smallest integers that are smaller than or equal to $\frac{\gamma}{2}(M-S)$ and $(1-\frac{\gamma}{2})(M-S)$, respectively.

3.1. Order-Restricted Bayesian Inference

Let's keep in mind the order restriction on the scale parameters as $\lambda_1 < \lambda_2 < \dots < \lambda_m$. Then, we obtain the joint prior of $(\lambda_1, \lambda_2, \dots, \lambda_m)$ as

$$\pi_{or}(\lambda_1, \lambda_2, \dots, \lambda_m) = \frac{\Gamma(c_{1m})}{\Gamma(a_0)\Gamma(c_1)\dots\Gamma(c_m)} b_0^{a_0} e^{-b_0 \lambda_{1m}} \lambda_{1m}^{a_0-c_{1m}} \sum_{\mathcal{P}} (\lambda_{i_1}^{c_1-1} \lambda_{i_2}^{c_2-1} \dots \lambda_{i_m}^{c_m-1}), \quad (16)$$

where $0 < \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m < \infty$. Here, π_{or} is an order-restricted gamma-Dirichlet (OGD) prior and $\mathcal{P}$ denotes the set of all the $m!$ permutations of I . Similarly to the unrestricted case, we propose gamma prior for α as $\mathbf{GA}(a_1, b_1)$. Hence, the joint posterior density function of $\alpha > 0$ and $0 < \lambda_1 < \lambda_2 < \dots < \lambda_m < \infty$ can be written as

$$\begin{aligned} L_{or}(\alpha, \lambda, \mathbf{X}) &\propto \lambda_{1m}^{a_0-c_{1m}} \prod_{i=1}^m \lambda_i^{k_i} e^{-\lambda_i(b_0+\Delta_i(\alpha))} \sum_{(i_1, i_2, \dots, i_m) \in \mathcal{P}} (\lambda_{i_1}^{c_1-1} \lambda_{i_2}^{c_2-1} \dots \lambda_{i_m}^{c_m-1}) \\ &\times \alpha^{a_1+\sum_{i=1}^m k_i-1} e^{-b_1 \alpha} \left(\prod_{i=1}^m \prod_{j=1}^{k_i} x_{ij}^\alpha \right). \end{aligned} \quad (17)$$

Similar to the results presented in Pal, Mitra, and Kundu (2021), we introduce an importance sampling algorithm to derive Bayesian estimates and their corresponding approximate confidence intervals. This is achieved by modifying the joint posterior density in (17) as follows:

$$L_{or}(\alpha, \boldsymbol{\lambda}, X) \propto \lambda_{1m}^{a_0 - c_{1m}} \prod_{i=1}^m \left(\lambda_i^{k_i - J} e^{-\lambda_i(\Delta_i(\alpha) - S(\alpha))} \right) e^{-\lambda_{1m}(b_0 + S(\alpha))} \alpha^{a_1 + \sum_{i=1}^m k_i - 1} \\ \times e^{-\alpha(b_1 - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln x_{ij})} \sum_{(i_1, i_2, \dots, i_m) \in \mathcal{P}} \left(\lambda_{i_1}^{c_1 + J - 1} \lambda_{i_2}^{c_2 + J - 1} \dots \lambda_{i_m}^{c_m + J - 1} \right), \quad (18)$$

where $J = \min(k_i)$ and $S(\alpha) = \min(\Delta_i(\alpha))$ for $i = 1, 2, \dots, m$. Then, we can obtain the posterior densities of the parameters as

$$\pi(\lambda, \alpha | \mathbf{data}) \propto \pi_1(\lambda | \alpha) \pi_2(\alpha | \lambda) \mathcal{H}(\lambda, \alpha),$$

where

$$\pi_1(\lambda | \alpha, \mathbf{data}) \propto \mathbf{OGD}(\mathbf{a}_0 + \mathbf{mJ}, \mathbf{b}_0 + \mathbf{S}(\alpha), \mathbf{c}_1 + \mathbf{J}, \mathbf{c}_2 + \mathbf{J}, \dots, \mathbf{c}_m + \mathbf{J}), \quad (19)$$

$$\pi_2(\alpha | \lambda) \propto \alpha^{a_1 + \sum_{i=1}^m k_i - 1} e^{-\alpha(b_1 - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln x_{ij})}, \quad (20)$$

$$\mathcal{H}(\lambda, \alpha) = \frac{\prod_{i=1}^m \lambda_i^{k_i - J} e^{-\lambda_i(\Delta_i(\alpha) - S(\alpha))}}{[b_0 + S(\alpha)]^{a_0 + mJ}}. \quad (21)$$

The condition of $b_1 > \sum_{i=1}^m \sum_{j=1}^{k_i} \ln x_{ij}$ makes $\pi_2(\alpha | \lambda)$ a proper and well-known log-concave density and generation from $\pi_1(\lambda | \alpha, \mathbf{data})$ is quite convenient to perform. Then, we can generate λ_i 's using the algorithm given by Pal et al. (2021) (page 8). Thus, we can perform the following importance sampling algorithm to obtain all unknown parameters

Step 1: Generate λ_i from $\pi_1(\lambda | \alpha, \mathbf{data})$.

Step 2: Generate α from $\pi_2(\alpha | \lambda, \mathbf{data})$.

Step 3: Compute $\mathcal{H} = \mathcal{H}(\lambda, \alpha)$.

Step 4: Repeat steps 1-3 M times, and generate M estimates of α_j , λ_{ij} , and $\mathcal{H}_j$, for $j = 1, 2, \dots, m$.

Thus, an approximate Bayes estimate of λ_i , $i = 1, 2, \dots, m$, under a squared error loss function can be obtained as

$$\hat{\lambda}_{i_B} = \frac{\frac{1}{M} \sum_{j=1}^M \lambda_{ij} \mathcal{H}_j}{\frac{1}{M} \sum_{j=1}^M \mathcal{H}_j}. \quad (22)$$

Bayesian estimate of α can be obtained similarly. Then, the credible interval of α and λ_i 's can be obtained by using the idea of Chen and Shao (1999). Let assume that $\nabla_j = \mathcal{H}(\alpha_j, \lambda_{ij})$ where α_j and λ_{ij} for $j = 1, 2, \dots, M$ is posterior samples generated respectively from (19) and (20) for α and λ_i . Let U_j be the ordered values of U_j and

define

$$\omega_j = \frac{\mathcal{H}(\alpha_j, \lambda_{ij})}{\sum_{j=1}^M \mathcal{H}(\alpha_j, \lambda_{ij})}.$$

To compute a $100(1 - \gamma)\%$ CRI of ∇_j , arrange ∇_{ji} in ascending order to obtain $(\nabla_{j1}, \nabla_{j2}, \dots, \nabla_{jM})$ and record the corresponding ω_j as $(\omega_{j1}, \omega_{j2}, \dots, \omega_{jM})$. It should be noted that ∇_j must be ordered but ω_j does not need to be ordered. A $100(1 - \gamma)\%$ CRI can be obtained as $(\nabla_{j_{L_0}}, \nabla_{j_{U_0}})$ where L_0 and U_0 such that,

$$L_0 < U_0, \quad L_0, U_0 \in \{1, \dots, M\}, \quad \sum_{i=L_0}^{U_0} \omega_i \leq 1 - \gamma < \sum_{i=L_0}^{U_0+1} \omega_i. \quad (23)$$

The $100(1 - \gamma)\%$ highest posterior density (HPD) CRI can be obtained as $(\nabla_{j_{L_0}}^*, \nabla_{j_{U_0}}^*)$, such that $\nabla_{j_{U_0}}^* - \nabla_{j_{L_0}}^* \leq \nabla_{j_{U_0}} - \nabla_{j_{L_0}}$ and L_0^*, U_0^* satisfying (23) for all L_0, U_0 satisfying (23).

