

ON SEMI-PARAMETRIC PROGRESSIVE CENSORING

Deepak Prajapati ^{*} and Debasis Kundu [†]

Abstract

This article concerns classical and Bayesian inference of the unknown parameters of a semi-parametric proportional hazard's model when the data are progressively censored. Wang et al. (2010) considered this problem for specific parametric proportional hazard's model with Weibull, Lomax and Gompertz baseline distributions. The aim of this paper is not to assume any specific parametric family of distributions. Instead, it is assumed that the baseline distribution has a piecewise constant hazard function with a given number of cut points. It becomes more flexible than any specific parametric class of distribution functions. The maximum likelihood estimators of the unknown parameters can be obtained in closed form. Order restricted maximum likelihood estimators have also been proposed. A very flexible Dirichlet gamma priors has been assumed on the unknown parameters. Based on these priors, the Bayes estimates, ordered restricted Bayes estimates, and the associated credible intervals are also provided. Simulation experiments have been performed to illustrate the effectiveness of the proposed method. In practice, the cut points may not be known. The choice of the cut points is an important issue. We have discussed the choice of the cut points based on a non-homogeneous Poisson process model. Finally, two data sets are analyzed for illustrative purposes.

KEY WORDS AND PHRASES: Type-II censoring; Bayes estimators; Credible intervals; Bayes factor; Maximum likelihood estimators.

AMS SUBJECT CLASSIFICATIONS: 62F10, 62F03, 62H12.

^{*}Decision Sciences Area, Indian Institute of Management Lucknow, IIM Rd, Prabandh Nagar, Mubarakpur, Lucknow, Uttar Pradesh 226013. Email id: deepak.prajapati@iiml.ac.in

[†]Department of Mathematics and Statistics, Indian Institute of Technology Kanpur, Pin 208016, India. E-mail: kundu@iitk.ac.in, Phone no. 91-512-2597141, Fax no. 91-512-2597500.

1 INTRODUCTION

Censoring is inevitable in any life testing experiment. For several reasons, it is usually not possible to wait till all the items fail. Type-I and Type-II censoring are the most popular censoring schemes which are being used in practice. In a Type-I censoring scheme the experiment is stopped at a prespecified time point, due to this reason it is also known as the time censoring. The main advantage of the Type-I censoring scheme is that the experimental time is decided before hand. The main disadvantage of the Type-I censoring scheme is that one may observe very few failed items during the experiment, hence the inference based on the lifetime of the failed items may not be very reliable. In Type-II censoring scheme, the experiment is stopped when a pre-fixed number of failures is observed. Hence, in this case the number of failures is fixed and it is based on experimenter's choice, but in this case the experimental time can be quite large. A mixture of Type-I and Type-II censoring schemes is known as hybrid censoring scheme. For a detailed discussions of different censoring schemes such as left censoring, right censoring, interval censoring etc., interested readers are referred to the book by Balakrishnan et al. (2023).

In the last two decades, progressive censoring scheme has received a considerable amount of attention in the statistical literature. One major advantage of the progressive censoring scheme over other censoring schemes is that it allows the experimenter to remove live items during the experiment, which may be used for other purposes. It is a more general censoring scheme than traditional Type-II censoring, and the Type-II censoring scheme can be obtained as a special case of the progressive censoring scheme. Estimation of the unknown parameters, generation of samples from a progressive censoring scheme, different optimal censoring schemes along with different issues associated with the progressive censoring schemes have been discussed by several authors. The book by Balakrishnan and Cramer (2014) provided a detailed review on different aspects of the progressive censoring schemes. For more recent

work, see [Maurya et al. \(2019\)](#), [Alshenawy et al. \(2022\)](#), and the references cited therein.

Interestingly, most of the work related to progressive censoring scheme are parametric in nature. Not much work has been done under the non-parametric or semi-parametric setup. [Wang et al. \(2010\)](#) considered a specific class of parametric family of distributions, and have developed the classical inferential procedures of the unknown parameters for these classes as proposed by [Cox \(1959\)](#). They considered three specific members of this class, namely Weibull, Lomax (Pareto II) and Gompertz distributions. [Wang et al. \(2010\)](#) developed different inferential procedures of the unknown parameters based on Type-II progressively censored samples for these distributions.

The main aim of this paper is not to assume any specific parametric forms of the baseline survival function. On the other hand, we do not want to make it purely non-parametric. We assume that the baseline survival function has piecewise-constant hazard function with a given number of known cut points. This makes the model very flexible. One need not choose the baseline distribution before hand so that it has a specific hazard function shape, for example increasing, decreasing, bathtub shaped, unimodal etc. The aim is to develop a data-driven approach, and from the data one should be able to choose the baseline distribution in such a manner that it has the appropriate hazard function, which matches with the data. It makes the baseline distribution functions to be a very flexible one, and the corresponding probability density function (PDF) can take variety of shapes. It may be mentioned that the assumption of piecewise-constant hazard function common in the lifetime data analysis, see for example, [Murphy \(1996\)](#), [Balakrishnan et al. \(2016\)](#), [Goodman et al. \(2011\)](#), [Kundu \(2022a\)](#), [Samanta and Kundu \(2023\)](#), [Ganguly et al. \(2024\)](#) and the references cited therein.

We develop the classical inference of the unknown parameters. The maximum likelihood estimators (MLEs) of the unknown parameters can be obtained in closed form. We provide a convenient way to generate samples from this model that can be used to construct bootstrap

confidence intervals of the unknown parameters. It may be known apriori that there is an ordering in the baseline hazard function. We provide the order-restricted inference of the unknown parameters. An algorithm, [similar to the algorithm proposed by Samanta et al. \(2019\) for the step-stress model](#), has been provided which converges in a finite number of steps to compute the MLEs of the unknown parameters under the ordered restriction. We prove that the algorithm provides the solutions which maximize the likelihood function.

We further develop the Bayesian inference of the unknown parameters under a very flexible set of priors. It is assumed that for the given cut points, the piecewise constant hazards follow a Dirichlet-gamma (DG) distribution. Based on the above priors, the Bayes estimates and the associated credible intervals can not be obtained in closed form. We develop the Bayes estimates and the associated credible intervals based on the importance sampling technique. We also provide ordered-restricted Bayesian inference based on ordered DG prior. Extensive simulations have been performed to see the effectiveness of the proposed method.

Note that the choice of the cut points is an important issue. In this paper we assume [similar to the assumption made by Samanta and Kundu \(2025\)](#) that the cut points occur according to a non-homogeneous Poisson process (NHPP) with the intensity function follows a power law process (PLP). The cut points can be obtained as the maximum likelihood predictor of the occurrence of the events of a non-homogeneous Poisson process. Two data sets are analyzed for illustrative purposes.

The rest of the paper is organized as follows. We discuss the problem and provide the model and the priors in Section 2. The maximum likelihood estimators with and without ordering, along with the associated confidence intervals, are provided in Section 3. The posterior analysis, the Bayes estimates and the associated credible intervals are provided in Section 4. Simulation results are presented in Section 5. In Section 6 we discuss the choice of the cut points based on non-homogeneous Poisson process and analyze two data sets to

show how the proposed method can be used in practice. Finally we conclude the paper in Section 7.

2 PROBLEM FORMULATION AND ASSUMPTION

2.1 PROGRESSIVE CENSORING SCHEME

Suppose n identical units are put on a life testing experiment. Denote the lifetimes of the n items as $T_1, \dots, T_n$, respectively, and it is assumed that $T_1, \dots, T_n$ are independent and identically distributed (i.i.d.) absolutely continuous random variables with the probability density function (PDF) $f(t)$, cumulative distribution function (CDF) $F(t)$ and survival function (SF) $S(t)$. The integer $m < n$ and the integers $R_1, \dots, R_m$ are prefixed such that $R_1 + \dots + R_m = n - m$. At the time of the first failure at time t_1 , R_1 out of $(n - 1)$ remaining surviving units are removed. Similarly, at the time of the second failure at time t_2 , R_2 out of $(n - R_1 - 2)$ remaining surviving units are removed, and so on. Finally at the time of the m -th failure, at time t_m , all the remaining $R_m = n - R_1 - \dots - R_{m-1} - m$ surviving units are removed and the life test stops. It may be observed that when $R_1 = \dots = R_{m-1} = 0$, then $R_m = n - m$ and Type-II progressive censoring scheme becomes the traditional Type-II censoring scheme.

2.2 PROPORTIONAL HAZARD CLASS

Recall that a class of distribution functions $\{F(t; \theta) : \theta > 0\}$ is said to be a proportional hazards class if its survival function $S(t; \theta)$ can be written as

$$S(t; \theta) = (S_0(t))^\theta; \quad t > 0.$$

The associated PDF and the hazard function can be written as

$$f(t; \theta) = \theta(S_0(t))^{\theta-1} f_0(t) = \theta h_0(t)(S_0(t))^\theta \quad \text{and} \quad h(t; \theta) = \theta h_0(t),$$

respectively. Here, $S_0(t)$ is known as the baseline survival function; similarly, $f_0(t)$ and $h_0(t)$ are the baseline PDF and the hazard function, respectively. Throughout this article it is assumed that $S_0(t) = 1$, for $t \leq 0$, although all the results can be extended for general S_0 also. Most of the results are valid even when $S_0(t)$ has bounded support also. The baseline distribution often has one or more parameters. Several well-known distributions functions belong to the proportional hazard class, for example exponential, Weibull, Rayleigh, generalized Pareto, Lomax, Gompertz, etc. An extensive amount of work has been done in developing properties and establishing inference procedures of several univariate proportional hazard classes, see for example [Cox and Oakes \(1984\)](#) in this respect. [It may be mentioned that although the proportional hazard class of distributions have been used extensively in analyzing lifetime data in presence of covariates, it can be used as a flexible model even without covariates, see for example Wang et al. \(2010\), Kundu \(2022b\) and the references cited there in.](#)

2.3 SEMIPARAMETRIC MODEL

It is assumed that the baseline hazard function $h_0(t)$ is piecewise constant, and it is based on the following: the number of constants, M and the cut points $\tau_0 = 0 < \tau_1 < \dots < \tau_{M-1} < \tau_M = \tau$ are assumed to be known. [Here \$\tau \geq t_m\$ and \$M < m\$.](#) Later we will discuss how to estimate these constants. Further, it is assumed that $h_0(t)$ has the following form:

$$h_0(t) = \sum_{j=1}^M \theta_j I_{[\tau_{j-1}, \tau_j)}(t), \tag{1}$$

where the θ_j 's are unknown positive constants and $I_{[\tau_{j-1}, \tau_j)}(t)$ is the indicator function, i.e.

