

Statistical Inference for Pareto Distribution under Generalized Progressive Hybrid Censoring Scheme

M. M. Mohie El-Din^{1,*}, M. Nagy^{2,3} and A. R. Shafay³

¹ Department of Mathematics, Faculty of Science, Al-Azhar University, Cairo, Egypt

² Department of Statistics and Operation Research, Faculty of Science, King Saud University, Riyadh, Saudi Arabia

³ Department of Mathematics, Faculty of Science, Fayoum University, Fayoum, Egypt

Received: 25 Jun. 2018, Revised: 30 Mar. 2019, Accepted: 31 Mar. 2019

Published online: 1 May 2019

Abstract: In this paper, the maximum likelihood and Bayesian estimations are developed based on generalized progressive hybrid censoring sample from the Pareto distribution. The Bayesian estimates are computed under three different loss functions. Also, the point and interval Bayesian predictions for unobserved failures from the same sample and that from the future sample are derived. Moreover, a Monte Carlo simulation study is carried out to compare the performance of the maximum likelihood and Bayesian estimates. Finally, numerical example is presented for illustrating all the inferential procedures developed here.

Keywords: Bayesian estimation; Bayesian prediction; Pareto distribution; Maximum likelihood estimation; generalized progressive hybrid censoring sample.

1 Introduction

Consider an experiment in which n units are placed on life test. In progressive censoring schemes, m units complete failures are going to be observed. When the first failure is observed, R_1 of the $n - 1$ surviving units are randomly selected and removed. At the second observed failure, R_2 of the $n - R_1 - 2$ surviving units are randomly selected and removed. The experiment finally terminates at the time of the m^{th} failure when all remaining $R_m = n - R_1 - \dots - R_{m-1} - m$ surviving units are removed. The censoring numbers $\{R_i, i = 1, \dots, m - 1\}$ are prefixed. We will denote the m ordered failure times thus observed by $X_{1:m:n}, \dots, X_{m:m:n}$. It is evident that $n = m + \sum_{k=1}^m R_k$. The resulting m ordered values which are obtained from this type of censoring are referred to as progressively Type-II right censored order statistics. Several authors have studied progressive Type-II censoring and properties of order statistics arising from such a progressively censored life test. Some key references are Aggarwala and Balakrishnan [1], Cramer and Iliopoulos [2], Raqab et al. [3], Mohie El-Din and Shafay [4], and Balakrishnan and Cohen [5].

The disadvantages of the progressive Type-II censoring scheme are that the time of the experiment can be very long if the units are highly reliable. Therefore, Kundu and Joarder [6] and Childs et al. [7] proposed a progressive hybrid censoring scheme (PHCS), in this life-testing the experiment is terminated at time $\min\{X_{m:m:n}, T\}$, where $T \in (0, \infty)$ pre-fixed in advance. Under PHCS, the time on experiment will be no more than T . Some recent studies on PHCS have been carried out by many authors including Lin et al. [8], Lin and Huang. [9], and Hemmati and Khorram [10]. On the other hand, the disadvantages of the PHCS is that it cannot be applied when very few failures may occur before time T . For this reason, Cho et al. [11] propose a generalized PHCS which allows us to observe a pre-specified number of failures. So, the certain number of failures and their lifetimes are always provided under the generalized PHCS. The life-testing experiment based on this censoring scheme can save both the total time on tests and the cost induced by failures of the units. Moreover, the efficiency of statistical estimation is increased due to more failed observations. The detailed description and its advantages will be described in the next section.

* Corresponding author e-mail: mn112@fayoum.edu.eg

In this paper, the underlying distribution is assumed to be the Pareto distribution which introduced by Pareto [12] as a model for the distribution of income, with the probability density function (PDF) and cumulative distribution function (CDF) given by

$$f(x|\alpha, \beta) = \frac{\alpha}{\beta} \left(\frac{\beta}{x}\right)^{(\alpha+1)}, x \geq \beta, \quad (1)$$

and

$$F(x|\alpha, \beta) = 1 - \left(\frac{\beta}{x}\right)^\alpha. \quad (2)$$

where $\alpha > 0$ and $\beta > 0$.

In recent years, its models in several different forms have been studied by many authors including Davis and Feldstein[13], Cohen and Whitten[14], and Grimshaw [15].

The rest of this paper is organized as follows. In Section 2, the description of the model of the generalized PHCS is presented. The maximum likelihood (ML) estimator and the Bayesian estimator under the squared error loss function for the unknown parameters are derived in Section 3. In Section 4, we derive the Bayesian prediction for the failure times of all units that are removed in all stages of censoring. Bayesian prediction for progressive order statistics from an unobserved future sample from the same distribution is derived in Section 5. Finally, in Section 6, a Monte Carlo simulation study is carried out to compare the performance of the ML and Bayesian estimates, for illustrating all the inferential methods developed here, numerical example is then presented.