4. Optimal Censoring Plan

This section outlines optimality criteria for selecting an optimal censoring plan. These criteria, outlined in Table 1, are commonly used by various authors to determine optimal censoring schemes (see Dey, Elshahhat, and Nassar (2023)). The A-optimality and D-optimality criteria aim to minimize the trace and determinant of approximate variance-covariance matrices, as detailed in Section 2.2. Similarly, the F-optimality criterion seeks to maximize the observed values of the Fisher information matrix concerning the MLEs. By minimizing variance for any reliability problem, these methods can be used to compare different censoring plans

Table 1. Some criteria for optimal censoring plan

Criterion	Aim
D-optimality	Minimize $\det(I^{-1}(\hat{\theta}))$
A-optimality	Minimize $tr(I^{-1}(\hat{\theta}))$
F-optimality	Maximize $tr(I(\hat{\theta}))$

Listing all possible combinations of censoring schemes is impractical due to their huge number. Hence, we provide a few example schemes to illustrate how varying failure numbers and pre-fixed time points impact block censoring plans. Setting $\alpha = 2$ and $\lambda_i = 1.25$, with equal sample sizes $n_i = 50$ for $m = 3$, we consider two different approaches for identifying optimal plans. In the first approach, we experiment with various combinations of failure numbers (s_i), computing the corresponding optimality values and optimal time points (T_i), presented in Table 2. In the second approach, we explore different selections of pre-fixed time points (T_i) to determine optimal failure numbers (s_i), as shown in Table 3.

The results indicate that higher time points yield optimal outcomes when dealing with a substantial number of observations. Conversely, smaller optimal values for s_i are evident with early pre-fixed time points. As T_i values increase, corresponding optimal s_i values also grow. These findings are consistent with expectations, suggesting that a higher number of observed items leads to improved results.

Table 2. Optimal values based on different optimality criteria and optimal T_i values under different schemes of the failure numbers for $m = 3$.

s_i	D-opt	T_i	A-opt	T_i	F-opt	T_i
70% $\times n_i$	0.0000*	1.20	0.1226	1.20	134.1904	1.20
80% $\times n_i$	0.0000*	1.19	0.1140	1.19	147.1584	1.19
90% $\times n_i$	0.0000*	1.40	0.0966	1.40	173.3461	1.40
(70% $\times n_1$, 90% $\times n_2$, 80% $\times n_3$)	0.0000*	1.43	0.1046	1.43	161.4981	1.43
(90% $\times n_1$, 90% $\times n_2$, 70% $\times n_3$)	0.0000*	1.53	0.0959	1.53	172.1172	1.53
(90% $\times n_1$, 80% $\times n_2$, 70% $\times n_3$)	0.0000*	1.41	0.0976	1.41	177.0472	1.41
(70% $\times n_1$, 80% $\times n_2$, 90% $\times n_3$)	0.0000*	1.56	0.0974	1.56	167.8341	1.56

* Smaller than 10^{-6} .

Table 3. Optimal values based on different optimality criteria and optimal s_i values under different schemes of the fixed time point, for $m = 3$.

T_i	D-opt	s_i	A-opt	s_i	F-opt	s_i
$T_1 = 2.00, T_2 = 2.25, T_3 = 1.75$	0.0000*	49	0.1039	49	163.8839	49
$T_1 = 1.75, T_2 = 1.80, T_3 = 1.25$	0.0000*	47	0.0825	47	202.1385	47
$T_1 = 1.35, T_2 = 1.50, T_3 = 1.75$	0.0000*	46	0.1056	46	162.7841	46
$T_1 = 1.50, T_2 = 1.40, T_3 = 1.30$	0.0000*	48	0.1168	48	153.4679	48
$T_1 = 1.30, T_2 = 1.40, T_3 = 1.50$	0.0000*	48	0.0973	48	170.9275	48
$T_1 = 1.10, T_2 = 1.20, T_3 = 1.15$	0.0000*	41	0.1254	41	137.4952	41
$T_1 = 0.75, T_2 = 0.75, T_3 = 1.00$	0.0000*	40	0.1366	40	135.0147	40
$T_1 = 0.50, T_2 = 0.75, T_3 = 0.50$	0.0000*	37	0.2788	37	89.3541	40

* Smaller than 10^{-6} .

5. Simulation Studies

In this section, the theoretical findings of this study are illustrated by considering different simulation schemes. For this purpose, we consider four different censoring scheme and names them as $CS - I$, $CS - II$, $CS - III$ and $CS - IV$. In these plans, we consider specific values for $m = 3, 5, 7$, and 10, respectively. The total sample size remains constant at $N = 100$ over all plans. For the first two plans, we assign varying sets of λ_i values. On the other hand, we take λ_i values as the same. These selections aim to evaluate the methods performance across all scenarios. For $CS - II$ and $CS - IV$, we maintain equal sample sizes in each set, while in $CS - I$ and $CS - III$, we use varying sample sizes. Refer to Table 4 for the details of all censoring schemes used in this simulation study. We consider three different failure rate values (s), namely 70%, 80%, and 90%, in all schemes to investigate the performance of the proposed method under varying levels of data availability. These values correspond to moderate-to-high censoring scenarios commonly encountered in reliability and survival analysis. By gradually reducing the censoring level from 30% to 10%, we aim to evaluate how the method performs across a range of practical situations, from relatively limited to more informative data. For convenience, we denote these cases as subsets a , b , and c of each censoring scheme. For example, we take s_i values as 7, 8 and 9 in $CS - IV$ and we denote these schemes as $CS - IV_a$, IV_b and IV_c , respectively. We determined prefixed time points T_i , $i = 1, 2, \dots, m$ by considering the actual parameter values of samples. For this purpose, we generate 10,000 samples with these actual parameters and observe the ranges of the samples. Then, we determine T_i values as being between these ranges. We use the same censoring rates and denotations in other schemes. On the other hand, we illustrate the case of order restriction. For this purpose, we take $\alpha = 0.50$, $\lambda_1 = 0.75$, $\lambda_2 = 1.25$, $\lambda_3 = 1.50$ with $n_1 = 20$, $n_2 = 20$, $n_3 = 20$. In this

scheme, we use the same censoring rates and denotations. The simulation steps can be described as follows.

- Step 1:* Determine the initial parameter values (α, λ_i) and the experimental setup values (T_i, s_i) for $i = 1, 2, \dots, m$.
- Step 2:* Generate m different set with sizes n_i from $WE(\alpha, \lambda_i)$.
- Step 3:* Determine U_i values for each set where $U_i = \min(X_{is_i}, T_i)$.
- Step 4:* Calculate parameter estimations and their approximate confidence intervals of the parameters.