$$I_{[\tau_{j-1}, \tau_j)}(t) = \begin{cases} 1 & t \in [\tau_{j-1}, \tau_j) \\ 0 & \text{otherwise.} \end{cases}$$

It is clear that under this set up $\theta_1, \dots, \theta_M$ are identifiable only up to a multiplicative constant. Hence, without loss of generality, it is assumed that $\theta_1 + \dots + \theta_M = 1$. Notice that when $M = 1$, $h_0(t)$ becomes an exponential baseline hazard. Now let us compute $S_0(t)$ and $f_0(t)$ which will be needed later. The cumulative baseline hazard function $H_0(t)$ can be written as

$$H_0(t) = \sum_{j=1}^M \theta_j (\min\{t, \tau_j\} - \tau_{j-1}) I_{[\tau_{j-1}, \tau_M)}(t).$$

From now on we will use the notation $A_j = [\tau_{j-1}, \tau_j)$, for $j = 1, \dots, M$. The cumulative hazard function can be written as

$$H_0(t) = \begin{cases} \theta_1 t & \text{if } t \in A_1 \\ \theta_1 \tau_1 + \theta_2 (t - \tau_1) & \text{if } t \in A_2 \\ \vdots & \vdots \\ \sum_{j=1}^{k-1} \theta_j (\tau_j - \tau_{j-1}) + \theta_k (t - \tau_{k-1}) & \text{if } t \in A_k \\ \vdots & \vdots \\ \sum_{j=1}^{M-1} \theta_j (\tau_j - \tau_{j-1}) + \theta_M (t - \tau_{M-1}) & \text{if } t \in A_M. \end{cases} \quad (2)$$

Once $H_0(t)$ is obtained, $S_0(t)$ and $f_0(t)$ can be obtained as

$$S_0(t) = e^{-H_0(t)} \quad \text{and} \quad f_0(t) = h_0(t) e^{-H_0(t)}.$$

2.4 PRIOR ASSUMPTIONS

Dirichlet-Gamma Prior

It is assumed that $\theta_1, \dots, \theta_M$ has the following Dirichlet prior

$$\pi_1(\theta_1, \dots, \theta_M | c) = \frac{1}{B(c)} \prod_{i=1}^M \theta_i^{c_i-1}, \quad (3)$$

for $0 < \theta_1, \dots, \theta_M < 1$ and $\sum_{i=1}^M \theta_i = 1$. Here $c = (c_1, \dots, c_M)$, $c_1 > 0, \dots, c_M > 0$, and $B(c) = \left\{ \prod_{i=1}^M \Gamma(c_i) \right\} / \Gamma\left(\sum_{i=1}^M c_i\right)$ is the normalizing constant. It will be denoted by $D(c_1, \dots, c_M)$.

It is assumed that θ has a gamma(a, b) prior where $a > 0$, $b > 0$, and it has the following probability density function

$$\pi_2(\theta|a, b) = \frac{b^a}{\Gamma(a)} \theta^{a-1} e^{-b\theta}; \quad \theta > 0.$$

Further, $(\theta_1, \dots, \theta_M)$ and θ are assumed to be independent. If we denote, $\alpha_1 = \theta\theta_1, \dots, \alpha_M = \theta\theta_M$, then $0 < \alpha_1, \dots, \alpha_M < \infty$. The prior on $\alpha_1, \dots, \alpha_M$ becomes

$$\pi_0(\alpha_1, \dots, \alpha_M|a, b, c) = \frac{\Gamma(\bar{c})}{\Gamma(a)} (b\bar{\alpha})^{a-\bar{c}} \times \prod_{i=1}^M \frac{b^{c_i}}{\Gamma(c_i)} \alpha_i^{c_i-1} e^{-b\alpha_i}. \quad (4)$$

Here, $\bar{c} = c_1 + \dots + c_M$ and $\bar{\alpha} = \alpha_1 + \dots + \alpha_M$. The prior (4) is known as the Dirichlet-gamma (DG) prior, and it will be denoted by $DG(a, b, c_1, \dots, c_M)$. It is a very flexible prior. Depending on the hyper-parameters, α_i 's can be independent, positively and negatively correlated. DG prior was originally proposed by [Peña and Gupta \(1990\)](#) as a conjugate prior of the parameters of the Marshall–Olkin bivariate exponential distribution. Later, [Kundu and Gupta \(2013\)](#) provided an effective way of generating samples from the DG distribution.

Ordered Dirichlet-Gamma Prior

Recently, [Mondal and Kundu \(2019\)](#) proposed an ordered Dirichlet-gamma distribution as an effective prior in case of an order-restricted parameter set. Suppose the parameters $\alpha_1, \dots, \alpha_M$ has the following order restriction $\alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_M$. In this case the ordered Dirichlet-gamma prior on the parameter space has the following form:

$$\pi_1(\alpha_1, \dots, \alpha_M|a, b, c) = \frac{b^a}{\Gamma(a)} \times \frac{\Gamma(\bar{c})(\bar{\alpha})^{a-\bar{c}} e^{-b\bar{\alpha}}}{\prod_{i=1}^M \Gamma(c_i)} \times \sum_{i_1, \dots, i_M \in \mathcal{P}} \left(\alpha_1^{c_{i_1}-1} \alpha_2^{c_{i_2}-1} \dots \alpha_M^{c_{i_M}-1} \right), \quad (5)$$

for $\infty > \alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_M > 0$, and zero, otherwise. Here $\mathcal{P}$ denotes the class of all permutations of $\{1, 2, \dots, M\}$ and $(i_1, \dots, i_M)$ is an element of $\mathcal{P}$. We call this prior (5) as the ordered Dirichlet-gamma (ODG) prior, and it will be denoted by $ODG(a, b, c_1, \dots, c_M)$. It may be noted that if $(\alpha_1, \dots, \alpha_M)$ has the joint PDF (4), then the order statistics of $(\alpha_1, \dots, \alpha_M)$ has the joint PDF (5). [Mondal and Kundu \(2019\)](#) provided an effective method

of generating random samples from (5), and it will be used to compute order-restricted Bayes estimators of $\alpha_1, \dots, \alpha_M$ and their associated credible intervals.

3 MAXIMUM LIKELIHOOD ESTIMATORS

3.1 WITHOUT ORDERING

In this section it is assumed there is no ordering on the unknown parameters. It is assumed that n, m and $R_1, \dots, R_m$ are fixed and we have the following observations $\mathcal{D} = \{t_1 < \dots < t_m\}$. Let us assume $n_1 + \dots + n_M = m$, and

$$\begin{aligned} 0 &< t_1 < \dots < t_{n_1} < \tau_1 \\ \tau_1 &< t_{n_1+1} < \dots < t_{n_1+n_2} < \tau_2 \\ &\vdots \\ \tau_{M-1} &< t_{n_1+\dots+n_{M-1}+1} < \dots < t_{n_1+\dots+n_M} < \tau_M. \end{aligned}$$

In this section, denote $H_0(t)$ as in (2) as $H_0(t; \theta_1, \dots, \theta_M)$. Hence, the likelihood function can be written as

$$\begin{aligned} L(\mathcal{D}|\theta_1, \dots, \theta_M, \theta) &\propto \prod_{i=1}^m \{f(t_i)(S(t_i))^{R_i}\} \\ &= \theta^m \prod_{i=1}^m \{h_0(t_i)(S_0(t_i))^{\theta(R_i+1)}\} = \theta^m \theta_1^{n_1} \dots \theta_M^{n_M} \prod_{i=1}^m e^{-\theta(R_i+1)H_0(t_i; \theta_1, \dots, \theta_M)}. \end{aligned} \tag{6}$$

Hence, the MLEs of $\theta_1, \dots, \theta_M$ and θ can be obtained by maximizing (6) with respect to the unknown parameters, subject to the restriction $\theta_1 + \dots + \theta_M = 1$. Let $\alpha_1 = \theta\theta_1, \dots, \alpha_M = \theta\theta_M$. The likelihood function can be written as

$$L(\mathcal{D}|\alpha_1, \dots, \alpha_M) \propto \alpha_1^{n_1} \dots \alpha_M^{n_M} \prod_{i=1}^m e^{-(R_i+1)H_0(t_i; \alpha_1, \dots, \alpha_M)}. \tag{7}$$

Here $0 < \alpha_1, \dots, \alpha_M < \infty$. Let $t_{n_0+n_1+\dots+n_{j-1}+k} = t_{jk}$. Note that t_{jk} represents the k -th observation at the j -th cut point. The associated $R_{n_0+n_1+\dots+n_{j-1}+k} = R_{jk}$. Here $j = 1, \dots, M$, $k = 1, \dots, n_j$ and $n_0 = 0$.

Hence, the log-likelihood function of $\alpha_1, \dots, \alpha_M$, without the additive constant can be written as

$$\begin{aligned} l(\mathcal{D}|\alpha_1, \dots, \alpha_M) &= \sum_{j=1}^M n_j \ln \alpha_j - \sum_{j=1}^M \sum_{k=1}^{n_j} (R_{jk} + 1) H_0(t_{jk}; \alpha_1, \dots, \alpha_M) \\ &= \sum_{j=1}^M n_j \ln \alpha_j - \sum_{j=1}^M \alpha_j \left\{ \sum_{k=1}^{n_j} (R_{jk} + 1)(t_{jk} - \tau_{j-1}) + (\tau_j - \tau_{j-1}) \sum_{k=j+1}^M n_k \right\}, \end{aligned} \quad (8)$$

assuming $\sum_{k=M+1}^M n_k = 0$. Let us use the following notations for brevity $\bar{n}_j = \sum_{k=j+1}^M n_k$, for $j = 1, \dots, M-1$, $\bar{n}_M = 0$ and

$$D_j = \sum_{k=1}^{n_j} (R_{jk} + 1)(t_{jk} - \tau_{j-1}) + (\tau_j - \tau_{j-1})\bar{n}_j; \quad j = 1, \dots, M. \quad (9)$$

Although the D_j 's depend of the data, we do not make it explicit for brevity. Based on the log-likelihood function (8), the MLEs of $\alpha_1, \dots, \alpha_M$ can be obtained in explicit forms, and they are

$$\hat{\alpha}_j = \frac{n_j}{D_j}; \quad j = 1, \dots, M. \quad (10)$$

The MLEs of $\alpha_1, \dots, \alpha_M$ if and only if $n_1 > 0, \dots, n_M > 0$. It should be also mentioned that it is not a very serious restriction in practice. As in practice M is usually not very large say for example 2,3, or 4 and the effective sample size m is quite larger than M . Hence, in most practical situation these restrictions are being satisfied. The construction of the confidence intervals of the unknown parameters can be obtained using the parametric bootstrap method.