2 The Model Description

Consider a life-testing experiment in which n identical units are put on test. Assume that $X_1, X_2, \dots, X_n$ denote the corresponding lifetimes from a distribution with CDF $F(x|\alpha, \beta)$ and PDF $f(x|\alpha, \beta)$. Generalized PHCS may be described as follows. For $T \in (0, \infty)$ and integers $k, m \in \{1, 2, \dots, n\}$ are pre-fixed such that $k < m$ with $R = (R_1, R_2, \dots, R_m)$ is also pre-fixed integers satisfying $n = m + R_1 + \dots + R_m$. At the time of first failure, R_1 of the remaining units are randomly removed. Similarly at the time of the second failure R_2 , of the remaining units are removed and so on. This process continues until, immediately following the terminated time $T^* = \max\{X_{k:m:n}, \min\{X_{m:m:n}, T\}\}$, at this time all the remaining units are removed from the experiment. This generalized PHCS modifies PHCS by allowing the experiment to continue beyond time T if very few failures had been observed up to time T . Under this scheme, the experimenter would ideally like to observe m failures, but is willing to accept a bare minimum of k failures. Let D denote the number of observed failures up to time T (see Fig. 1).

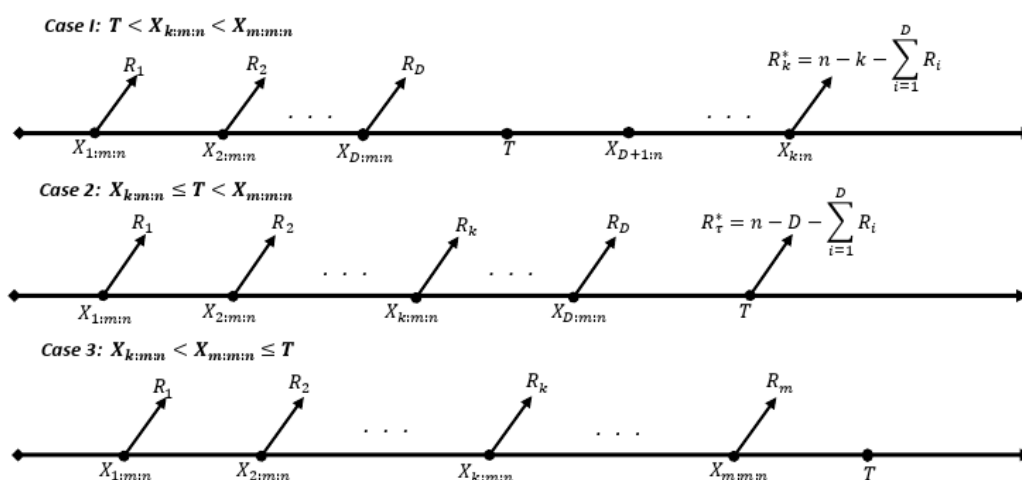


Fig. 1. Schematic representation of generalized progressive hybrid censoring scheme

Under generalized PHCS described above, we have one of the following types of observations:

1. Suppose that the k^{th} failure occurs after T , then the experiment terminates at $X_{k:m:n}$ and we will observe $\{X_{1:m:n} < \dots < X_{k:m:n}\}$.

2. Suppose that the k^{th} failure occurs before T and the m^{th} failure occurs after T , then the experiment terminates at T and we will observe $\{X_{1:m:n} < \dots < X_{k:m:n} < X_{k+1:n} < \dots < X_{D:n}\}$.
3. Suppose that the m^{th} failure occurs before T , then the experiment terminates at $X_{m:m:n}$ and we will observe $\{X_{1:m:n} < \dots < X_{k:m:n} < X_{k+1:n} < \dots < X_{m:m:n}\}$.

Given a generalized PHCS, the joint density function for three different cases are as follows:

$$f_{\underline{\mathbf{x}}}(\underline{\mathbf{x}}) = \left[\prod_{i=1}^{D^*} \sum_{j=i}^m (R_j^* + 1) \right] \prod_{i=1}^{D^*} f(x_{i:D^*:n}) [\bar{F}(x_{i:D^*:n})]^{R_i^*} [\bar{F}(T)]^{R_\tau^*}, \quad (3)$$

where

$$D^* = \begin{cases} k & \text{if } T < X_{k:m:n} < X_{m:m:n}, \\ D & \text{if } X_{k:m:n} \leq T < X_{m:m:n}, \\ m & \text{if } X_{k:m:n} < X_{m:m:n} \leq T, \end{cases} \quad (4)$$

$$R^* = \begin{cases} \left(R_1, \dots, R_D, 0, \dots, 0, R_k^* = n - k - \sum_{j=1}^D R_j \right) & \text{if } T < X_{k:m:n} < X_{m:m:n}, \\ (R_1, \dots, R_D) & \text{if } X_{k:m:n} \leq T < X_{m:m:n}, \\ (R_1, \dots, R_m) & \text{if } X_{k:m:n} < X_{m:m:n} \leq T, \end{cases} \quad (5)$$

with R_τ^* is the number of surviving units that are removed at T , given by

$$R_\tau^* = \begin{cases} 0 & \text{if } T < X_{k:m:n} < X_{m:m:n}, \\ n - D - \sum_{j=1}^D R_j & \text{if } X_{k:m:n} \leq T < X_{m:m:n}, \\ 0 & \text{if } X_{k:m:n} < X_{m:m:n} \leq T, \end{cases} \quad (6)$$

and

$$\underline{\mathbf{x}} = \begin{cases} (x_{1:m:n}, \dots, x_{k:m:n}) & \text{if } T < X_{k:m:n} < X_{m:m:n}, \\ (x_{1:m:n}, \dots, x_{D:m:n}) & \text{if } X_{k:m:n} \leq T < X_{m:m:n}, \\ (x_{1:m:n}, \dots, x_{m:m:n}) & \text{if } X_{k:m:n} < X_{m:m:n} \leq T. \end{cases} \quad (7)$$