Table 4. The predetermined censoring schemes for the simulation studies

CS	θ	n_i	s_i	T_i
I_a	(2.00, 1.25, 1.00, 1.75)	(35, 30, 35)	(25, 21, 25)	(1.75, 1.80, 1.25)
I_b	(2.00, 1.25, 1.00, 1.75)	(35, 30, 35)	(28, 24, 28)	(1.75, 1.80, 1.25)
I_c	(2.00, 1.25, 1.00, 1.75)	(35, 30, 35)	(32, 27, 32)	(1.75, 1.80, 1.25)
II_a	(1.50, 0.75, 1.25, 1.50, 1.00, 0.80)	20 ₍₅₎	14 ₍₅₎	(2.50, 1.75, 1.50, 2.00, 2.50)
II_b	(1.50, 0.75, 1.25, 1.50, 1.00, 0.80)	20 ₍₅₎	16 ₍₅₎	(2.50, 1.75, 1.50, 2.00, 2.50)
II_c	(1.50, 0.75, 1.25, 1.50, 1.00, 0.80)	20 ₍₅₎	18 ₍₅₎	(2.50, 1.75, 1.50, 2.00, 2.50)
III_a	(1.50, 1.20 ₍₇₎)	(14, 15, 14 ₍₃₎ , 15, 14)	(10, 11, 10 ₍₃₎ , 11, 10)	1.75 ₍₇₎
III_b	(1.50, 1.20 ₍₇₎)	(14, 15, 14 ₍₃₎ , 15, 14)	(10, 11, 10 ₍₃₎ , 11, 10)	1.75 ₍₇₎
III_c	(1.50, 1.20 ₍₇₎)	(14, 15, 14 ₍₃₎ , 15, 14)	(10, 11, 10 ₍₃₎ , 11, 10)	1.75 ₍₇₎
IV_a	(2.00, 1.50 ₍₁₀₎)	10 ₍₁₀₎	7 ₍₁₀₎	1.25 ₍₁₀₎
IV_b	(2.00, 1.50 ₍₁₀₎)	10 ₍₁₀₎	8 ₍₁₀₎	1.25 ₍₁₀₎
IV_c	(2.00, 1.50 ₍₁₀₎)	10 ₍₁₀₎	9 ₍₁₀₎	1.25 ₍₁₀₎

For computations, we use the R software R Core Team (2023). The simulations are repeated 2000 times, employing Markov chains with a sample size of 3500 in each iteration. A burn-in number of $K = 500$ is used in the Markov chains, with a thinning procedure that involves selecting every third variate as uncorrelated samples. This process enables the derivation of absolute biases (AB) and mean squared errors (MSE) for the point estimations. Additionally, the average lengths (AL) of the approximate confidence intervals are computed, along with their corresponding coverage probabilities (CP). The performance of unrestricted parameter cases are presented in Tables 5 to 8. The results for the restricted case are outlined in Table 9.

Simulation results in Tables 5-8 show that quite small MSE, AB, and AL values are observed. The Bayesian method provides better results, especially in small sample sizes. Differences between the results of the two inference methods decrease with increasing sample size. The performance assessment between the censoring scheme is not possible since the observed sample size differs according to the failure number and fixed time point. However, it is clearly seen that the estimation procedures perform well in all cases. In addition to the good performances of the point estimations, approximate confidence intervals show adequate performance. These intervals have quite close short average lengths and also they provide quite close coverage probabilities to their actual value of 0.95. Consequently, these simulation results show us theoretical findings of this study provide the theoretical expectations and perform well. The results in Tables 9 also show good performances similar to the unrestricted case. Since the sample sizes in this case are quite small, the Bayesian method shows better performance in most cases.

Table 5. Inferences of the parameters under $CS - I$ scheme.

CS	η	AB		MSE		AL		CP	
		MLE	$Bayes$	MLE	$Bayes$	ACI	HPD	ACI	HPD
I_a	α	0.0850	0.0167	0.0577	0.0419	0.8653	0.8246	0.9295	0.9575
	λ_1	0.0734	0.0093	0.0856	0.0481	1.0963	0.9673	0.9530	0.9680
	λ_2	0.0581	0.0121	0.0685	0.0380	0.9447	0.8362	0.9490	0.9670
	λ_3	0.1588	0.0091	0.2482	0.1113	1.6764	1.4105	0.9370	0.9655
I_b	α	0.0729	0.0268	0.0476	0.0376	0.7893	0.7634	0.9365	0.9525
	λ_1	0.0575	0.0067	0.0699	0.0436	1.0039	0.9035	0.9495	0.9725
	λ_2	0.0414	0.0119	0.0534	0.0339	0.8606	0.7767	0.9475	0.9690
	λ_3	0.1292	0.0146	0.1980	0.1029	1.5021	1.3031	0.9355	0.9595
I_c	α	0.0551	0.0285	0.0362	0.0313	0.7057	0.6915	0.9410	0.9555
	λ_1	0.0423	0.0073	0.0569	0.0392	0.9179	0.8383	0.9540	0.9680
	λ_2	0.0305	0.0131	0.0459	0.0315	0.8027	0.7331	0.9485	0.9650
	λ_3	0.0889	0.0056	0.1450	0.0887	1.3342	1.1897	0.9395	0.9580

Table 6. Inferences of the parameters under $CS - II$ scheme.

CS	η	AB		MSE		AL		CP	
		MLE	$Bayes$	MLE	$Bayes$	ACI	HPD	ACI	HPD
II_a	α	0.1105	0.0454	0.0443	0.0278	0.6747	0.6396	0.8935	0.9500
	λ_1	0.0583	0.0183	0.0701	0.0291	0.8937	0.7407	0.9285	0.9535
	λ_2	0.1431	0.0250	0.2289	0.0682	1.5743	1.1971	0.9200	0.9655
	λ_3	0.1796	0.0296	0.3260	0.0824	1.9314	1.4228	0.9385	0.9720
	λ_4	0.1128	0.0124	0.1375	0.0463	1.2361	0.9789	0.9255	0.9625
	λ_5	0.0787	0.0108	0.0825	0.0324	0.9715	0.7945	0.9250	0.9630
II_b	α	0.0899	0.0534	0.0337	0.0252	0.6033	0.5869	0.8975	0.9415
	λ_1	0.0397	0.0209	0.0539	0.0271	0.8147	0.6947	0.9290	0.9545
	λ_2	0.1064	0.0227	0.1718	0.0640	1.4023	1.1211	0.9180	0.9605
	λ_3	0.1297	0.0275	0.2233	0.0756	1.7018	1.3299	0.9435	0.9715
	λ_4	0.0809	0.0149	0.1028	0.0436	1.1053	0.9155	0.9270	0.9640
	λ_5	0.0605	0.0110	0.0643	0.0303	0.8854	0.7482	0.9290	0.9595
II_c	α	0.0699	0.0537	0.0262	0.0224	0.5425	0.5376	0.9125	0.9410
	λ_1	0.0300	0.0225	0.0440	0.0254	0.7626	0.6630	0.9305	0.9475
	λ_2	0.0743	0.0286	0.1285	0.0587	1.2783	1.0573	0.9245	0.9590
	λ_3	0.0993	0.0286	0.1803	0.0733	1.5523	1.2572	0.9410	0.9660
	λ_4	0.0580	0.0207	0.0835	0.0414	1.0154	0.8635	0.9285	0.9605
	λ_5	0.0456	0.0149	0.0518	0.0282	0.8226	0.7092	0.9330	0.9520

6. Numerical Example

Applications of the BSC on the real-data examples can be performed using the following steps.

Step 1: Gather m datasets from the same experimental environment.

Step 2: Assign s_i and T_i values according to the experimental conditions. ($i = 1, 2, \dots, m$).

Step 3: Determine U_i values for each of m -set.

Step 4: Calculate parameter estimations and their approximate confidence intervals of the parameters.

In this section, we use an illustrative data set to illustrate the theoretical outcomes of the study. The data set involves jute fiber measurements from Xia, Yu, Cheng, Liu, and Wang (2009). The experiment measured the diameters of jute fibers using an XSP-8CA digital biological microscope from the Shanghai Optical Instrument Factory,

Table 7. Inferences of the parameters under $CS - III$ scheme.