3.2 ORDER RESTRICTED

It may be known before hand that the hazard function is monotone. In that case, we have the natural restrictions as

$$O1 : \alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_M \quad \text{or} \quad O2 : \alpha_1 \leq \alpha_2 \leq \dots \leq \alpha_M.$$

Hence, the problem is to maximize (8) with respect to the unknown parameters under the order restriction $O1$ or $O2$. We provide the method under $O1$; the other one can be performed in a similar fashion. Let us make the following reparametrization of the parameters for $i = 2, \dots, M$.

$$\alpha_i = \beta_{i-1} \alpha_{i-1} = \alpha_1 \prod_{j=1}^{i-1} \beta_j.$$

Here $0 < \beta_1, \beta_2, \dots, \beta_{M-1} \leq 1$. Since, there is a one-to-one correspondence between the ordered $\{\alpha_1, \dots, \alpha_M\}$ and the unordered $\{\alpha_1, \beta_1, \dots, \beta_{M-1}\}$, the statistical inferences based on these two sets of parameters will be equivalent.

First let us assume that $n_1 > 0, \dots, n_M > 0$. The log-likelihood function of $\alpha_1, \beta_1, \dots, \beta_{M-1}$ without the additive constant becomes:

$$l(\mathcal{D} | \alpha_1, \beta_1, \dots, \beta_{M-1}) = m \ln \alpha_1 + \sum_{j=1}^{M-1} \bar{n}_j \ln \beta_j - \alpha_1 \left[\sum_{j=1}^M D_j \left\{ \prod_{k=1}^{j-1} \beta_k \right\} \right]. \quad (11)$$

Here $\prod_{k=1}^0 \beta_k = 1$. If $\hat{\alpha}_1^*, \hat{\beta}_1^*, \dots, \hat{\beta}_{M-1}^*$ maximize (11), then they are the solutions of the

following set of equations

$$\begin{aligned}
\frac{m}{\alpha_1} &= \sum_{j=1}^M D_j \left\{ \prod_{k=1}^{j-1} \beta_k \right\} \\
\frac{\bar{n}_1}{\beta_1} &= \alpha_1 \left[D_2 + D_3 \beta_2 + \dots + D_M \prod_{k=2}^{M-1} \beta_k \right] \\
\frac{\bar{n}_2}{\beta_2} &= \alpha_1 \beta_1 \left[D_3 + D_4 \beta_3 + \dots + D_M \prod_{k=1}^{M-1} \beta_k \right] \\
&\vdots \\
\frac{\bar{n}_{M-1}}{\beta_{M-1}} &= \alpha_1 \left\{ \prod_{k=1}^{M-2} \beta_k \right\} D_M.
\end{aligned}$$

Hence

$$\hat{\alpha}_1^* = \frac{n_1}{D_1}, \quad \hat{\beta}_1^* = \frac{n_2}{\hat{\alpha}_1^* D_2}, \quad \hat{\beta}_2^* = \frac{n_3}{\hat{\alpha}_1^* \hat{\beta}_1^* D_3}, \dots, \hat{\beta}_{M-1}^* = \frac{n_M}{\hat{\alpha}_1^* \hat{\beta}_1^* \dots \hat{\beta}_{M-2}^* D_M}.$$

Note that if $0 < \hat{\beta}_i^* \leq 1$, for $i = 1, \dots, M-1$, then $\hat{\alpha}_1^*, \hat{\beta}_1^*, \dots, \hat{\beta}_{M-1}^*$ are the MLEs of $\alpha_1, \beta_1, \dots, \beta_{M-1}$, respectively. If some of the $\hat{\beta}_i^* > 1$, then the following algorithm can be used to obtain the MLEs of the unknown parameters.

Algorithm 1.

Step 1: For the given data set first obtain $\hat{\alpha}_1^, \hat{\beta}_1^*, \dots, \hat{\beta}_{M-1}^*$. If $0 < \hat{\beta}_i^* \leq 1$, for $i = 1, \dots, M-1$, then stop, and they are the MLEs.*

Step 2: If some of the $\hat{\beta}_i^ > 1$, then replace all of them by 1 in (11) and re-estimate the remaining parameters by maximizing with respect to the rest of the parameters.*

Step 3: Continue the process until all the $0 < \hat{\beta}_i^ \leq 1$, and in that case we will denote the MLEs as $\hat{\alpha}_1, \hat{\beta}_1, \dots, \hat{\beta}_{M-1}$.*

Theorem 1: Algorithm 1 stops in finite number of steps.

Proof: Since M is finite, Algorithm 1 stops in a finite number of steps. ■

Theorem 2: Algorithm 1 provides the MLEs of $\alpha_1, \beta_1, \dots, \beta_{M-1}$, under the constraints $0 < \beta_1 \leq \beta_2 \leq \dots \leq \beta_{M-1} \leq 1$.

Proof: See Appendix A.

4 BAYES ESTIMATES AND CREDIBLE INTERVALS

4.1 WITHOUT ORDERING

In this section we provide the Bayes estimates and the associated credible intervals when there is no ordering on the unknown parameters. The likelihood function of the observed data $\mathcal{D}$ can be written as

$$L(\mathcal{D}|\alpha_1, \dots, \alpha_M) = \alpha_1^{n_1} \dots \alpha_M^{n_M} e^{-\sum_{j=1}^M \alpha_j D_j}, \quad (12)$$

where D_j 's are same as defined in (9). Now based on the prior π_0 as defined in (4), the joint posterior of $\alpha_1, \dots, \alpha_M$ can be written as

$$\begin{aligned} \pi(\alpha_1, \dots, \alpha_M|\mathcal{D}) &\propto (\bar{\alpha})^{a-\bar{c}} \prod_{i=1}^M \{\alpha_i^{n_i+c_i-1} e^{-\alpha_i(b+D_i)}\} \\ &= K(\bar{\alpha})^{a-\bar{c}} \prod_{i=1}^M f_{\text{Gamma}}(\alpha_i; c_i + n_i, b + D_i), \end{aligned} \quad (13)$$

here $f_{\text{Gamma}}(\alpha_i; c_i + n_i, (b + D_i))$ denotes the PDF of $\text{Gamma}(c_i + n_i, b + D_i)$, for $i = 1, \dots, M$, and K is the normalizing constant. Hence, if we want to compute the Bayes estimate of any function of $\alpha_1, \dots, \alpha_M$, say $g(\alpha_1, \dots, \alpha_M)$, then with respect to the squared error loss function, it becomes

$$\hat{g}_{\text{Bayes}}(\alpha_1, \dots, \alpha_M) = K \int_0^\infty \dots \int_0^\infty g(\alpha_1, \dots, \alpha_M) (\bar{\alpha})^{a-\bar{c}} \prod_{i=1}^M f_{\text{Gamma}}(\alpha_i; c_i + n_i, b + D_i) d\alpha_i. \quad (14)$$

Now it is clear from the structure of (14) that $\hat{g}_{\text{Bayes}}(\alpha_1, \dots, \alpha_M)$ and the associated credible intervals can be computed quite conveniently based on importance sampling technique using

the following algorithm.

Algorithm 2.

Step 1: Generate independent samples $\{\alpha_{i1}, \dots, \alpha_{iB}\}$, from $\text{Gamma}(c_i + n_i, b + D_i)$, for $i = 1, \dots, M$

Step 2: Calculate $g_i = g(\alpha_{1i}, \dots, \alpha_{Mi})$ and

$$w_i = \frac{(\alpha_{1i} + \dots + \alpha_{Mi})^{a-\bar{c}}}{\sum_{j=1}^B (\alpha_{1j} + \dots + \alpha_{Mj})^{a-\bar{c}}}$$

Step 3: $\widehat{g}_{\text{Bayes}}(\alpha_1, \dots, \alpha_M)$ can be approximated as $\sum_{i=1}^B g_i w_i$.

Step 4: Rearrange $(g_1, w_1), (g_2, w_2), \dots, (g_B, w_B)$ as $(g_{(1)}, w_{(1)}), (g_{(2)}, w_{(2)}), \dots, (g_{(B)}, w_{(B)})$, where $g_{(1)} \leq g_{(2)} \leq \dots \leq g_{(B)}$. Note that $w_{(i)}$'s are not ordered, they are just associated with $g_{(i)}$'s.

A $100(1-\gamma)\%$ credible interval for $g(\alpha_1, \dots, \alpha_M)$ can be obtained as (g_{j_1}, g_{j_2}) , where j_1 and j_2 satisfy

$$j_1, j_2 \in \{1, 2, \dots, B\}, j_1 < j_2, \sum_{i=j_1}^{j_2} w_{(i)} \leq 1 - \gamma < \sum_{i=j_1}^{j_2+1} w_{(i)}. \quad (15)$$

The $100(1-\gamma)\%$ HPD credible interval (CRI) of $g(\alpha_1, \dots, \alpha_M)$ becomes $(g_{(j_1^)}, g_{(j_2^*)})$, where $1 \leq j_1^* < j_2^* \leq B$ satisfy*

$$\sum_{i=j_1^*}^{j_2^*} w_{(i)} \leq 1 - \gamma < \sum_{i=j_1^*}^{j_2^*+1} w_{(i)} \quad \text{and} \quad g_{(j_2^*)} - g_{(j_1^*)} \leq g_{(j_2)} - g_{(j_1)}$$

for all j_1 and j_2 satisfying (15).