Upon using (2) and (1) in (3), the likelihood function of α, β based on generalized PHCS can be obtained as

$$L(\alpha, \beta | \underline{\mathbf{x}}) = \left[\prod_{i=1}^{D^*} \sum_{j=i}^m (R_j^* + 1) \right] \alpha^{D^*} \left(\prod_{i=1}^{D^*} \frac{1}{x_i} \right) \exp \{ -\alpha [W(\underline{\mathbf{x}}) - n \ln \beta] \}, \quad (8)$$

where $W(\underline{\mathbf{x}}) = \sum_{i=1}^{D^*} (R_i^* + 1) \ln x_i + R_\tau^* \ln T$ and $x_i = x_{i:D^*:n}$ for simplicity of notation.

3 The ML and Bayesian Estimations

In this section, we derive the ML and Bayesian estimators of the unknown parameters α and β . It is clear that the likelihood function is monotone increasing function in β , so its maximum value $\hat{\beta}_{ML}$ will be attained at the minimum value x_1 of β .

From (8), The log-likelihood function of (α, β) is given by

$$\ln [L(\alpha, \beta | \underline{\mathbf{x}})] = \text{const.} + D^* \ln \alpha - \alpha [W(\underline{\mathbf{x}}) - n \ln \beta], \quad (9)$$

to maximize relative to α , differentiate (9) with respect to α and solve the equation

$$\frac{\partial \ln [L(\alpha, \beta | \underline{\mathbf{x}})]}{\partial \alpha} = 0,$$

and the ML estimator $\hat{\alpha}_{ML}$ of α is then obtained as

$$\hat{\alpha}_{ML} = \frac{D^*}{W(\underline{\mathbf{x}}) - n \ln x_1}. \quad (10)$$

The observed Fisher information matrix of the parameters α and β is given by

$$I(\hat{\alpha}, \hat{\beta}) = \begin{bmatrix} -\frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \alpha^2} & -\frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \alpha \partial \beta} \\ -\frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \beta \partial \alpha} & -\frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \beta^2} \end{bmatrix}_{(\hat{\alpha}, \hat{\beta})} \quad (11)$$

where

$$\frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \alpha^2} = -\frac{D^*}{\alpha^2}, \quad \frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \alpha \partial \beta} = \frac{n}{\beta}, \quad \text{and} \quad \frac{\partial^2 \ln L(\alpha, \beta | \underline{x})}{\partial \beta^2} = -\frac{n\alpha}{\beta^2},$$

and a $100(1 - \gamma)\%$ two-sided approximate confidence intervals for the parameters α and β are then

$$\left(\hat{\alpha} - z_{\gamma/2} \sqrt{V(\hat{\alpha})}, \hat{\alpha} + z_{\gamma/2} \sqrt{V(\hat{\alpha})} \right) \quad (12)$$

and

$$\left(\hat{\beta} - z_{\gamma/2} \sqrt{V(\hat{\beta})}, \hat{\beta} + z_{\gamma/2} \sqrt{V(\hat{\beta})} \right) \quad (13)$$

respectively, where $V(\hat{\alpha})$ and $V(\hat{\beta})$ are the estimated variances of $\hat{\alpha}_{ML}$ and $\hat{\beta}_{ML}$, which are given by the first and the second, diagonal element of $I^{-1}(\hat{\alpha}, \hat{\beta})$, and $z_{\gamma/2}$ is the upper $(\gamma/2)$ percentile of standard normal distribution.

For the Bayesian estimations, under the assumption that both parameters α and β are unknown, we may consider the joint prior density function of α and β which was suggested by Lwin [16] and generalized by Arnold and Press [17]. The generalized Lwin prior or the power-gamma prior is given by

$$\pi(\alpha, \beta) \propto \alpha^a \beta^{-1} \exp[-\alpha(\ln g - b \ln \beta)], \quad \alpha > 0, \quad 0 < \beta < h, \quad (14)$$

where a, b, g, h are positive constants and $h^b < g$.

Upon combining (8) and (14), given generalized PHCS, the posterior density function of α, β is obtained as

$$\begin{aligned} \pi^*(\alpha, \beta | \underline{x}) &= L(\alpha, \beta | \underline{x}) \pi(\alpha, \beta) / \int L(\alpha, \beta | \underline{x}) \pi(\alpha, \beta) d\alpha d\beta \\ &= I^{-1} \alpha^{D^*+a} \beta^{-1} \exp\{-\alpha[W(\underline{x}) - (n+b)\ln\beta + \ln g]\}, \end{aligned} \quad (15)$$

where

$$\begin{aligned} I &= \int_0^{x_0} \int_0^\infty \alpha^{D^*+a} \beta^{-1} \exp\{-\alpha[W(\underline{x}) - (n+b)\ln\beta + \ln g]\} d\alpha d\beta \\ &= \frac{\Gamma(D^*+a)}{n+b} [W(\underline{x}) - (n+b)\ln x_0 + \ln g]^{-(D^*+a)}, \end{aligned}$$

with $x_0 = \min(x_1, g)$.