CS	η	AB		MSE		AL		CP	
		MLE	$Bayes$	MLE	$Bayes$	ACI	HPD	ACI	HPD
III_a	α	0.1417	0.0285	0.0517	0.0239	0.6730	0.6217	0.8575	0.9595
	λ_1	0.1624	0.0604	0.2934	0.0461	1.8390	1.1474	0.9335	0.9740
	λ_2	0.1582	0.0561	0.2718	0.0462	1.7357	1.1191	0.9195	0.9740
	λ_3	0.1912	0.0564	0.3832	0.0486	1.8850	1.1539	0.9195	0.9720
	λ_4	0.1676	0.0608	0.3005	0.0458	1.8283	1.1482	0.9305	0.9815
	λ_5	0.1813	0.0580	0.3495	0.0505	1.8690	1.1521	0.9210	0.9785
	λ_6	0.1586	0.0558	0.2771	0.0466	1.7365	1.1200	0.9295	0.9775
	λ_7	0.1694	0.0619	0.3192	0.0490	1.8503	1.1455	0.9240	0.9725
III_b	α	0.1136	0.0462	0.0375	0.0224	0.5923	0.5675	0.8755	0.9520
	λ_1	0.1090	0.0546	0.1950	0.0434	1.5763	1.0787	0.9345	0.9735
	λ_2	0.1201	0.0534	0.2033	0.0441	1.5920	1.0802	0.9315	0.9680
	λ_3	0.1234	0.0525	0.2531	0.0477	1.5969	1.0826	0.9285	0.9725
	λ_4	0.1104	0.0551	0.1999	0.0450	1.5655	1.0801	0.9295	0.9765
	λ_5	0.1152	0.0534	0.2184	0.0482	1.5852	1.0808	0.9285	0.9730
	λ_6	0.1308	0.0495	0.2328	0.0456	1.6064	1.0848	0.9305	0.9785
	λ_7	0.1114	0.0555	0.2112	0.0462	1.5794	1.0784	0.9340	0.9715
III_c	α	0.0919	0.0496	0.0299	0.0213	0.5447	0.5316	0.8890	0.9360
	λ_1	0.0899	0.0548	0.1738	0.0440	1.4918	1.0515	0.9395	0.9710
	λ_2	0.0775	0.0547	0.1534	0.0441	1.4195	1.0241	0.9375	0.9695
	λ_3	0.0982	0.0545	0.2137	0.0480	1.5034	1.0521	0.9360	0.9715
	λ_4	0.0868	0.0566	0.1720	0.0456	1.4805	1.0506	0.9355	0.9730
	λ_5	0.0914	0.0548	0.1908	0.0488	1.4954	1.0518	0.9285	0.9660
	λ_6	0.0909	0.0488	0.1667	0.0452	1.4338	1.0310	0.9340	0.9755
	λ_7	0.0911	0.0554	0.1762	0.0456	1.4925	1.0499	0.9420	0.9730

China. For simplicity, the fibers were thought to be completely spherical. Twenty fibre samples were tested at four gauge lengths (5 mm, 10 mm, 15 mm, and 20 mm). The coefficient of variations of fibre diameter is the average value of diameter variation coefficients of 20 fibres at each gauge length. We consider each case of the gauge lengths as a set in the block and investigate the modeling of this system under the Type-I censoring scheme. In this study, we use the values of breaking strengths at gauge lengths 10mm, 15mm, and 20mm by considering them three sets such that X_1 , X_2 , and X_3 with sample sizes $n_i = 30$ for $i = 1, 2, 3$, respectively. That is, we take $m = 3$ and $N = 90$ in this example. We divide all samples to 1000 for better fitting to the reparameterized Weibull distributions. Different authors have been recently used this data set in stress-strength reliability problems. For example Saraçoğlu, Kınacı, and Kundu (2012) fitted this data to the exponential distribution and studied it under the progressive type-II censoring scheme. Then, Krishna, Dube, and Garg (2019) handled this data for the inverse Weibull distribution under progressive first failure censoring. Recently, Çetinkaya (2021) considered it for generalized progressive hybrid censored Weibull populations. We present the scaled data in the following

We first fit these data samples to the Weibull distributions by using graphical tools and the Kolmogorov-Smirnov (K-S) goodness of fit test. For this purpose, we fit Weibull distributions to these data samples using the MLE procedure, and we obtain the corresponding K-S test statistics and associated p-values are obtained as 0.1333 and 0.9578 for X_1 , 0.1667 and 0.8080 for X_2 , and 0.2333 and 0.3929 for X_3 , respectively. Thus, the null hypothesis that these data sets come from the Weibull distributions cannot be rejected. The QQ plots and the empirical cdf plots also support this observation (see Fig. 3 and 4).

Table 8. Inferences of the parameters under $CS - IV$ scheme.

		AB		MSE		AL		CP	
		<i>MLE</i>	<i>Bayes</i>	<i>MLE</i>	<i>Bayes</i>	<i>ACI</i>	<i>HPD</i>	<i>ACI</i>	<i>HPD</i>
<i>CS</i>	η								
	α	0.2204	0.0320	0.1106	0.0010	0.9309	0.7966	0.8390	0.9545
	λ_1	0.3363	0.1268	0.9263	0.0161	3.1053	1.2357	0.9220	0.9855
	λ_2	0.3331	0.1267	0.9561	0.0160	3.1022	1.2354	0.9150	0.9810
	λ_3	0.3116	0.1295	0.8803	0.0168	3.0610	1.2312	0.9310	0.9885
	λ_4	0.3160	0.1295	0.8992	0.0168	3.0563	1.2321	0.9280	0.9805
<i>IV_a</i>	λ_5	0.3357	0.1279	0.9464	0.0164	3.1051	1.2352	0.9245	0.9825
	λ_6	0.3060	0.1321	0.8571	0.0175	3.0503	1.2292	0.9145	0.9755
	λ_7	0.3374	0.1247	0.8952	0.0155	3.1047	1.2381	0.9315	0.9820
	λ_8	0.3371	0.1299	1.0303	0.0169	3.1149	1.2326	0.9105	0.9805
	λ_9	0.2847	0.1338	0.8548	0.0179	3.0139	1.2272	0.9305	0.9795
	λ_{10}	0.3320	0.1280	0.9037	0.0164	3.0984	1.2333	0.9175	0.9840
	α	0.1903	0.0074	0.0860	0.0001	0.8417	0.7511	0.8535	0.9485
	λ_1	0.2574	0.1152	0.6390	0.0133	2.7032	1.1958	0.9320	0.9845
	λ_2	0.2582	0.1142	0.6155	0.0130	2.7050	1.1965	0.9210	0.9805
	λ_3	0.2456	0.1176	0.6361	0.0138	2.6872	1.1923	0.9350	0.9830
	λ_4	0.2349	0.1193	0.6333	0.0142	2.6490	1.1886	0.9280	0.9745
<i>IV_b</i>	λ_5	0.2554	0.1158	0.6784	0.0134	2.7018	1.1940	0.9225	0.9805
	λ_6	0.2293	0.1208	0.6097	0.0146	2.6611	1.1874	0.9200	0.9745
	λ_7	0.2669	0.1132	0.6596	0.0128	2.7198	1.1962	0.9280	0.9805
	λ_8	0.2506	0.1196	0.6955	0.0143	2.6980	1.1895	0.9220	0.9770
	λ_9	0.2216	0.1228	0.6271	0.0151	2.6480	1.1869	0.9335	0.9780
	λ_{10}	0.2507	0.1174	0.7238	0.0138	2.6995	1.1924	0.9240	0.9830
	α	0.1557	0.0055	0.0673	0.0000	0.7713	0.7134	0.8710	0.9535
	λ_1	0.2017	0.1122	0.5291	0.0126	2.4474	1.1599	0.9295	0.9840
	λ_2	0.2075	0.1101	0.5160	0.0121	2.4557	1.1634	0.9270	0.9835
	λ_3	0.1787	0.1154	0.4705	0.0133	2.4157	1.1592	0.9385	0.9835
	λ_4	0.1784	0.1163	0.4964	0.0135	2.4002	1.1567	0.9300	0.9790
<i>IV_c</i>	λ_5	0.1977	0.1133	0.5276	0.0128	2.4432	1.1618	0.9290	0.9840
	λ_6	0.1839	0.1165	0.5248	0.0136	2.4245	1.1564	0.9285	0.9750
	λ_7	0.1967	0.1100	0.4909	0.0121	2.4405	1.1633	0.9320	0.9825
	λ_8	0.1935	0.1155	0.5377	0.0133	2.4381	1.1587	0.9165	0.9795
	λ_9	0.1680	0.1188	0.4886	0.0141	2.4013	1.1571	0.9315	0.9785
	λ_{10}	0.1876	0.1147	0.5289	0.0132	2.4295	1.1566	0.9330	0.9820

Table 9. Inferences of the parameters under order restriction.