4.2 ORDER RESTRICTED CASE

In this section it is assumed that $\infty > \alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_M > 0$. The parameters $\{\alpha_1, \dots, \alpha_M\}$ has the ODG prior as defined in (5). Based on the likelihood function (12), the posterior density function of $\alpha_1, \dots, \alpha_M$ for $\infty > \alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_M > 0$ can be written as

$$\begin{aligned} \pi_0(\alpha_1, \dots, \alpha_M | \mathcal{D}) &\propto e^{-(b+D)\bar{\alpha}} (\bar{\alpha})^{(a+MJ) - (\bar{c}+MJ)} \sum_{i_1, \dots, i_M \in \mathcal{P}} \left(\alpha_1^{c_{i_1}+J-1} \alpha_2^{c_{i_2}+J-1} \dots \alpha_M^{c_{i_M}+J-1} \right) \\ &\times \left\{ \prod_{i=1}^M \alpha_i^{n_i-J} \right\} e^{-\sum_{j=1}^M (D_j-D)}, \end{aligned} \quad (16)$$

where $J = \min\{n_1, \dots, n_M\}$ and $D = \min\{D_1, \dots, D_M\}$. Hence, (16) can be written as

$$\pi_0(\alpha_1, \dots, \alpha_M | \mathcal{D}) = K \pi^*(\alpha_1, \dots, \alpha_M | \mathcal{D}) h(\alpha_1, \dots, \alpha_M | \mathcal{D}), \quad (17)$$

where $\pi^*(\alpha_1, \dots, \alpha_M | \mathcal{D})$ is the PDF of ODG $(a + MJ, b + D, c_1 + J, \dots, c_M + J)$,

$$h(\alpha_1, \dots, \alpha_M | \mathcal{D}) = \left\{ \prod_{i=1}^M \alpha_i^{n_i-J} \right\} e^{-\sum_{j=1}^M (D_j-D)},$$

and K is the normalizing constant. The generation of $(\alpha_1, \dots, \alpha_M)$ from $\pi^*(\alpha_1, \dots, \alpha_M | \mathcal{D})$ is quite straight forward as follows, see [Mondal and Kundu \(2019\)](#). Generate $\theta \sim \text{Gamma}(a + MJ, b + D)$ and $(\theta_1, \dots, \theta_M) \sim \text{D}(c_1 + J, \dots, c_M + J)$ independently. Compute $\theta\theta_1, \dots, \theta\theta_M$, order them from largest to smallest and let us denote them as $(\alpha_1, \dots, \alpha_M)$. Then $(\alpha_1, \dots, \alpha_M)$ is the desired sample. Hence, if we want to compute the Bayes estimate of any function of $\alpha_1, \dots, \alpha_M$, say $g(\alpha_1, \dots, \alpha_M)$, with respect to the squared error loss function and the associated credible interval as before, we can use Algorithm 3 similar to Algorithm 2.

Algorithm 3.

Step 1: Generate independent samples $\{\alpha_{1j}, \dots, \alpha_{Mj}\}$, from ODG $(a + MJ, b + D, c_1 + J, \dots, c_M + J)$ for $j = 1, \dots, B$. In this case $\alpha_{1j} \geq \dots \geq \alpha_{Mj}$, for $j = 1, \dots, B$.

Step 2: Calculate $g_i = g(\alpha_{1i}, \dots, \alpha_{Mi})$ and

$$w_i = \frac{h(\alpha_{1i}, \dots, \alpha_{Mi})}{\sum_{j=1}^B h(\alpha_{1j}, \dots, \alpha_{Mj})}.$$

Step 3 and Step 4: Same as in Algorithm 2.

5 SIMULATION RESULTS

In this section we present the results of numerical experiments that were performed to assess how the suggested approaches behaved with varying sample sizes, sampling schemes, parameter values, and priors. All of the computations were performed using *R* software. We have considered $M = 2, 3$, with different sample sizes (n), different number of failures (m), different hyperparameters, and different sampling schemes (i.e., different R_i values). To generate progressively censored samples, we used the algorithm given in [Balakrishnan and Cramer \(2014\)](#) with the proposed semi-parametric model defined in Section 2. We report the average estimates (AEs) and mean squared errors (MSEs) based on 1000 replications. We consider the following sampling schemes:

Scheme 1: $n = 20, m = 15, R_1 = \dots = R_{14} = 0, R_{15} = 5$, i.e., $(20, 15, 14 * 0, 5)$

Scheme 2: $n = 20, m = 15, R_1 = 5, R_2 = \dots = R_{15} = 0$, i.e., $(20, 15, 5, 14 * 0)$

Scheme 3: $n = 30, m = 15, R_1 = \dots = R_{14} = 0, R_{15} = 15$, i.e., $(30, 15, 14 * 0, 15)$

Scheme 4: $n = 30, m = 15, R_1 = 15, R_2 = \dots = R_{15} = 0$, i.e., $(30, 15, 15, 14 * 0)$

Scheme 5: $n = 30, m = 15, R_1 = \dots = R_{15} = 1$, i.e., $(30, 15, 1 * 15)$

Scheme 6: $n = 30, m = 20, R_1 = \dots = R_{19} = 0, R_{20} = 10$, i.e., $(30, 20, 19 * 0, 10)$

Scheme 7: $n = 30, m = 20, R_1 = 10, R_2 = \dots = R_{20} = 0$, i.e., $(30, 20, 10, 19 * 0)$

Take note that Schemes 1, 3, and 6 are the standard Type-II censoring schemes for fixed n and m , meaning that at the time of the m -th failure occurs, $n - m$ items are eliminated. Schemes 2, 4, and 7 are essentially the opposite of Type-II censoring schemes in that, for

specified n and m , $n - m$ items are eliminated at the first failure. It is commonly known that the estimated experimental time of the Type-II censoring scheme (i.e., for Schemes 1, 3, and 6) is less than that of Schemes 2, 4, and 7 for fixed n and m . Any other censoring scheme (for fixed m and n) actually has an expected time that falls between these two extremes; for example, the expected experimental time of Scheme 5 lies between Schemes 3 and 4.

5.1 Case when $M = 2$

For $M = 2$, we set parameter values $\theta_1 = 2/3$, and $\theta_2 = 1/3$ with varying $\tau_1 = 0.3, 0.5$. For Bayes estimation without ordering the hyperparameters the settings are: $a = 1.5, b = 1, c_1 = 1.5, c_2 = 0.2$. For Bayes estimation with ordering the hyperparameters the settings are: $a = 0.8, b = 0.5, c_1 = 2.0, c_2 = 1.8$. The results are reported in Tables 1 to 2 for $\tau_1 = 0.3$ and in Tables 3 to 4 for $\tau_1 = 0.5$.

Table 1: AEs and MSEs (in parenthesis) of the estimates of θ_1 , and θ_2 with different values of n and m ($\theta_1 = 2/3 = 0.6667$, and $\theta_2 = 1/3 = 0.3333$, $\tau_1 = 0.3$).

Sampling scheme	Without Ordering			
	MLE		Bayes	
	θ_1	θ_2	θ_1	θ_2
(20, 15, 14 * 0, 5)	0.6907 (0.0177)	0.3093 (0.0177)	0.7282 (0.0122)	0.2718 (0.0122)
(20, 15, 5, 14 * 0)	0.6232 (0.0244)	0.3768 (0.0244)	0.6836 (0.0108)	0.3164 (0.0108)
(30, 15, 14 * 0, 15)	0.7763 (0.0233)	0.2237 (0.0233)	0.7878 (0.0212)	0.2122 (0.0212)
(30, 15, 15, 14 * 0)	0.6175 (0.0232)	0.3825 (0.0232)	0.6764 (0.0102)	0.3236 (0.0102)
(30, 15, 1 * 15)	0.7375 (0.0162)	0.2625 (0.0162)	0.7575 (0.0149)	0.2425 (0.0149)
(30, 20, 19 * 0, 10)	0.7241 (0.0173)	0.2759 (0.0173)	0.7458 (0.0142)	0.2542 (0.0142)
(30, 20, 10, 19 * 0)	0.6362 (0.0188)	0.3638 (0.0188)	0.6819 (0.0095)	0.3181 (0.0095)

Table 2: AEs and MSEs (in parenthesis) of the estimates of θ_1 , and θ_2 with different values of n and m ($\theta_1 = 2/3 = 0.6667$, and $\theta_2 = 1/3 = 0.3333$, $\tau_1 = 0.3$).

Sampling scheme	With Ordering			
	MLE		Bayes	
	θ_1	θ_2	θ_1	θ_2
(20, 15, 14 * 0, 5)	0.6991 (0.0139)	0.3009 (0.0139)	0.5653 (0.0115)	0.4347 (0.0115)
(20, 15, 5, 14 * 0)	0.6461 (0.0133)	0.3539 (0.0133)	0.5559 (0.0128)	0.4441 (0.0128)
(30, 15, 14 * 0, 15)	0.7780 (0.0225)	0.2220 (0.0225)	0.6035 (0.0088)	0.3965 (0.0088)
(30, 15, 15, 14 * 0)	0.6384 (0.0130)	0.3616 (0.0130)	0.5559 (0.0129)	0.4441 (0.0129)
(30, 15, 1 * 15)	0.7405 (0.0148)	0.2595 (0.0148)	0.5916 (0.0089)	0.4084 (0.0089)
(30, 20, 19 * 0, 10)	0.7292 (0.0148)	0.2708 (0.0148)	0.5616 (0.0126)	0.4384 (0.0126)
(30, 20, 10, 19 * 0)	0.6515 (0.0112)	0.3485 (0.0112)	0.5475 (0.0147)	0.4525 (0.0147)

Table 3: AEs and MSEs (in parenthesis) of the estimates of θ_1 , and θ_2 with different values of n and m ($\theta_1 = 2/3 = 0.6667$, and $\theta_2 = 1/3 = 0.3333$, $\tau_1 = 0.5$).

Sampling scheme	Without Ordering			
	MLE		Bayes	
	θ_1	θ_2	θ_1	θ_2
(20, 15, 14 * 0, 5)	0.7003 (0.0160)	0.2997 (0.0160)	0.7260 (0.0128)	0.2740 (0.0128)
(20, 15, 5, 14 * 0)	0.6341 (0.0188)	0.3659 (0.0188)	0.6756 (0.0109)	0.3244 (0.0109)
(30, 15, 14 * 0, 15)	0.7869 (0.0239)	0.2131 (0.0239)	0.7954 (0.0233)	0.2046 (0.0233)
(30, 15, 15, 14 * 0)	0.6272 (0.0186)	0.3728 (0.0186)	0.6694 (0.0103)	0.3306 (0.0103)
(30, 15, 1 * 15)	0.7256 (0.0130)	0.2744 (0.0130)	0.7426 (0.0128)	0.2574 (0.0128)
(30, 20, 19 * 0, 10)	0.7397 (0.0149)	0.2603 (0.0149)	0.7526 (0.0143)	0.2474 (0.0143)
(30, 20, 10, 19 * 0)	0.6384 (0.0134)	0.3616 (0.0134)	0.6704 (0.0085)	0.3296 (0.0085)

Table 4: AEs and MSEs (in parenthesis) of the estimates of θ_1 , and θ_2 with different values of n and m ($\theta_1 = 2/3 = 0.6667$, and $\theta_2 = 1/3 = 0.3333$, $\tau_1 = 0.5$).