If $\hat{\phi}$ is the estimator of the parameter ϕ , then the squared error loss function (SEL) is defined as,

$$L_1(\phi, \hat{\phi}) = (\hat{\phi} - \phi)^2, \quad (16)$$

then, the Bayesian estimate for a function $u(\alpha, \beta)$ under the loss function L_1 is the posterior mean and given by,

$$\begin{aligned} \hat{u}(\alpha, \beta)_{BS} &= E[u(\alpha, \beta) | \underline{x}] \\ &= \int_0^{x_0} \int_0^\infty u(\alpha, \beta) \pi^*(\alpha, \beta | \underline{x}) d\alpha d\beta. \end{aligned} \quad (17)$$

The LINEX loss function for ϕ and $\hat{\phi}$ is given by,

$$L_2(\phi, \hat{\phi}) = e^{v(\hat{\phi}-\phi)} - v(\hat{\phi}-\phi) - 1, \quad v \neq 0, \quad (18)$$

then, the Bayes estimate for the function of parameters $u(\alpha, \beta)$ under the loss function L_2 is given as,

$$\begin{aligned}\hat{u}(\alpha, \beta)_{BL} &= -\frac{1}{v} \ln \left[E(e^{-vu(\alpha, \beta)}) \right], \\ &= -\frac{1}{v} \ln \left[\int_0^{x_0} \int_0^\infty e^{-hu(\alpha, \beta)} \pi^*(\alpha, \beta | \underline{x}) d\alpha d\beta \right].\end{aligned}\quad (19)$$

Also, the generalized entropy loss function (GE) defined as,

$$L_3(\phi, \hat{\phi}) = \left(\frac{\hat{\phi}}{\phi}\right)^\kappa - \kappa \ln\left(\frac{\hat{\phi}}{\phi}\right) - 1, \kappa \neq 0, \quad (20)$$

and the corresponding Bayesian estimate for $u(\alpha, \beta)$ is given by,

$$\begin{aligned}\hat{u}(\alpha, \beta)_{BE} &= (E[(u(\alpha, \beta))^{-\kappa}])^{-1/\kappa}, \\ &= \left(\int_0^{x_0} \int_0^\infty (u(\alpha, \beta))^{-\kappa} \pi^*(\alpha, \beta) d\alpha d\beta \right)^{-1/\kappa}.\end{aligned}\quad (21)$$

By using (15), the Bayesian estimator of α under the squared error loss function is the mean of the posterior density function, given by

$$\begin{aligned}\hat{\alpha}_{BS} &= \int_0^{x_0} \int_0^\infty \alpha \pi^*(\alpha, \beta | \underline{x}) d\alpha d\beta \\ &= \frac{D^* + a}{W(\underline{x}) - (n+b) \ln x_0 + \ln g},\end{aligned}\quad (22)$$

and the Bayesian estimator of β under the squared error loss function is obtained as

$$\begin{aligned}\hat{\beta}_{BS} &= \int_0^{x_0} \int_0^\infty \beta \pi^*(\alpha, \beta | \underline{x}) d\alpha d\beta \\ &= I^{-1} x_0 \int_0^\infty \frac{\alpha^{D^*+a}}{\alpha(n+b)+1} \exp\{-\alpha[W(\underline{x}) - (n+b) \ln x_0 + \ln g]\} d\alpha \\ &= \frac{I^{-1} x_0}{(n+b)} [W(\underline{x}) - (n+b) \ln x_0 + \ln g]^{-(D^*+a)} \\ &\quad \times \int_0^\infty \frac{t^{D^*+a} e^{-t}}{t + [W(\underline{x}) - (n+b) \ln x_0 + \ln g] / (n+b)} dt \\ &= \frac{x_0}{\Gamma(D^*+a)} \Phi\left(D^*+a, \frac{[W(\underline{x}) - (n+b) \ln x_0 + \ln g]}{(n+b)}\right),\end{aligned}\quad (23)$$

where

$$\Phi(x, y) = \int_0^\infty \frac{t^x e^{-t}}{t+y} dt.$$

A partial tabulation of $\psi(x, y) = (y/\Gamma(x))\Phi(x-1, y)$ has been provided by Arnold and Press [17].

The Bayesian estimator of α and β under the LINEX loss function are obtained, respectively, as

$$\hat{\alpha}_{BL} = \frac{-1}{v} \ln \left\{ \frac{[W_2(\underline{x}) - (n+b) \ln x_0 + \ln g]^{(D^*+a)}}{[W_2(\underline{x}) - (n+b) \ln x_0 + \ln g + v]^{(D^*+a)}} \right\}, \quad (24)$$

$$\begin{aligned}\hat{\beta}_{BL} &= \frac{-1}{v} \ln \left\{ I^{-1} \int_0^{x_0} \frac{\Gamma(D^*+a+1)}{\beta} \exp(-v\beta) \right. \\ &\quad \times [W_2(\underline{x}) - (n+b) \ln \beta + \ln g]^{-(D^*+a+1)} d\beta \left. \right\}.\end{aligned}\quad (25)$$