		AB		MSE		AL		CP	
		<i>MLE</i>	<i>Bayes</i>	<i>MLE</i>	<i>Bayes</i>	<i>ACI</i>	<i>HPD</i>	<i>ACI</i>	<i>HPD</i>
<i>CS</i>	η								
	α	0.0379	0.0017	0.0077	0.0050	0.2954	0.2234	0.9135	0.8545
	λ_1	0.0414	0.0909	0.0557	0.0564	0.8894	0.6629	0.9470	0.8378
	λ_2	0.0785	0.0907	0.1326	0.1035	1.5176	1.1721	0.9700	0.8910
	λ_3	0.2696	0.0763	0.3703	0.1233	2.1130	1.4005	0.9475	0.9210
	α	0.0254	0.0066	0.0059	0.0047	0.2692	0.2445	0.9175	0.8940
	λ_1	0.0111	0.0494	0.0442	0.0446	0.8255	0.6376	0.9480	0.8744
	λ_2	0.0389	0.0659	0.0972	0.0943	1.3520	1.1607	0.9750	0.9055
	λ_3	0.1972	0.0614	0.2785	0.1166	1.8278	1.4210	0.9420	0.9215
	α	0.0186	0.0093	0.0048	0.0043	0.2564	0.2592	0.9365	0.9250
	λ_1	0.0057	0.0249	0.0391	0.0344	0.8125	0.6474	0.9610	0.9129
	λ_2	0.0239	0.0542	0.0821	0.0851	1.2802	1.1875	0.9735	0.9245
	λ_3	0.1311	0.0323	0.1942	0.1028	1.6611	1.4481	0.9500	0.9440

Since we assume common shape parameter α within the scope of this study, it should be checked these data samples have equal shape parameters. Therefore, we

Table 10. The breaking strengths (scaled by dividing to 1000) of jute fiber at different gauge lengths 10mm, 15mm, and 20mm given by Xia et al. (2009)

	0.69373	0.70466	0.32383	0.77817	0.12306	0.63766	0.38343	0.15148	0.10894	0.05016
X_1	0.67149	0.18316	0.25744	0.72723	0.29127	0.10115	0.37642	0.16340	0.14138	0.70074
	0.26290	0.35324	0.42211	0.04393	0.59048	0.21213	0.30390	0.50660	0.53055	0.17725
X_2	0.59440	0.20275	0.16837	0.57486	0.22565	0.07638	0.15667	0.12781	0.81387	0.56239
	0.46847	0.13509	0.07224	0.49794	0.35556	0.56907	0.64048	0.20076	0.55042	0.74875
	0.48966	0.67806	0.45771	0.10673	0.71630	0.04266	0.08040	0.33922	0.07009	0.19342
X_3	0.07146	0.41902	0.28464	0.58557	0.45660	0.11385	0.18785	0.68816	0.66266	0.04558
	0.57862	0.75670	0.59429	0.16649	0.09972	0.70736	0.76514	0.18713	0.14596	0.35070
	0.54744	0.11699	0.37581	0.58160	0.11986	0.04801	0.20016	0.03675	0.24453	0.08355

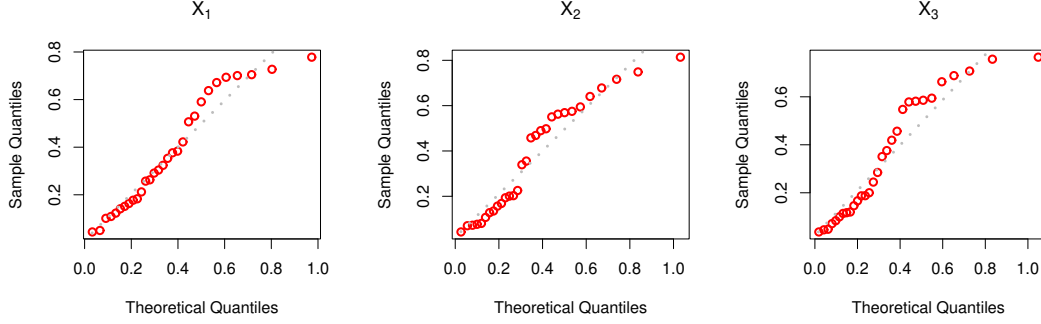


Figure 3. Q-Q plots of the fitted Weibull distributions for real data sets.

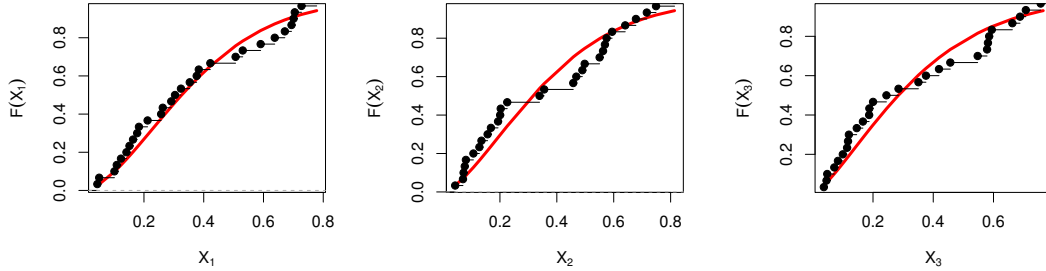


Figure 4. Theoretical and empirical CDF plots of the fitted Weibull distributions for real data sets.

provide a likelihood-ratio test to evaluate the equivalence of the shape parameters $\alpha_1, \alpha_2, \dots, \alpha_m$ from the Weibull populations under the hybrid censoring scheme. In this way, the following hypothesis testing is proposed as

$$H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_m = \alpha \text{ versus } H_1 : \alpha_i \neq \alpha_j \text{ for } i \neq j.$$

Under a large sample size N , the likelihood-ratio statistics

$$-2\{\ell(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m, \hat{\alpha}) - \ell(\hat{\lambda}_1, \dots, \hat{\lambda}_m, \hat{\alpha}_1, \dots, \hat{\alpha}_m)\} \sim \chi_1^2,$$

where $\ell(\lambda, \alpha)$ is the log-likelihood function at (6) when parameters are common. Thus, the asymptotic distribution of $-2\{\ell(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m, \hat{\alpha}) -$

$\ell(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m, \hat{\alpha}_1, \hat{\alpha}_2, \dots, \hat{\alpha}_m) \sim \chi_1^2$ can be used to construct the likelihood ratio test and reject H_0 under this case if $-2\{\ell(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m, \hat{\alpha}) - \ell(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m, \hat{\alpha}_1, \hat{\alpha}_2, \dots, \hat{\alpha}_m)\} \sim \chi_1^2 > \xi^*$ where ξ^* is such that $P(\chi_1^2 > \xi^*)$ is equal to the size of the test.