Sampling scheme	With Ordering			
	MLE		Bayes	
	θ_1	θ_2	θ_1	θ_2
(20, 15, 14 * 0, 5)	0.7067 (0.0129)	0.2933 (0.0129)	0.6095 (0.0081)	0.3905 (0.0081)
(20, 15, 5, 14 * 0)	0.6484 (0.0116)	0.3516 (0.0116)	0.5755 (0.0101)	0.4245 (0.0101)
(30, 15, 14 * 0, 15)	0.7879 (0.0234)	0.2121 (0.0234)	0.7125 (0.0109)	0.2875 (0.0109)
(30, 15, 15, 14 * 0)	0.6423 (0.0113)	0.3577 (0.0113)	0.5800 (0.0097)	0.4200 (0.0097)
(30, 15, 1 * 15)	0.7266 (0.0125)	0.2734 (0.0125)	0.6752 (0.0073)	0.3248 (0.0073)
(30, 20, 19 * 0, 10)	0.7416 (0.0140)	0.2584 (0.0140)	0.6286 (0.0091)	0.3714 (0.0091)
(30, 20, 10, 19 * 0)	0.6462 (0.0098)	0.3538 (0.0098)	0.5625 (0.0122)	0.4375 (0.0122)

Some conclusions are quite clear from these table values. Both with and without ordering cases it is observed that the performance of the Bayes estimators are better than MLEs in terms of MSEs. It is as expected due to the involvement of priors information in case of Bayes estimators. For fixed n as m increases the MSEs of all the estimators decrease, as expected. It can also be seen that τ_1 and τ_2 also affect the MSEs of the MLE and the Bayes estimators. Additionally, it is evident that τ_1 and τ_2 have an impact on the Bayes estimators and the MLEs in terms of the MSEs. Increasing sample size n may not reduce MSEs for fixed R_i values (within given restrictions) and m , if we compare schemes 1 and 3 or 2 and 4. Next, a comparison of schemes 3 and 6 or 4 and 7 shows that, under all conditions, a larger effective sample size m results in lower MSEs for both estimators. This implies that the number of failures m matters more than the actual sample size n when determining the efficiency of the parameters.

5.2 Case when $M = 3$

For $M = 3$, we set parameter values $\theta_1 = 5/12$, $\theta_2 = 1/3$, and $\theta_3 = 1/4$ with various combination of $(\tau_1, \tau_2) = (0.5, 1.0), (0.5, 1.5), (1.0, 1.5)$. For Bayes estimation without ordering the hyperparameters settings are: $a = 1.5, b = 1, c_1 = 1.5, c_2 = 0.2, c_3 = 1.0$. For Bayes estimation with ordering the hyperparameters settings are: $a = 0.8, b = 0.5, c_1 = 2.0, c_2 = 1.8, c_3 = 1.7$. The results are reported in Tables 7 to 8 for $(\tau_1, \tau_2) = (0.5, 1.0)$, in Tables 9 to 10 for $(\tau_1, \tau_2) = (0.5, 1.5)$, and in Tables 11 to 12 for $(\tau_1, \tau_2) = (1.0, 1.5)$. [The Tables 9 to 12 have been presented in Appendix B.](#)

It can also be seen that τ_1 and τ_2 also affect the MSEs of the MLE and the Bayes estimators. Comparing Schemes 1 and 3 or Schemes 2 and 4 reveals that for fixed R_i values (within stated limitations) and m , increasing sample size n may not reduce MSEs. Next, comparing Schemes 3 and 6 or Schemes 4 and 7 reveals that increasing the effective sample size m leads to lower MSEs for both estimators in all circumstances. This suggests that when calculating the efficiency of the parameters, the number of failures m matters more than the sample size n .

If we compare the Bayes estimators for the Schemes 1 and 2, Schemes 3 and 4, and Schemes 6 and 7, it is clear that overall the better estimates are provided by schemes 2, 4, and 7. This is mainly due to two reasons: (i) proper choice of hyperparameter of the prior distribution and (ii) the estimated duration of experiments is longer for censoring schemes 2, 4, and 7 compared to schemes 1, 3, and 6. Thus, data obtained from censoring schemes 2, 4, and 7 should reveal more information about the unknown parameters than data obtained from censoring schemes 1, 3, and 6.

6 CHOICE OF CUT POINTS & DATA ANALYSIS

6.1 CHOICE OF CUT POINTS

So far it has been assumed that the number of cut points and also the position of the cut points are known. In the literature so far no one had discussed this issue. In this section we propose a method which can be implemented quite easily in practice. We are mainly interested about the shape of the hazard function in $[0, \tau]$, where $\tau = t_m$. It is assumed that the position of the cut points are modeled as the arrival times of the events in $[0, \tau]$, from a non-homogeneous Poisson process $\{N(t); t \geq 0\}$ with intensity function $u(t) = \alpha t^{\alpha-1}$, for $\alpha > 0$. But the arrival times of the events from $N(t)$ are not observed, hence they are being treated as missing values. For the given number of cut points $M - 1$ in $[0, \tau]$, the likelihood function of the observation is treated as a function of the unknown parameters $\theta_1, \dots, \theta_M, \alpha$, when the arrival times of the events from $N(t)$, which are missing, are replaced by their corresponding expected values. For a given M we compute the maximum likelihood estimators of the unknown parameters, and then based on some information theoretic criteria, like Akaike Information Criterion (AIC) or Bayes Information Criterion (BIC), we estimate M also.

Now we provide the joint PDF of the arrival times of a non-homogeneous Poisson process. Suppose $\{N(t); t \geq 0\}$ is a non-homogeneous Poisson process with intensity function $u(t) = \alpha \lambda t^{\alpha-1}$, for $\alpha > 0$. Let $\tilde{T}_1, \tilde{T}_2, \dots$ denote the arrival times of the first, second etc. arrival times of the events from $N(t)$. Then the joint PDF of $\tilde{T}_1, \dots, \tilde{T}_n$, given that $N(\tau) = n$, can be written as

$$f_{\tilde{T}_1, \dots, \tilde{T}_n | N(\tau)=n}(t_1, \dots, t_n) = \frac{n! \alpha^n}{\tau^{n\alpha}} t_1^{\alpha-1} \dots t_n^{\alpha-1}; \quad 0 < t_1 < t_2 < \dots < t_n < \tau, \quad (18)$$

and zero, otherwise. Hence, for $j = 1, \dots, n$,

$$\begin{aligned}
E(\tilde{T}_j | N(\tau) = n) &= \frac{n! \alpha^n}{\tau^{n\alpha}} \int_0^{t_2} \int_0^{t_3} \dots \int_0^{t_n} \int_0^\tau t_1^{\alpha-1} \dots t_{j-1}^{\alpha-1} t_j^\alpha t_{j+1}^{\alpha-1} \dots t_n^{\alpha-1} dt_1 dt_2 \dots dt_{n-1} dt_n \\
&= \frac{n! \alpha^n}{\tau^{n\alpha}} \times \frac{\tau^{n\alpha+1}}{\alpha(2\alpha) \dots ((j-1)\alpha)(j\alpha+1)((j+1)\alpha+1) \dots (n\alpha+1)} \\
&= \frac{n!}{(j-1)!} \times \frac{\alpha^{n-j+1} \tau}{(j\alpha+1)((j+1)\alpha+1) \dots (n\alpha+1)} \\
&= \frac{1}{(1 + \frac{1}{1+j\alpha}) \dots (1 + \frac{1}{1+n\alpha})}
\end{aligned}$$

Therefore, for homogeneous Poisson process i.e. when $\alpha = 1$,

$$E(\tilde{T}_j | N(\tau) = n) = \frac{j\tau}{n+1}.$$

In this case the cut points are equidistant.

6.2 DATA ANALYSIS

Data Set 1: Lawless electrical appliances dataset

For illustrative purposes, we have performed the analyses of two real-life data sets. The first consists of failure times of 36 electrical appliances subjected to a life test is taken from [Lawless \(2011\)](#). The failure times, in cycles, are given below:

11, 35, 49, 170, 329, 381, 708, 958, 1062, 1167, 1594, 1925, 1990, 2223, 2327, 2400, 2451, 2471, 2551, 2565, 2568, 2694, 2702, 2761, 2831, 3034, 3059, 3112, 3214, 3478, 3504, 4329, 6367, 6976, 7846, 13403.

We begin the analysis of data set by checking the shape of its hazard rate. This can be done by using the total time on test transform (TTT-Transform) plot, as a non-parametric test. From Fig. 1, we can see that the hazard rate shape of this data set is bathtub shaped. A progressively censored sample of the scheme $n = 36$, $m = 10$, $R_1 = \dots = R_9 = 2$, and $R_{10} = 8$ was generated and analyzed by [Kundu and Joarder \(2006\)](#) and [Kundu \(2008\)](#)

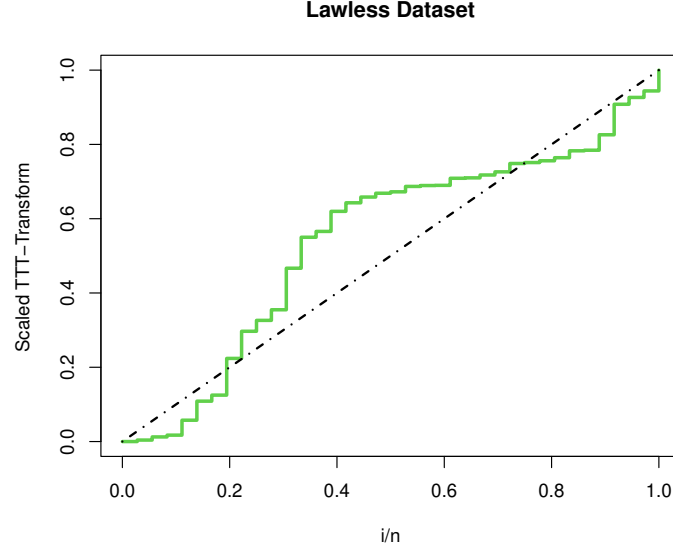


Figure 1: The TTT-Transform plots for the Lawless electrical appliances dataset.

using an exponential and Weibull model. The following progressively censored data were observed: 11, 35, 49, 170, 329, 958, 1925, 2223, 2400, and 2568. Here we reanalyze the data using the semiparametric model. For computational ease, we divide all of the values by 100.