The Bayesian estimator of α and β under the GE loss function are obtained, respectively, as

$$\hat{\alpha}_{BE} = \left\{ \frac{\Gamma(D^* + a - \kappa) [W_2(\underline{\mathbf{x}}) - (n+b)\ln x_0 + \ln g]^{(D^*+a)}}{\Gamma(D^* + a) [W_2(\underline{\mathbf{x}}) - (n+b)\ln x_0 + \ln g]^{(D^*+a-\kappa)}} \right\}^{\frac{-1}{\kappa}}, \quad (26)$$

$$\hat{\beta}_{BE} = \left\{ I^{-1} \int_0^{x_0} \frac{\Gamma(D^* + a + 1)}{\beta^{1-\kappa}} [W_2(\underline{\mathbf{x}}) - (n+b)\ln \beta + \ln g]^{-(D^*+a+1)} d\beta \right\}^{\frac{-1}{\kappa}}. \quad (27)$$

The marginal posterior distributions $\pi_1^*(\alpha|\underline{\mathbf{x}})$ and $\pi_2^*(\beta|\underline{\mathbf{x}})$ of the parameters α and β can be computed and written, respectively, as

$$\begin{aligned} \pi_1^*(\alpha|\underline{\mathbf{x}}) &= I^{-1} \int_0^{x_0} \alpha^{D^*+a} \beta^{-1} \exp\{-\alpha[W(\underline{\mathbf{x}}) - (n+b)\ln \beta + \ln g]\} d\beta \\ &= \frac{I^{-1} \alpha^{D^*+a-1}}{(n+b)} \exp\{-\alpha[W(\underline{\mathbf{x}}) - (n+b)\ln x_0 + \ln g]\} \end{aligned}$$

and

$$\begin{aligned} \pi_2^*(\beta|\underline{\mathbf{x}}) &= I^{-1} \int_0^\infty \alpha^{D^*+a} \beta^{-1} \exp\{-\alpha[W(\underline{\mathbf{x}}) - (n+b)\ln \beta + \ln g]\} d\alpha \\ &= \frac{I^{-1} \Gamma(D^* + a + 1)}{\beta} \{[W(\underline{\mathbf{x}}) - (n+b)\ln \beta + \ln g]\}^{-(D^*+a+1)}. \end{aligned}$$

In Bayesian statistics, we use the posterior distribution to determine Bayesian confidence intervals. Bayesian confidence intervals are also called credible intervals in order to make clear the distinction between the interpretation of Bayesian confidence intervals and classical confidence intervals. An interval (L, U) is a $100(1 - \gamma)\%$ credible interval for parameter θ (θ is α or β) if

$$\int_L^U \pi^*(\theta|\underline{\mathbf{x}}) d\theta = 1 - \gamma, \quad (28)$$

and the $100(1 - \gamma)\%$ HPD interval (L, U) for the parameter θ is given by the simultaneous solution of the equations

$$\int_L^U \pi^*(\theta|\underline{\mathbf{x}}) d\theta = 1 - \gamma \text{ and } \pi^*(L|\underline{\mathbf{x}}) = \pi^*(U|\underline{\mathbf{x}}) \quad (29)$$

If $\pi^*(\theta|\underline{\mathbf{x}})$ is unimodal, then the shortest interval will happen at the two quantiles a and b such that $\pi^*(a|\underline{\mathbf{x}}) = \pi^*(b|\underline{\mathbf{x}})$. So the HPD interval is basically the shortest interval for some given confidence level.

4 One-Sample Bayesian Prediction

For $\rho = 1, 2, \dots, R_j^*$, let $X_{\rho:R_j^*}$ denote the ρ^{th} order statistic out of R_j^* removed units at stage j . Then, the conditional density function of $X_{\rho:R_j^*}$, given the observed generalized PHCS, is given, see Basak et al. [18], by

$$f(X_{\rho:R_j^*}|\underline{\mathbf{x}}) = f(x|\underline{\mathbf{x}}) = \frac{R_j^*!}{(\rho-1)!(R_j^*-\rho)!} \frac{[F(x) - F(x_j)]^{\rho-1} [1 - F(x)]^{R_j^*-\rho} f(x)}{[1 - F(x_j)]^{R_j^*}}, \quad x > x_j, \quad (30)$$

where

$$j = \begin{cases} 1, \dots, k & \text{if } T < X_{k:m:n} < X_{m:m:n}, \\ 1, \dots, D, \tau & \text{if } X_{k:m:n} < T < X_{m:m:n}, \\ 1, \dots, m & \text{if } X_{k:m:n} < X_{m:m:n} < T, \end{cases}$$

with $x_\tau = T$.

By using (2) and (1) in (30), given generalized PHCS, the conditional density function of $X_{\rho:R_j^*}$ is then given as follows:

$$f(x|\underline{\mathbf{x}}) = \sum_{q=0}^{\rho-1} C_q \frac{\alpha}{x} \exp \{ -\alpha [\varpi_q (\ln x - \ln x_j)] \}, \quad x > x_j, \quad (31)$$

where $C_q = \frac{(-1)^q (\rho-1) R_j^*!}{(\rho-1)! (R_j^* - \rho)!}$ and $\varpi_q = q + R_j^* - \rho + 1$ for $q = 0, \dots, \rho - 1$.