We use this likelihood ratio testing to compare the equality of the shape parameters. We test the hypothesis $H_0 : \alpha_1 = \alpha_2 = \alpha_3 = \alpha$ versus $H_1 : \alpha_i \neq \alpha_j$ for $i, j = 1, 2, 3$. The likelihood ratio statistics and the associated p-value for the hypothesis $\alpha_1 = \alpha_2 = \alpha_3$ with %5 level of significance are obtained as 1.3604 and 0.2435 (> 0.05). These results imply that the null hypothesis cannot be rejected. Therefore, it is recommended that the shape parameters α_1 , α_2 , and α_3 are equal for this breaking strength of jute fiber data.

After proving the availability of the data, we consider the following nine different censoring schemes given in Table 11 and we perform the estimation methods. In Bayesian calculations, we prefer informative hyper-parameter values and use the MLEs of the parameters. We take $a_1 = \hat{\alpha}$, $b_1 = 1$, $c_i = \hat{\lambda}_i / \sum_{i=1}^3 \hat{\lambda}_i$, ($i = 1, 2, 3$) to provide expected values of the prior distributions as discussed in Section 3. For Bayesian estimations, we employ MCMC with 100,000 iterations. The burn-in phase is not used in Bayesian calculations because the initial values of the parameters are their MLEs. Using the thinning process, we construct uncorrelated Markov Chains by taking each 10th variate. Then, we report the estimation results in Table 12.

Table 11. Different censoring schemes for jute fiber data.

CS	s_i	T_i
I, II, III	$\%60 \times n_i$	$(0.5, 0.6, 0.5), (0.6, 0.5, 0.7)$ and $(0.7, 0.7, 0.7)$
IV, V, VI	$\%75 \times n_i$	$(0.5, 0.6, 0.5), (0.6, 0.5, 0.7)$ and $(0.7, 0.7, 0.7)$
$VII, VIII, IX$	$\%90 \times n_i$	$(0.5, 0.6, 0.5), (0.6, 0.5, 0.7)$ and $(0.7, 0.7, 0.7)$

Table 12. Inferences for the parameters of jute fiber data under different censoring schemes.

	MLE	Bayes	ACI	HPD	MLE	Bayes	ACI	HPD	MLE	Bayes	ACI	HPD
	CS_I				CS_{II}				CS_{III}			
α	1.2779	0.9480	0.6118	0.4448	1.2779	1.2037	0.6118	0.5260	1.2779	1.2030	0.6118	0.5373
λ_1	3.0673	1.5370	3.7651	1.3949	3.0673	2.6827	3.7651	2.3924	3.0673	2.7364	3.7651	2.4915
λ_2	2.6025	1.3421	2.9842	1.2353	2.6025	2.3081	2.9842	2.0325	2.6025	2.3148	2.9842	2.0851
λ_3	3.4106	1.6961	4.2217	1.6115	3.4106	3.0028	4.2217	2.5751	3.4106	3.0488	4.2217	2.7532
	CS_{IV}				CS_V				CS_{VI}			
α	1.2371	1.0345	0.5361	0.4480	1.2692	1.3318	0.5346	0.5302	1.2921	1.3121	0.5308	0.5078
λ_1	2.6805	1.7566	2.8414	1.5256	3.0778	3.4700	3.0736	2.7569	3.1438	3.3238	3.1296	2.5506
λ_2	2.8192	1.8746	2.7437	1.4984	2.7187	3.0922	2.8340	2.4951	2.9508	3.0935	2.8556	2.2433
λ_3	2.8876	1.8982	3.0547	1.5681	3.0916	3.4786	3.0107	2.6596	3.1513	3.3192	3.0586	2.3373
	CS_{VII}				CS_{VIII}				CS_{IX}			
α	1.2630	1.0959	0.5348	0.4569	1.2785	1.2757	0.5125	0.4889	1.3505	1.4993	0.4966	0.4929
λ_1	2.7481	1.9822	2.9076	1.6828	2.8520	2.8625	2.7190	2.1720	3.1329	4.0453	2.7353	2.7010
λ_2	3.0071	2.2116	2.7616	1.7308	2.7421	2.7771	2.8284	2.1178	3.3037	4.3126	2.8335	2.7918
λ_3	2.9602	2.1178	3.1246	1.7388	3.3397	3.3575	2.9008	2.4163	3.5202	4.5673	3.0225	2.8934

Table (12) shows us consistent results. We see from the average lengths of the approximate confidence intervals that the Bayesian method provides better results.

We also verify that the Markov chains produced for the jute fiber data are convergent. First, we create trace plots that plot the parameter values against the number of iterations at each stage. The trace plots for the parameters shown in Figure 5 demonstrate how similarly the Markov chains for each parameter vary about their centers. In addition, the density plots appear to have symmetric and unimodal forms, as seen

in Fig. 6. The values having the strongest support from the data and the selected priors are the unimodal peaks in these plots, which identify the mode of the posterior distributions of the parameters. The convergency control plots of a Markov chain can be drawn using the package "mcmcplots" Curtis, Goldin, and Evangelou (2018) in R R Core Team (2023) software.

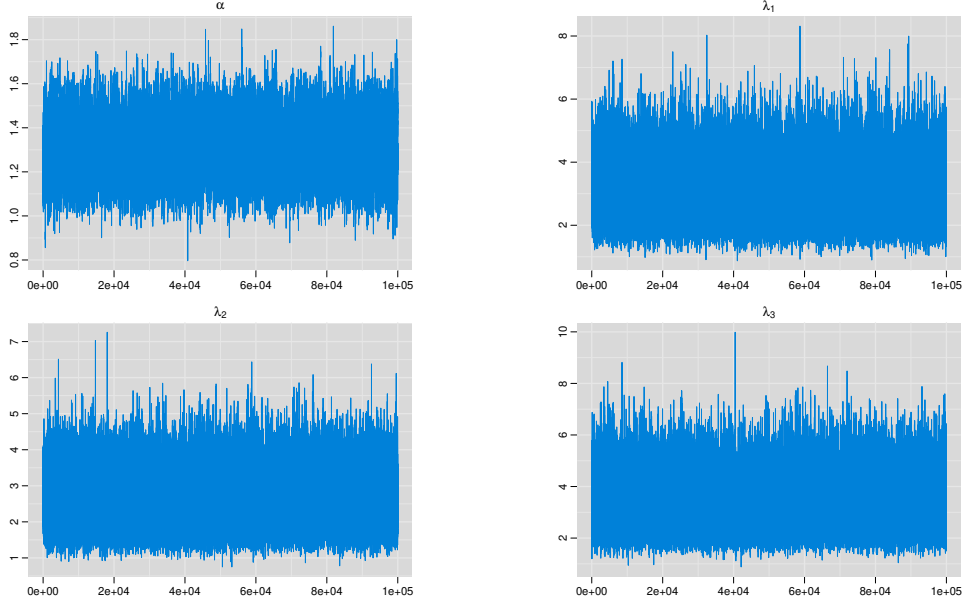


Figure 5. Trace plots of the posterior distribution the parameters.