It has been observed by several authors, see for example [Balakrishnan et al. \(2016\)](#) or [Ganguly et al. \(2024\)](#), that for all practical purposes $M = 2, 3$ or 4 works quite well. Hence, here also we have considered $M = 2, 3$ and 4 only. We have obtained the cut points, the

Table 5: Choice of cut points for the Lawless electrical appliances dataset.

	α	Cut points	θ_1	θ_2	θ_3	θ_4	Log L	AIC	BIC
M = 2	9.1651	22.2300	0.3316	0.6684			-39.8549	83.7098	86.8768
M = 3	9.1651	22.2300 24.4169	0.1035	0.6742	0.2223		-37.4438	80.8877	85.6382
M = 4	5.1250	19.1068 22.2263 24.2020	0.0976	0.2044	0.5476	0.1504	-35.5825	79.1651	85.4991
Parametric model									
Weibull			$\hat{\lambda} = 0.0627$	$\hat{\nu} = 0.6297$			-46.9359	97.8719	101.0389
log-normal			$\hat{\mu} = 3.4722$	$\hat{\tau} = 3.6935$			-55.5372	115.0745	118.2415

estimates of the unknown parameters, the log-likelihood values and the associated AIC and BIC values. The results are reported in Table 5. It is clear from the table that $M = 3$ with

cut points $\tau_1 = 22.2300$ and $\tau_2 = 24.4169$ is better choice than $M = 2$, because both the AIC and BIC values are less in this case. The results for the $M = 4$ are provided in the Table 5. It is clear that AIC and BIC values are smaller for $M = 4$ compared to $M = 2$ and 3. The associated 95% confidence intervals of θ_1 , θ_2 , θ_3 and θ_4 based on parametric bootstrap are (0.0194,0.1242), (0.0977,0.4619), (0.3373,0.8102) and (0.0129,0.4034), respectively.

We have also considered two parametric models namely Weibull and log-normal mainly for comparison purposes. The corresponding PDFs are as follows;

$$\begin{aligned} f_{WE}(x|\lambda, \nu) &= \lambda \nu x^{\nu-1} \exp -\lambda x^\nu, \quad x > 0, \quad \lambda > 0, \quad \nu > 0, \\ f_{LN}(x|\mu, \tau) &= \frac{1}{x\sqrt{\tau}} \phi\left(\frac{\ln x - \mu}{\sqrt{\tau}}\right); \quad x > 0, -\infty < \mu < \infty, \tau > 0. \end{aligned}$$

We have reported the MLEs of the unknown parameters, the corresponding log-likelihood, AIC and BIC values in Table 5. It is clear that for the Lawless dataset proposed semiparametric model performs better than the Weibull model.

Data Set 2: Meeker and Escobar U.S. Grampus dataset

This data set consists of 57 times (in thousands of operating hours) of unscheduled maintenance actions for the number 4 diesel engine of the U.S. Grampus, up to 16 thousand hours of operation [Meeker et al. \(2022\)](#). The observation are:

0.860, 0.258, 1.317, 1.442, 1.897, 2.011, 2.122, 2.439, 3.203, 3.298, 3.902, 3.910, 4.000, 4.247, 4.411, 4.456, 4.517, 4.899, 4.910, 5.676, 5.755, 6.137, 6.221, 6.311, 6.613, 6.975, 7.335, 8.158, 8.498, 8.690, 9.042, 9.330, 9.394, 9.426, 9.872, 10.191, 11.511, 11.575, 12.100, 12.126, 12.368, 12.681, 12.795, 13.399, 13.668, 13.780, 13.877, 14.007, 14.028, 14.035, 14.173, 14.173, 14.449, 14.587, 14.610, 15.070, 16.000.

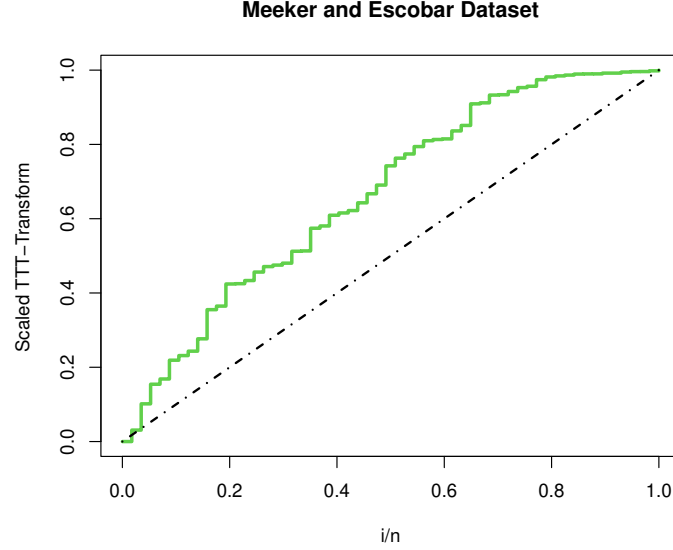


Figure 2: The TTT-Transform plots for the Meeker and Escobar U.S. Grampus dataset.

Based on the TTT- Transform hazard plot for this data set, shown in Fig. 2, we can see that it has an increasing failure rate. A progressively censored sample of the scheme $n = 57, m = 26, R_1 = 11, R_2 = 3, R_3 = 5, R_4 = 3, R_5 = 5, R_6 = 1, R_7 = 0, R_8 = 0, R_9 = 0, R_{10} = 1, R_{11} = \dots = R_{26} = 0$ was generated and analyzed by [Sarhan et al. \(2023\)](#). The following progressively censored data were observed: 0.258, 4.000, 4.517, 6.221, 7.335, 9.394, 9.872, 10.191, 11.511, 11.575, 12.126, 12.368, 12.681, 12.795, 13.399, 13.668, 13.78, 13.877, 14.007, 14.028, 14.035, 14.173, 14.173, 14.449, 14.587, 14.61.

Here also we analyze the data using the semi-parametric model with $M = 2, 3$ and 4 . The results are reported in Table 6. It is cleared that log-likelihood and AIC are smaller for $M = 4$ compare to the $M = 2$ and 3 . For $M = 4$, BIC is slightly higher in comparison of $M = 3$. But if we look at the survival function at any point, say $t = 13$, $S(t) = 0.5555$ for $M = 4$ and $S(t) = 0.5206$ for $M = 3$, which are close to each other. Comparing all these we prefer $M = 3$, because it involves smaller number of parameters. The associated 95% confidence intervals of θ_1, θ_2 and θ_3 based on parametric bootstrap are (0.0056,0.0294),

(0.0346,0.2609) and (0.7190,0.9487), respectively.

We have also considered a parametric model for the comparison by assuming the Weibull distribution with the PDF defined above. We analyzed the progressively censored Mekker dataset given above and the loglikelihood, AIC and BIC is obtained which is given in Table 6. It is cleared that proposed semiparametric model also performs better than the parametric models for this dataset.

Table 6: Choice of cut points for the Meeker and Escobar U.S. Grampus dataset.

	α	Cut points	θ_1	θ_2	θ_3	θ_4	LogL	AIC	BIC
M = 2	21.2600	13.6678	0.0168	0.9832			-66.7546	137.5092	141.5953
M = 3	5.0285	11.4919 13.3983	0.0127	0.1272	0.8600		-60.6542	127.3083	133.4375
M = 4	6.6063	11.5107 13.0240 13.9404	0.0058	0.0702	0.1045	0.8195	-58.7997	125.5994	133.7716
Parametric model									
Weibull			$\hat{\lambda} = 0.0001$	$\hat{\nu} = 3.5402$			-77.5526	159.1051	163.1912
log-normal			$\hat{\mu} = 2.3796$	$\hat{\tau} = 0.3069$			-83.6169	171.2340	175.3201

7 CONCLUSIONS

In this paper we have considered both the classical and Bayesian inference of the unknown parameters of a proportional hazard model based on Type-II progressively censored data. We have considered a flexible proportional hazard class where the baseline distribution has a piecewise-constant hazard function. We have provided both with and without order-restricted inference of the unknown parameters. Extensive simulations have been performed to show the effectiveness of the proposed methods and it is observed that the proposed methods work quite well in practice. We have analyzed two real data sets for illustrative purposes, and the results are quite satisfactory.

ACKNOWLEDGEMENTS

The authors would like to thank the unknown reviewers for their constructive comments which have helped to improve the paper significantly.

Disclosure of interest

We would like to mention that we do not have any conflict of interests in publishing this paper.

Data Availability Statement

Data is available within the article.

Funding

No funding was received.

APPENDIX A

Proof of Theorem 2:

We provide the proof for $M = 4$, for brevity. The same argument holds for the general case. For $M = 4$, the log-likelihood function becomes

$$\begin{aligned} l(\alpha_1, \beta_1, \beta_2, \beta_3 | \mathcal{D}) &= m \ln \alpha_1 + \bar{n}_1 \ln \beta_1 + \bar{n}_2 \ln \beta_2 + \bar{n}_3 \ln \beta_3 \\ &\quad - \alpha_1 [D_1 + D_2 \beta_1 + D_3 \beta_1 \beta_2 + D_4 \beta_1 \beta_2 \beta_3]. \end{aligned} \quad (19)$$

The log-likelihood function in (19) has a unique maximum at $\hat{\alpha}_1^*, \hat{\beta}_1^*, \hat{\beta}_2^*, \hat{\beta}_3^*$, where

$$\hat{\alpha}_1^* = \frac{n_1}{D_1}, \quad \hat{\beta}_1^* = \frac{n_2 D_1}{n_1 D_2}, \quad \hat{\beta}_2^* = \frac{n_3 D_2}{n_2 D_3}, \quad \hat{\beta}_3^* = \frac{n_4 D_3}{n_3 D_4}. \quad (20)$$

It is clear that the function (19) does not have any other local maximum. Observe that for a given $(\beta_1, \beta_2, \beta_3)$, the log-likelihood function in (19) attains its maximum when

$$\hat{\alpha}_1(\beta_1, \beta_2, \beta_3) = \frac{m}{D_1 + D_2 \beta_1 + D_3 \beta_1 \beta_2 + D_4 \beta_1 \beta_2 \beta_3}. \quad (21)$$

Substituting $\hat{\alpha}_1(\beta_1, \beta_2, \beta_3)$ in (19), and ignoring the additive constant, the profile log-likelihood function of $\beta_1, \beta_2, \beta_3$ can be written as

$$p(\beta_1, \beta_2, \beta_3) = -m \ln(D_1 + D_2\beta_1 + D_3\beta_1\beta_2 + D_4\beta_1\beta_2\beta_3) + \bar{n}_1 \ln \beta_1 + \bar{n}_2 \ln \beta_2 + \bar{n}_3 \ln \beta_3. \quad (22)$$

Hence

$$\sup_{\substack{\alpha_1 > 0 \\ 0 \leq \beta_1, \beta_2, \beta_3 \leq 1}} l(\alpha_1, \beta_1, \beta_2, \beta_3) = \sup_{0 \leq \beta_1, \beta_2, \beta_3 \leq 1} p(\beta_1, \beta_2, \beta_3)$$

From (22), it can be observed that the function $p(\beta_1, \beta_2, \beta_3)$ has a unique maximum at $(\hat{\beta}_1^*, \hat{\beta}_2^*, \hat{\beta}_3^*)$, where $\hat{\beta}_1^*$, $\hat{\beta}_2^*$ and $\hat{\beta}_3^*$ are same as defined in (20), and the function (22) does not have any other local maximum.