Upon combining (15) and (31), the Bayesian predictive density function of $X_{\rho:R_j^*}$, given generalized PHCS, is obtained as

$$\begin{aligned} f^*(x|\underline{\mathbf{x}}) &= \int_0^{x_0} \int_0^\infty f(x|\underline{\mathbf{x}}) \pi^*(\alpha, \beta|\underline{\mathbf{x}}) d\alpha d\beta \\ &= \frac{I^{-1} \Gamma(D^* + a + 1)}{(n+b)} \sum_{q=0}^{\rho-1} \frac{C_q}{x} [W(\underline{\mathbf{x}}) - (n+b) \ln x_0 + \ln g + \varpi_q (\ln x - \ln x_j)]^{-(D^*+a+1)}. \end{aligned} \quad (32)$$

The Bayesian predictive survival function of $X_{\rho:R_j^*}$, given generalized PHCS, is given as

$$\begin{aligned} \bar{F}^*(t|\underline{\mathbf{x}}) &= \int_t^\infty f^*(x|\underline{\mathbf{x}}) dx \\ &= \frac{I^{-1} \Gamma(D^* + a)}{(n+b)} \sum_{q=0}^{\rho-1} \frac{C_q}{\varpi_q} [W(\underline{\mathbf{x}}) - (n+b) \ln x_0 + \ln g + \varpi_q (\ln t - \ln x_j)]^{-(D^*+a)}. \end{aligned} \quad (33)$$

The Bayesian point predictor of $X_{\rho:R_j^*}$ under the squared error loss function is the mean of the predictive density, given by

$$\hat{X}_{\rho:R_j^*} = \int_0^\infty x f^*(x|\underline{\mathbf{x}}) dx, \quad (34)$$

where $f^*(x|\underline{\mathbf{x}})$ is given as in (32). The Bayesian predictive bounds of $100(1-\gamma)\%$ two-sided equi-tailed (ET) interval for $X_{\rho:R_j^*}$ can be obtained by solving the following two equations:

$$\bar{F}^*(L_{ET}|\underline{\mathbf{x}}) = \frac{\gamma}{2} \quad \text{and} \quad \bar{F}^*(U_{ET}|\underline{\mathbf{x}}) = 1 - \frac{\gamma}{2}, \quad (35)$$

where $\bar{F}^*(t|\underline{\mathbf{x}})$ is given as in (33), and L_{ET} and U_{ET} denote the lower and upper bounds, respectively. On the other hand, for the highest posterior density (HPD) method, the following two equations need to be solved:

$$\bar{F}^*(L_{HPD}|\underline{\mathbf{x}}) - \bar{F}^*(U_{HPD}|\underline{\mathbf{x}}) = 1 - \gamma$$

and

$$f^*(L_{HPD}|\underline{\mathbf{x}}) - f^*(U_{HPD}|\underline{\mathbf{x}}) = 0,$$

where $f^*(x|\underline{\mathbf{x}})$ is as in (32), and L_{HPD} and U_{HPD} denote the HPD lower and upper bounds, respectively.

5 Two-Sample Bayesian Prediction

Let $Y_{1:\ell:N} \leq Y_{2:\ell:N} \leq \dots \leq Y_{\ell:\ell:N}$ be a future independent progressive Type-II censored sample from the same population with censoring scheme $\underline{\mathbf{S}} = (S_1, \dots, S_\ell)$. In this section, we develop a general procedure for deriving the point and interval predictions for $Y_{s:\ell:N}$, $1 \leq s \leq \ell$, based on the observed generalized PHCS. The marginal density function of $Y_{s:\ell:N}$ is given by Balakrishnan et al. [19] as

$$f_{Y_{s:\ell:N}}(y_s) = c(N, s) \sum_{q=0}^{s-1} c_{q,s-1} [1 - F(y_s)]^{M_{q,s-1}} f(y_s), \quad (36)$$

where $1 \leq s \leq \rho$, $c(N, s) = N(N - S_1 - 1) \dots (N - S_1 \dots - S_{s-1} + 1)$, $M_{q,s} = N - S_1 - \dots - S_{s-q-1} - s + q + 1$, and $c_{q,s-1} = (-1)^q \left\{ \left[\prod_{u=1}^q \sum_{v=s-q}^{s-q+u-1} (S_v + 1) \right] \left[\prod_{u=1}^{s-q-1} \sum_{v=u}^{s-q-1} (S_v + 1) \right] \right\}^{-1}$.

Upon substituting (2) and (1) in (36), the marginal density function of $Y_{s:\ell:N}$ is then obtained as

$$f_{Y_{s:\ell:N}}(y_s) = c(N, s) \sum_{q=0}^{s-1} c_{q,s-1} \frac{\alpha}{y_s} \exp \left\{ -\alpha \left[M_{q,s} \ln \left(\frac{y_s}{\beta} \right) \right] \right\}, \quad y_s > 0. \quad (37)$$