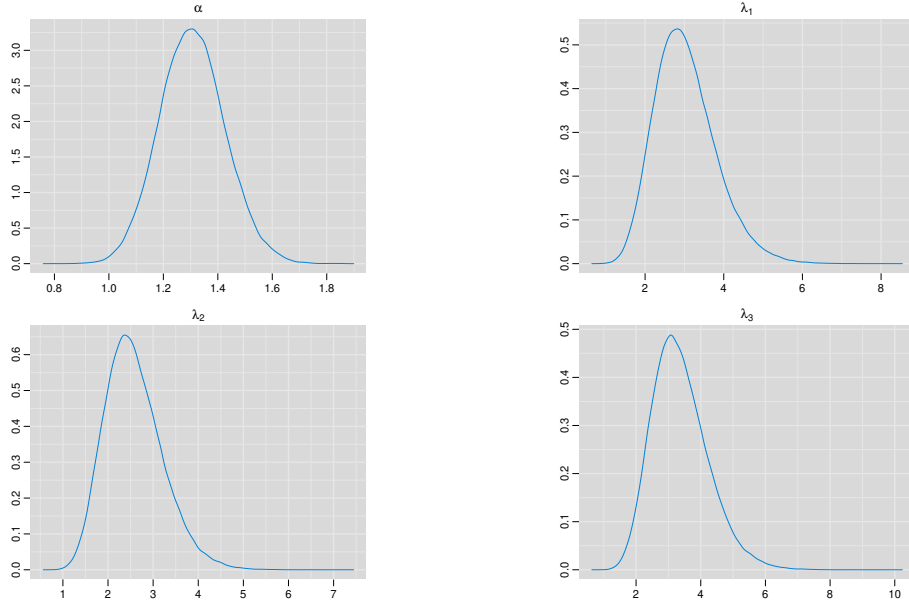


Figure 6. Density plots of the posterior distributions of α and λ_i 's.

7. Conclusion

In this study, we considered m blocks using the Type-I censoring method. We assumed Weibull distributions with a common shape but distinct scale parameters for each set. Our analysis primarily centered on parameter estimation in this context, specifically examining the scale parameters in two distinct scenarios. The first case involved considering the scale parameters without any restrictions, while the second case incorporated order restrictions. The main concepts of the study are constructed in the unrestricted case, but we also provide results and examples for the order-restricted scenario. In addition to employing the classical inference method, maximum likelihood estimation, we also explore Bayesian inference techniques. For Bayesian analysis, we adopt a dependent prior for the parameters based on the approach outlined by Peña and Gupta (1990). Thus, we employ a gamma-Dirichlet prior for the Bayesian analysis. All estimates are obtained with their approximate confidence intervals. Our comprehensive analysis indicates that the inference procedures are thoroughly examined through simulations and numerical examples. Theoretical outcomes demonstrate robust performance across various censoring schemes and sample sizes. Comparing the results between the unrestricted and restricted cases, we find that the order-restricted cases provide smaller biases, MSEs, and average length of the confidence intervals, but with slightly lower coverage probabilities compared to the unrestricted cases.

Acknowledgements

The authors would like to thank two unknown reviewers and the associate editor for their constructive suggestions that have helped improve the earlier version of the manuscript.

References

- Ahmadi, M. V., Doostparast, M., & Ahmadi, J. (2018). Block censoring scheme with two-parameter exponential distribution. *Journal of Statistical Computation and Simulation*, 88(7), 1229-1251.
- Arnold, B. C., & Press, S. J. (1983). Bayesian inference for Pareto populations. *Journal of Econometrics*, 21(3), 287-306.
- Brunk, H., Barlow, R. E., Bartholomew, D. J., & Bremner, J. M. (1972). Statistical inference under order restrictions.(the theory and application of isotonic regression). *International Statistical Review*, 41, 395.
- Çetinkaya, Ç. (2021). Reliability estimation of a stress-strength model with non-identical component strengths under generalized progressive hybrid censoring scheme. *Statistics*, 55(2), 250-275.
- Chen, M.-H., & Shao, Q.-M. (1999). Monte carlo estimation of Bayesian credible and HPD intervals. *Journal of Computational and Graphical Statistics*, 8(1), 69-92.
- Congdon, P. (2001). *Bayesian statistical modelling*. New York: John Wiley & Sons.
- Curtis, S. M., Goldin, I., & Evangelou, E. (2018). Package ‘mcmcplots’ [Computer software manual]. Retrieved from <https://CRAN.R-project.org/package=mcmcplots> (R package version 0.4.3)
- Dey, S., Elshahhat, A., & Nassar, M. (2023). Analysis of progressive type-II censored gamma distribution. *Computational Statistics*, 38(1), 481-508.
- Epstein, B. (1954). Truncated life tests in the exponential case. *The Annals of Mathematical Statistics*, 555-564.

- Greene, W. H. (2012). *Econometric analysis*. Upper Saddle River, USA.: Prentice Hall.
- Gupta, R. D., & Kundu, D. (2006). On the comparison of Fisher information of the Weibull and ge distributions. *Journal of Statistical Planning and Inference*, 136(9), 3130–3144.
- Johnson, L. G. (1964). *Theory and technique of variation research*. Amsterdam: Elsevier.
- Krishna, H., Dube, M., & Garg, R. (2019). Estimation of stress strength reliability of inverse Weibull distribution under progressive first failure censoring. *Austrian Journal of Statistics*, 48(1), 14–37.
- Kumari, R., Tripathi, Y. M., Sinha, R. K., & Wang, L. (2023). Reliability estimation for bathtub-shaped distribution under block progressive censoring. *Mathematics and Computers in Simulation*.
- Pal, A., Mitra, S., & Kundu, D. (2021). Order restricted classical inference of a Weibull multiple step-stress model. *Journal of Applied Statistics*, 48(4), 623–645.
- Peña, E. A., & Gupta, A. K. (1990). Bayes estimation for the Marshall–Olkin exponential distribution. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 52(2), 379–389.
- R Core Team. (2023). R: A language and environment for statistical computing [Computer software manual]. Vienna, Austria. Retrieved from <https://www.R-project.org/>
- Saraçoğlu, B., Kınacı, I., & Kundu, D. (2012). On estimation of $R = P(Y < X)$ for exponential distribution under progressive type-II censoring. *Journal of Statistical Computation and Simulation*, 82(5), 729–744.
- Sarhan, A. E., & Greenberg, B. G. (1962). *Contributions to order statistics*. Wiley New York.
- Singh, K., Tripathi, Y. M., Lodhi, C., & Wang, L. (2024). Inference for unit inverse weibull distribution under block progressive type-ii censoring. *Journal of Statistical Theory and Practice*, 18(3), 42.
- Singh, K., Tripathi, Y. M., Wang, L., & Wu, S.-J. (2024). Analysis of block adaptive type-ii progressive hybrid censoring with weibull distribution. *Mathematics*, 12(24), 1–21.
- Tierney, L. (1994). Markov chains for exploring posterior distributions. *the Annals of Statistics*, 1701–1728.
- US Department of Defense. (1977). *Reliability design qualification and production acceptance tests: Exponential distribution, MIL-STD-781C*. Washington, DC.
- Wang, L., Wu, S.-J., Lin, H., & Tripathi, Y. M. (2023). Inference for block progressive censored competing risks data from an inverted exponentiated exponential model. *Quality and Reliability Engineering International*.
- Wu, J.-W., Tsai, T.-R., & Ouyang, L.-Y. (2001). Limited failure-censored life test for the weibull distribution. *IEEE Transactions on Reliability*, 50(1), 107–111.
- Xia, Z., Yu, J., Cheng, L., Liu, L., & Wang, W. (2009). Study on the breaking strength of jute fibres using modified Weibull distribution. *Composites Part A: Applied Science and Manufacturing*, 40(1), 54–59.
- Zhu, T. (2020). Reliability estimation for two-parameter Weibull distribution under block censoring. *Reliability Engineering & System Safety*, 203, 107071.