Now we claim that if $\hat{\beta}_1^* > 1$, then

$$\sup_{\substack{0 \leq \beta_1 \leq 1 \\ \beta_2 \geq 0, \beta_3 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_2 \geq 0, \beta_3 \geq 0} p(1, \beta_2, \beta_3) \quad (23)$$

Suppose (23) is not true, then there exists $0 < \tilde{\beta}_1 < 1$, $0 < \tilde{\beta}_2 < \infty$ and $0 < \tilde{\beta}_3 < \infty$, such that

$$\sup_{\substack{0 \leq \beta_1 \leq 1 \\ \beta_2 \geq 0, \beta_3 \geq 0}} p(\beta_1, \beta_2, \beta_3) = p(\tilde{\beta}_1, \tilde{\beta}_2, \tilde{\beta}_3).$$

It implies that $(\hat{\beta}_1^*, \hat{\beta}_2^*, \hat{\beta}_3^*) \neq (\tilde{\beta}_1, \tilde{\beta}_2, \tilde{\beta}_3)$, is a local maximum of $p(\beta_1, \beta_2, \beta_3)$, as $p(\beta_1, \beta_2, \beta_3) \rightarrow -\infty$ as $\beta_2 \rightarrow \infty$ or $\beta_3 \rightarrow \infty$. This is clearly a contradiction.

Along the same line it follows that if $\hat{\beta}_2^* > 1$, then

$$\sup_{\substack{0 \leq \beta_2 \leq 1 \\ \beta_1 \geq 0, \beta_3 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_1 \geq 0, \beta_3 \geq 0} p(\beta_1, 1, \beta_3) \quad (24)$$

and if $\hat{\beta}_3^* > 1$, then

$$\sup_{\substack{0 \leq \beta_3 \leq 1 \\ \beta_1 \geq 0, \beta_2 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_1 \geq 0, \beta_2 \geq 0} p(\beta_1, \beta_2, 1). \quad (25)$$

Combining (23) and (24), we obtain if $\hat{\beta}_1^* > 1$ and $\hat{\beta}_2^* > 1$, then

$$\sup_{\substack{0 \leq \beta_1, \beta_2 \leq 1 \\ \beta_3 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_3 \geq 0} p(1, 1, \beta_3). \quad (26)$$

Similarly, combining (23) and (25), if $\widehat{\beta}_1^* > 1$ and $\widehat{\beta}_3^* > 1$, then

$$\sup_{\substack{0 \leq \beta_1, \beta_3 \leq 1 \\ \beta_2 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_2 \geq 0} p(1, \beta_2, 1), \quad (27)$$

combining (24) and (25), if $\widehat{\beta}_2^* > 1$ and $\widehat{\beta}_3^* > 1$, then

$$\sup_{\substack{0 \leq \beta_2, \beta_3 \leq 1 \\ \beta_1 \geq 0}} p(\beta_1, \beta_2, \beta_3) = \sup_{\beta_1 \geq 0} p(\beta_1, 1, 1), \quad (28)$$

and combining (23), (24) and (25), if $\widehat{\beta}_1^* > 1$, $\widehat{\beta}_2^* > 1$ and $\widehat{\beta}_3^* > 1$, then

$$\sup_{0 \leq \beta_1, \beta_2, \beta_3 \leq 1} p(\beta_1, \beta_2, \beta_3) = p(1, 1, 1). \quad (29)$$

Further, observe that

$$\sup_{0 \leq \beta_1, \beta_2, \beta_3 \leq 1} p(\beta_1, \beta_2, \beta_3) = \sup_{0 \leq \beta_i \leq 1} \sup_{0 \leq \beta_j \leq 1} \sup_{0 \leq \beta_k \leq 1} p(\beta_1, \beta_2, \beta_3), \quad (30)$$

for all $i \neq j \neq k$, and $1 \leq i, j, k \leq 3$.

Now let us consider different cases:

Case 1: $\widehat{\beta}_1^* > 1$, $\widehat{\beta}_2^* > 1$, $\widehat{\beta}_3^* > 1$. The MLEs of β_1 , β_2 and β_3 are 1, 1, and 1, respectively.

Case 2: $\widehat{\beta}_1^* > 1$, $\widehat{\beta}_2^* > 1$, $\widehat{\beta}_3^* \leq 1$. The MLEs of β_1 and β_2 are 1 and 1, respectively. The MLE of β_3 can be obtained as the $\arg\max_{0 \leq \beta_3 \leq 1} p(1, 1, \beta_3)$. Since $p(1, 1, \beta_3)$ is unimodal function, it has a unique maximum. Further, $\sup_{0 \leq \beta_1, \beta_2, \beta_3 \leq 1} p(\beta_1, \beta_2, \beta_3) = \sup_{0 \leq \beta_3 \leq 1} p(1, 1, \beta_3)$, due to (30).

Case 3: $\widehat{\beta}_1^* > 1$, $\widehat{\beta}_2^* \leq 1$, $\widehat{\beta}_3^* \leq 1$. The MLEs of β_1 is 1, the MLEs β_2 and β_3 can be obtained as the $\arg\max_{0 \leq \beta_2, \beta_3 \leq 1} p(1, \beta_2, \beta_3)$. The function $p(1, \beta_2, \beta_3)$ has a unique maximum and we repeat the argument as before.

The other cases can be considered along the same line, and that proves Theorem 2. ■

APPENDIX B

Table 7: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 0.5$, $\tau_2 = 1.0$).

Sampling scheme	Without Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.4105 (0.0204)	0.3803 (0.0224)	0.2092 (0.0098)	0.4336 (0.0174)	0.3221 (0.0156)	0.2443 (0.0086)
(20, 15, 5, 14 * 0)	0.3863 (0.0205)	0.3488 (0.0183)	0.2649 (0.0100)	0.4113 (0.0161)	0.2888 (0.0153)	0.2999 (0.0124)
(30, 15, 14 * 0, 15)	0.4053 (0.0190)	0.4490 (0.0349)	0.1457 (0.0173)	0.4350 (0.0168)	0.3851 (0.0201)	0.1799 (0.0108)
(30, 15, 15, 14 * 0)	0.3793 (0.0224)	0.3560 (0.0213)	0.2647 (0.0100)	0.4045 (0.0174)	0.2951 (0.0169)	0.3003 (0.0124)
(30, 15, 1 * 15)	0.4240 (0.0164)	0.3854 (0.0202)	0.1906 (0.0102)	0.4456 (0.0153)	0.3302 (0.0141)	0.2241 (0.0075)
(30, 20, 19 * 0, 10)	0.4183 (0.0183)	0.3943 (0.0228)	0.1873 (0.0094)	0.4370 (0.0164)	0.3487 (0.0160)	0.2144 (0.0071)
(30, 20, 10, 19 * 0)	0.3943 (0.0201)	0.3382 (0.0189)	0.2675 (0.0089)	0.4133 (0.0166)	0.2919 (0.0164)	0.2948 (0.0107)

These table values make certain conclusions very evident. The parameter values are $\theta_1 = 5/12$, $\theta_2 = 1/3$, and $\theta_3 = 1/4$, which follows $\theta_1 > \theta_2 > \theta_3$, as expected, from the tables that with ordered MLEs efficiency is better then the without ordered MLEs efficiency. Also, both without and with ordering, Bayes estimators outperform MLE estimators in terms of MSEs as before.

Table 8: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 0.5$, $\tau_2 = 1.0$).

Sampling scheme	With Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.4229 (0.0104)	0.3587 (0.0049)	0.2184 (0.0094)	0.4481 (0.0035)	0.3094 (0.0011)	0.2425 (0.0013)
(20, 15, 5, 14 * 0)	0.3803 (0.0115)	0.3501 (0.0029)	0.2697 (0.0078)	0.4310 (0.0018)	0.3129 (0.0008)	0.2561 (0.0007)
(30, 15, 14 * 0, 15)	0.4610 (0.0090)	0.3879 (0.0081)	0.1512 (0.0159)	0.5125 (0.0149)	0.2992 (0.0025)	0.1884 (0.0070)
(30, 15, 15, 14 * 0)	0.3745 (0.0123)	0.3515 (0.0032)	0.2740 (0.0083)	0.4348 (0.0021)	0.3127 (0.0009)	0.2525 (0.0008)
(30, 15, 1 * 15)	0.4408 (0.0100)	0.3620 (0.0056)	0.1971 (0.0100)	0.4841 (0.0084)	0.2981 (0.0020)	0.2178 (0.0028)
(30, 20, 19 * 0, 10)	0.4424 (0.0086)	0.3620 (0.0052)	0.1956 (0.0091)	0.4498 (0.0042)	0.3075 (0.0013)	0.2427 (0.0016)
(30, 20, 10, 19 * 0)	0.3774 (0.0114)	0.3483 (0.0028)	0.2743 (0.0073)	0.4222 (0.0016)	0.3144 (0.0008)	0.2634 (0.0009)

Table 9: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 0.5$, $\tau_2 = 1.5$).