Upon combining (15) and (37), given generalized PHCS, the Bayesian predictive density function of $Y_{s:\ell:N}$ is obtained as

$$f_{Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) = \begin{cases} f_{1Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}), & 0 < y_s \leq x_0, \\ f_{2Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}), & y_s > x_0, \end{cases} \quad (38)$$

where

$$\begin{aligned} f_{1Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) &= \int_0^{y_s} \int_0^\infty f_{Y_{s:\ell:N}}(y_s | \underline{\mathbf{x}}) \pi^*(\alpha, \beta | \underline{\mathbf{x}}) d\alpha d\beta \\ &= I^{-1} \Gamma(D^* + a + 1) c(N, s) \sum_{q=0}^{\rho-1} \frac{c_{q,s-1}}{(n+b+M_{q,s}) y_s} \\ &\quad \times [W(\underline{\mathbf{x}}) - (n+b) \ln y_s + \ln g]^{-(D^*+a+1)}, \end{aligned} \quad (39)$$

and

$$\begin{aligned} f_{2Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) &= \int_0^{x_0} \int_0^\infty f_{Y_{s:\ell:N}}(y_s | \underline{\mathbf{x}}) \pi^*(\alpha, \beta | \underline{\mathbf{x}}) d\alpha d\beta \\ &= I^{-1} \Gamma(D^* + a + 1) c(N, s) \sum_{q=0}^{\rho-1} \frac{c_{q,s-1}}{(n+b+M_{q,s}) y_s} \\ &\quad \times [W(\underline{\mathbf{x}}) - (n+b+M_{q,s}) \ln x_0 + M_{q,s} \ln y_s + \ln g]^{-(D^*+a+1)}. \end{aligned} \quad (40)$$

From (38), we simply obtain the predictive survival function of $Y_{s:\ell:N}$, given generalized PHCS, as

$$\bar{F}_{Y_{s:\ell:N}}^*(t | \underline{\mathbf{x}}) = \begin{cases} \bar{F}_{1Y_{s:\ell:N}}^*(t | \underline{\mathbf{x}}), & 0 < t \leq x_0, \\ \bar{F}_{2Y_{s:\ell:N}}^*(t | \underline{\mathbf{x}}), & t > x_0, \end{cases} \quad (41)$$

where

$$\begin{aligned} \bar{F}_{1Y_{s:\ell:N}}^*(t | \underline{\mathbf{x}}) &= \int_t^{x_0} f_{1Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) dy_s + \int_{x_0}^\infty f_{2Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) dy_s \\ &= I^{-1} \Gamma(D^* + a) c(N, s) \sum_{q=0}^{\rho-1} \frac{c_{q,s-1}}{(n+b)(n+b+M_{q,s}) M_{q,s}} \\ &\quad \times \left\{ (n+b+M_{q,s}) [W(\underline{\mathbf{x}}) - (n+b) \ln x_0 + \ln g]^{-(D^*+a)} \right. \\ &\quad \left. - M_{q,s} [W(\underline{\mathbf{x}}) - (n+b) \ln t + \ln g]^{-(D^*+a)} \right\}, \end{aligned} \quad (42)$$

and

$$\begin{aligned} \bar{F}_{2Y_{s:\ell:N}}^*(t | \underline{\mathbf{x}}) &= \int_t^\infty f_{2Y_{s:\ell:N}}^*(y_s | \underline{\mathbf{x}}) dy_s \\ &= I^{-1} \Gamma(D^* + a) c(N, s) \sum_{q=0}^{\rho-1} \frac{c_{q,s-1}}{M_{q,s} (n+b+M_{q,s})} \\ &\quad \times [W(\underline{\mathbf{x}}) - (n+b+M_{q,s}) \ln x_0 + M_{q,s} \ln t + \ln g]^{-(D^*+a)}. \end{aligned} \quad (43)$$

The Bayesian point predictor of $Y_{s:\ell:N}$, $1 \leq s \leq m$, under the squared error loss function is the mean of the predictive density, given by

$$\hat{Y}_{s:\ell:N} = \int_0^{\infty} y_s f_{Y_{s:\ell:N}}^*(y_s | \mathbf{x}) dy_s, \quad (44)$$

where $f_{Y_{s:\ell:N}}^*(y_s | \mathbf{x})$ is given as in (38).

6 Numerical Results

6.1 Monte Carlo Simulation

In this section, a Monte Carlo simulation study is carried out to compare the performance of the ML and the Bayesian estimates under different sampling schemes. We used different values for n , m , k and T to generate 1000 generalized progressive hybrid censored samples from the Pareto distribution (with $\alpha = 2$ and $\beta = 5$). For comparison, we computed the estimated risk (ER) for each estimate by using the root mean square error and also computed the estimated bias (EB) for each estimate. Tables 1 and 2 present the values of EB and ER of the ML and Bayesian estimates for α and β , respectively. Also, we computed the 95% confidence intervals and Bayes credible intervals for α and β and checked whether the true values lay within the interval and recorded width of the interval. The estimated coverage probability was computed as the number of intervals that covered the true values divided by 1000, whereas the estimated expected width of the intervals was computed as the sum of the widths for all intervals divided by 1000. Tables 3 and 4 present the average width (A.W.) of 95% confidence intervals and corresponding coverage portability (C.P.) for $\hat{\alpha}$ and $\hat{\beta}$, respectively.

We perform a Monte Carlo Simulation study using different sample sizes (n), different effective samples sizes (m, k) and the following two censoring schemes

1. Scheme 1: $R_i = \frac{2(n-m)}{m}$ if i is odd and $R_i = 0$ if i is even.
2. Scheme 2: $R_i = \frac{2(n-m)}{m}$ if i is even and $R_i = 0$ if i is odd.