Appendix A

It is demonstrable that, $\hat{\alpha}$ and $\hat{\lambda}_i$ maximize $\ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})$ for given α . For this purpose, let assume that $\xi(\lambda)$ be the Hessian matrix of $\ell(\alpha, \boldsymbol{\lambda}, \mathbf{X})$ at $(\hat{\alpha}, \hat{\lambda}_i)$. Thus,

$$\xi_{ii}(\lambda) = \frac{\partial^2 \ell}{\partial \lambda_i^2} = -\frac{k_i}{\lambda_i^2}$$

and

$$\xi_{ij}(\lambda) = \frac{\partial^2 \ell}{\partial \lambda_i \partial \lambda_j} = 0, \quad i, j = 1, 2, \dots, m.$$

Thus, the determinant of the Hessian matrix is acquired as

$$\det[\xi(\hat{\lambda})] = \xi_{11}(\hat{\lambda}_1) \xi_{22}(\hat{\lambda}_2) \cdots \xi_{mm}(\hat{\lambda}_m) = \frac{k_1 k_2 \cdots k_m}{\hat{\lambda}_1^2 \hat{\lambda}_1^2 \cdots \hat{\lambda}_m^2} > 0.$$

It is clear that $(\hat{\lambda}_1 \hat{\lambda}_2 \cdots \hat{\lambda}_m)$ is the local maximum of $\ell(\alpha, \mathbf{\lambda}, \mathbf{X})$ for given α . As there isn't a single point of $\ell(\alpha, \mathbf{\lambda}, \mathbf{X})$ and it has just one crucial point, $\hat{\alpha}$ and $\hat{\lambda}_i$ are the absolute maximum of the log-likelihood function.

On the other hand, by equating (7) equal to zero, we obtain

$$\frac{\sum_{i=1}^m k_i}{\alpha} = \sum_{i=1}^m \hat{\lambda}_i \Delta'_i(\alpha) - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}), \quad (1)$$

where

$$\Delta'_i(\alpha) = \sum_{j=1}^{k_i} x_{ij}^\alpha \ln x_{ij} + (n_i - k_i) U_i^\alpha \ln U_i.$$

We can denote the each sides of equation (1) by $\psi_1(\alpha, \mathbf{X})$ and $\psi_2(\alpha, \mathbf{X})$ respectively as follows

$$\psi_1(\alpha, \mathbf{X}) = \frac{\sum_{i=1}^m k_i}{\alpha}$$

and

$$\psi_2(\alpha, \mathbf{X}) = \sum_{i=1}^m \hat{\lambda}_i \Delta'_i(\alpha) - \sum_{i=1}^m \sum_{j=1}^{k_i} \ln(x_{ij}).$$

For a given sample of $\mathbf{X}$, it can be shown that $\psi_2(\alpha, \mathbf{X})$ is a monotone increasing function of α with a finite and positive limit as $\alpha \rightarrow \infty$. The plots of $\psi_1(\alpha, \mathbf{X})$ and $\psi_2(\alpha, \mathbf{X})$ would intersect exactly once at $\hat{\alpha}$ since $\frac{\sum_{i=1}^m k_i}{\alpha}$ is strictly decreasing with right limit ∞ at 0. For proof, it can be shown that $\frac{\partial \psi_2(\alpha, \mathbf{X})}{\partial \alpha} \geq 0$. Following the Cauchy-Schwarz inequality, it is clearly seen that

$$\frac{\partial \psi_2(\alpha, \mathbf{X})}{\partial \alpha} = \sum_{i=1}^m \hat{\lambda}_i \left(\sum_{j=1}^{k_i} x_{ij}^\alpha \ln^2 x_{ij} + (n_i - k_i) U_i^\alpha \ln^2 U_i \right) \geq 0.$$

This establishes the required property that $\psi_2(\alpha, \mathbf{X})$ is a monotone increasing function of α . It should be noted here that there exists a finite upper limit for $\psi_2(\alpha, \mathbf{X})$, namely $\lim_{\alpha \rightarrow \infty} \psi_2(\alpha, \mathbf{X}) > 0$ indicates that the curves of $\psi_1(\alpha, \mathbf{X})$ and $\psi_2(\alpha, \mathbf{X})$ have a unique intersection point. Therefore, $\hat{\alpha}$, as the root of equation $\psi_1(\alpha, \mathbf{X}) = \psi_2(\alpha, \mathbf{X})$ exist and

unique. Following the obtaining unique MLE of α , the MLE of α and λ_i can be obtained uniquely as $\hat{\lambda}_i = \hat{\lambda}_i(\hat{\alpha})$. Thus, the existence and uniqueness of the MLEs are proved.

Appendix B

Following the generalized isotonic regression problem introduced by Brunk, Barlow, Bartholomew, and Bremner (1972) and given in Theorem (A.1) by Pal et al. (2021), for fixed α we aim maximizing the following likelihood function

$$L(\lambda_1, \lambda_2, \dots, \lambda_m | \alpha, \Delta) \propto \alpha^{\sum_{i=1}^m k_i} \exp \left\{ - \sum_{i=1}^m \lambda_i \Delta_i(\alpha) \right\} \left(\prod_{i=1}^m \lambda_i^{k_i} \right) \left(\prod_{i=1}^m \prod_{j=1}^{k_i} x_{ij}^\alpha \right), \quad (2)$$

with respect to λ_i for $i = 1, 2, \dots, m$ based on the order restriction $\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m$. Then,

$$\max_{\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m} L(\lambda_1, \lambda_2, \dots, \lambda_m | \alpha, \Delta) = \max_{\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m} \left\{ \prod_{i=1}^m \lambda_i^{k_i} e^{-\lambda_i \Delta_i(\alpha)} \right\}.$$

As $\log(x)$ increases monotonically, this is equivalent to

$$\begin{aligned} \max_{\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m} \left\{ \sum_{i=1}^m k_i \ln \lambda_i - \lambda_i \Delta_i(\alpha) \right\} &= \max_{\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_m} \left\{ \sum_{i=1}^m \left[-k_i \ln \left(\frac{1}{\lambda_i} \right) - \frac{\Delta_i(\alpha)}{\frac{1}{\lambda_i}} \right] \right\} \\ &= \max_{1/\lambda_1 \geq 1/\lambda_2 \geq \dots \geq 1/\lambda_m} \left\{ \sum_{i=1}^m \left[-\ln \left(\frac{1}{\lambda_i} \right) - \frac{\Delta_i(\alpha)}{\frac{k_i}{\lambda_i}} \right] k_i \right\} \\ &= \min_{1/\lambda_1 \geq 1/\lambda_2 \geq \dots \geq 1/\lambda_m} \left\{ \sum_{i=1}^m \left[\ln \left(\frac{1}{\lambda_i} \right) + \frac{\Delta_i(\alpha)}{\frac{k_i}{\lambda_i}} \right] k_i \right\}. \end{aligned}$$

Then, consider the convex function $\Phi(u) = -\log u$, $\xi_\Phi(g(x), f(x)) = -\log g(x) + \log f(x) + (g(x) - f(x))/f(x)$. Now, we aim

$$\text{Minimize } \sum_{x \in A} \xi_\Phi[g(x), f(x)] w(x),$$

subject to $f(x)$ is isotonic, we can write it as

$$\text{Minimize } \sum_{x \in A} \left[\log f(x) + \frac{g(x)}{f(x)} \right] w(x),$$

subject to $f(x)$ is isotonic. If we take $A = 1, 2, \dots, m$, $g(i) = \Delta_i(\alpha)/n_i$, $w(i) = n_i$, $f(i) = 1/\lambda_i$, $i = I$. Using the Theorem (A.1) in Pal et al. (2021), the generalized isotonic regression of $g(x)$ solves the problem as

$$\frac{1}{\tilde{\lambda}_u(\alpha)} = \left[\min_{s \leq u} \max_{t \geq u} \frac{\sum_{r=s}^t \Delta_r(\alpha)}{\sum_{r=s}^t k_r} \right], \quad s, u, t \in I.$$

Thus, the result follows from the invariance property of the MLEs.