Sampling scheme	Without Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.4029 (0.0192)	0.3855 (0.0216)	0.2115 (0.0096)	0.4115 (0.0160)	0.3486 (0.0155)	0.2399 (0.0078)
(20, 15, 5, 14 * 0)	0.4006 (0.0212)	0.3272 (0.0175)	0.2722 (0.0102)	0.4107 (0.0172)	0.2914 (0.0153)	0.2980 (0.0115)
(30, 15, 14 * 0, 15)	0.3833 (0.0198)	0.4472 (0.0310)	0.1695 (0.0192)	0.3943 (0.0167)	0.4091 (0.0199)	0.1967 (0.0110)
(30, 15, 15, 14 * 0)	0.4038 (0.0202)	0.3258 (0.0171)	0.2703 (0.0103)	0.4139 (0.0165)	0.2900 (0.0150)	0.2962 (0.0114)
(30, 15, 1 * 15)	0.4585 (0.0187)	0.3364 (0.0122)	0.2051 (0.0109)	0.4601 (0.0166)	0.3076 (0.0105)	0.2323 (0.0081)
(30, 20, 19 * 0, 10)	0.4017 (0.0155)	0.4118 (0.0204)	0.1864 (0.0107)	0.4079 (0.0136)	0.3833 (0.0145)	0.2088 (0.0076)
(30, 20, 10, 19 * 0)	0.3945 (0.0192)	0.3275 (0.0145)	0.2780 (0.0093)	0.4024 (0.0162)	0.2992 (0.0130)	0.2984 (0.0106)

Table 10: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 0.5$, $\tau_2 = 1.5$).

Sampling scheme	With Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.4081 (0.0137)	0.3709 (0.0053)	0.2211 (0.0111)	0.4663 (0.0054)	0.3178 (0.0009)	0.2159 (0.0034)
(20, 15, 5, 14 * 0)	0.3546 (0.0205)	0.3598 (0.0040)	0.2856 (0.0108)	0.4365 (0.0020)	0.3171 (0.0006)	0.2464 (0.0010)
(30, 15, 14 * 0, 15)	0.4411 (0.0100)	0.3929 (0.0077)	0.1660 (0.0164)	0.5180 (0.0137)	0.3310 (0.0018)	0.1511 (0.0123)
(30, 15, 15, 14 * 0)	0.3568 (0.0202)	0.3593 (0.0039)	0.2840 (0.0107)	0.4398 (0.0022)	0.3154 (0.0007)	0.2449 (0.0010)
(30, 15, 1 * 15)	0.4361 (0.0186)	0.3488 (0.0050)	0.2151 (0.0118)	0.5069 (0.0120)	0.3105 (0.0017)	0.1827 (0.0069)
(30, 20, 19 * 0, 10)	0.4336 (0.0089)	0.3772 (0.0058)	0.1893 (0.0109)	0.4733 (0.0069)	0.3242 (0.0012)	0.2026 (0.0054)
(30, 20, 10, 19 * 0)	0.3513 (0.0192)	0.3587 (0.0034)	0.2900 (0.0099)	0.4236 (0.0013)	0.3181 (0.0006)	0.2583 (0.0007)

Table 11: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 1.0$, $\tau_2 = 1.5$).

Sampling scheme	Without Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.3906 (0.0162)	0.4030 (0.0260)	0.2064 (0.0110)	0.4306 (0.0144)	0.3222 (0.0164)	0.2472 (0.0088)
(20, 15, 5, 14 * 0)	0.3690 (0.0167)	0.3697 (0.0197)	0.2613 (0.0098)	0.4028 (0.0133)	0.2959 (0.0153)	0.3013 (0.0123)
(30, 15, 14 * 0, 15)	0.3577 (0.0199)	0.4888 (0.0468)	0.1534 (0.0218)	0.4255 (0.0150)	0.3796 (0.0192)	0.1949 (0.0105)
(30, 15, 15, 14 * 0)	0.3727 (0.0153)	0.3701 (0.0189)	0.2572 (0.0095)	0.4070 (0.0123)	0.2956 (0.0146)	0.2974 (0.0115)
(30, 15, 1 * 15)	0.3934 (0.0123)	0.4069 (0.0218)	0.1997 (0.0112)	0.4392 (0.0114)	0.3177 (0.0124)	0.2430 (0.0076)
(30, 20, 19 * 0, 10)	0.4064 (0.0141)	0.4186 (0.0274)	0.1750 (0.0118)	0.4441 (0.0131)	0.3483 (0.0158)	0.2076 (0.0077)
(30, 20, 10, 19 * 0)	0.3850 (0.0144)	0.3476 (0.0167)	0.2674 (0.0084)	0.4109 (0.0120)	0.2911 (0.0144)	0.2980 (0.0103)

Table 12: AEs and MSEs (in parenthesis) of the estimates of θ_1 , θ_2 , and θ_3 with different values of n and m ($\theta_1 = 5/12 = 0.4167$, $\theta_2 = 1/3 = 0.3333$, and $\theta_3 = 1/4 = 0.25$, $\tau_1 = 1.0$, $\tau_2 = 1.5$).

Sampling scheme	With Ordering					
	MLE			Bayes		
	θ_1	θ_2	θ_3	θ_1	θ_2	θ_3
(20, 15, 14 * 0, 5)	0.4196 (0.0084)	0.3621 (0.0050)	0.2182 (0.0091)	0.5255 (0.0189)	0.2796 (0.0044)	0.1948 (0.0058)
(20, 15, 5, 14 * 0)	0.3786 (0.0099)	0.3510 (0.0031)	0.2704 (0.0076)	0.4740 (0.0068)	0.2971 (0.0021)	0.2289 (0.0018)
(30, 15, 14 * 0, 15)	0.4381 (0.0065)	0.3959 (0.0080)	0.1660 (0.0163)	0.6173 (0.0469)	0.2527 (0.0092)	0.1300 (0.0167)
(30, 15, 15, 14 * 0)	0.3818 (0.0098)	0.3515 (0.0032)	0.2667 (0.0076)	0.4793 (0.0078)	0.2950 (0.0024)	0.2257 (0.0021)
(30, 15, 1 * 15)	0.4238 (0.0067)	0.3659 (0.0047)	0.2103 (0.0091)	0.5758 (0.0323)	0.2590 (0.0074)	0.1652 (0.0096)
(30, 20, 19 * 0, 10)	0.4413 (0.0068)	0.3713 (0.0063)	0.1874 (0.0103)	0.5625 (0.0302)	0.2630 (0.0072)	0.1745 (0.0090)
(30, 20, 10, 19 * 0)	0.3759 (0.0088)	0.3489 (0.0026)	0.2751 (0.0066)	0.4615 (0.0054)	0.2994 (0.0021)	0.2391 (0.0013)

References

- Alshenawy, R., Haj Ahmad, H., and Al-Alwan, A. (2022). Progressive censoring schemes for marshall-olkin pareto distribution with applications: Estimation and prediction. *Plos one*, 17(7):e0270750.
- Balakrishnan, N. and Cramer, E. (2014). The art of progressive censoring. *Statistics for industry and technology*, 138.
- Balakrishnan, N., Cramer, E., and Kundu, D. (2023). *Hybrid Censoring Know-How: Design and Implementation*. Academic Press, London, United Kingdom.
- Balakrishnan, N., Koutras, M., Milienos, F., and Pal, S. (2016). Piecewise linear approximations for cure rate models and associated inferential issues. *Methodology and Computing in Applied Probability*, 18(4):937–966.
- Cox, D. R. (1959). The analysis of exponentially distributed life-times with two types of failure. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 21(2):411–421.

- Cox, D. R. and Oakes, D. (1984). *Analysis of Survival Data*. Chapman & Hall.
- Ganguly, A., Mitra, D., Balakrishnan, N., and Kundu, D. (2024). A flexible model based on piecewise linear approximation for the analysis of left truncated right censored data with covariates, and applications to worcester heart attack study data and channing house data. *Statistics in Medicine*, 43(2):233–255.
- Goodman, M. S., Li, Y., and Tiwari, R. C. (2011). Detecting multiple change points in piecewise constant hazard functions. *Journal of Applied Statistics*, 38(11):2523–2532.
- Kundu, D. (2008). Bayesian inference and life testing plan for the weibull distribution in presence of progressive censoring. *Technometrics*, 50(2):144–154.
- Kundu, D. (2022a). Bivariate semi-parametric singular family of distributions and its applications. *Sankhya B*, 84(2):846–872.
- Kundu, D. (2022b). Bivariate semi-parametric singular family of distributions and its applications. *Sankhya, Ser B*, 84(Part 2):846 – 872.
- Kundu, D. and Gupta, A. K. (2013). Bayes estimation for the marshall–olkin bivariate weibull distribution. *Computational Statistics & Data Analysis*, 57(1):271–281.
- Kundu, D. and Joarder, A. (2006). Analysis of type-ii progressively hybrid censored data. *Computational Statistics & Data Analysis*, 50(10):2509–2528.
- Lawless, J. F. (2011). *Statistical models and methods for lifetime data*. John Wiley & Sons.
- Maurya, R. K., Tripathi, Y. M., Sen, T., and Rastogi, M. K. (2019). On progressively censored inverted exponentiated rayleigh distribution. *Journal of Statistical Computation and Simulation*, 89(3):492–518.
- Meeker, W. Q., Escobar, L. A., and Pascual, F. G. (2022). *Statistical methods for reliability data*. John Wiley & Sons.

- Mondal, S. and Kundu, D. (2019). Point and interval estimation of weibull parameters based on joint progressively censored data. *Sankhya b*, 81:1–25.
- Murphy, A. (1996). A piecewise-constant hazard-rate model for the duration of unemployment in single-interview samples of the stock of unemployed. *Economics Letters*, 51(2):177–183.
- Peña, E. A. and 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.
- Samanta, D., Ganguly, A., Gupta, A., and Kundu, D. (2019). On classical and bayesian order restricted inference for multiple exponential step stress model. *Statistics*, 53:177–195.
- Samanta, D. and Kundu, D. (2023). Bivariate semi-parametric model: Bayesian inference. *Methodology and Computing in Applied Probability*, 25(4):87.
- Samanta, D. and Kundu, D. (2025). Absolute continuous semi-parametric bivariate distribution. *Sankhya B*, 87:220–248.
- Sarhan, A. M., Manshi, T., and Smith, B. (2023). Statistical inference of reliability distributions based on progressively hybrid censoring data. *Scientific African*, 21:e01808.
- Wang, B. X., Yu, K., and Jones, M. (2010). Inference under progressively Type II right-censored sampling for certain lifetime distributions. *Technometrics*, 52(4):453–460.