All Bayesian results are computed based on two different choices of the hyperparameters (a, b, g, h), namely,

1. Informative prior (IP): $a = 3$, $b = 0.22$, $g = 5.10$ and $h = 15.50$ (by letting the mean of the marginal prior distribution of α is 3 and its variance is 3, and the median of the marginal prior distribution of β is 5 and its third quartile is 10).
2. Noninformative prior (NIP): $a = -1$, $b = 0$, $g = 1$ and $h = \infty$.

6.2 Numerical example

In this numerical example, we generated generalized progressive hybrid censoring data from a sample of size $n = 25$. Suppose $m = 15$ and $R = (0, 0, 2, 0, 0, 2, 0, 0, 2, 0, 0, 2, 0, 0, 2)$, then we would have the following generated data: 5.006, 5.103, 5.358, 5.400, 5.437, 5.492, 5.669, 5.670, 5.866, 5.916, 6.202, 6.289, 9.195, 11.293, and 13.180. We consider different values for k and T , then we have the following three different generalized HPCS:

1. Scheme 1: Suppose $k = 10$, $T = 5.5$, since $T < X_{10:15:25}$, then the experiment would have terminated at $X_{10:15:25}$, with $R^* = (0, 0, 2, 0, 0, 2, 0, 0, 0, 11)$, $R_{\tau}^* = 0$, and we would have the following data: 5.006, 5.103, 5.358, 5.400, 5.437, 5.492, 5.669, 5.670, 5.866, and 5.916.
2. Scheme 2: Suppose $k = 10$, $T = 8$, since $X_{10:15:25} < T < X_{15:15:25}$, then the experiment would have terminated at $T = 8$, with $R^* = (0, 0, 2, 0, 0, 2, 0, 0, 2, 0, 0, 2)$, $R_{\tau}^* = 5$, and we would have the following data: 5.006, 5.103, 5.358, 5.400, 5.437, 5.492, 5.669, 5.670, 5.866, 5.916, 6.202, and 6.289.
3. Scheme 3: Suppose $k = 10$, $T = 15$, since $X_{15:15:25} < T$, then the experiment would have terminated at $X_{15:15:25}$, with $R^* = R$, $R_{\tau}^* = 0$, and we would have the following data: 5.006, 5.103, 5.358, 5.400, 5.437, 5.492, 5.669, 5.670, 5.866, 5.916, 6.202, 6.289, 9.195, 11.293, and 13.180.

We assume these data to have come from Pareto distribution with α and β are unknown. Based on the above generated generalized HPCS, Table 5 presents the point predictor and 95% ET and HPD prediction intervals of $X_{p:R_j^*}$ and Table 6 presents the point predictor and 95% ET and HPD prediction intervals of $Y_{s:\ell:N}$ from the future progressive censored sample of size $\ell = 10$ from a sample of size $N = 20$ with progressive censoring scheme $S = (1, 1, 0, 0, 2, 1, 0, 2, 1, 2)$.

Table 1: The values of ER of ML and Bayesian estimates for α based on the different censoring schemes.

(n, m, k)	Sch.	T	$\hat{\alpha}_{ML}$	$\hat{\alpha}_B$					
				$\hat{\alpha}_{BS}$		$\hat{\alpha}_{BL}$		$\hat{\alpha}_{BE}$	
				IP	NIP	IP	NIP	IP	NIP
(15,10,7)	1	5.5	1.3591	0.8776	1.0615	0.7445	0.8397	0.7480	0.9017
	2		1.4467	0.8959	1.1351	0.7561	0.8829	0.7623	0.9604
(20,10,7)	1		1.1597	0.8021	0.9340	0.6978	0.7872	0.7017	0.8395
	2		1.2498	0.8188	0.9796	0.6918	0.7892	0.6966	0.8467
(25,10,7)	1		1.1872	0.9393	1.0174	0.8666	0.9205	0.8661	0.9574
	2		1.1704	0.9060	0.9716	0.8248	0.8559	0.8174	0.8942
(30,20,15)	1		0.7244	0.6016	0.6365	0.5406	0.5746	0.5447	0.5903
	2		0.6932	0.5817	0.6057	0.5217	0.5471	0.5243	0.5608
(40,20,15)	1		0.7491	0.6177	0.6530	0.5525	0.5838	0.5560	0.5996
	2		0.7550	0.6244	0.6607	0.5601	0.5932	0.5642	0.6087
(50,20,15)	1		0.7505	0.6301	0.6590	0.5689	0.5956	0.5718	0.6093
	2		0.6939	0.5848	0.6042	0.5241	0.5451	0.5261	0.5575
(15,10,7)	1	6.5	1.2912	0.8337	1.0085	0.7073	0.7977	0.7106	0.8566
	2		1.3744	0.8511	1.0783	0.7183	0.8387	0.7242	0.9124
(20,10,7)	1		1.1018	0.7620	0.8873	0.6629	0.7479	0.6666	0.7975
	2		1.1873	0.7779	0.9307	0.6572	0.7497	0.6618	0.8044
(25,10,7)	1		1.1278	0.8923	0.9666	0.8232	0.8744	0.8228	0.9095
	2		1.1118	0.8607	0.9230	0.7836	0.8131	0.7766	0.8495
(30,20,15)	1		0.6882	0.5715	0.6046	0.5136	0.5459	0.5174	0.5608
	2		0.6585	0.5526	0.5755	0.4956	0.5197	0.4981	0.5328
(40,20,15)	1		0.7116	0.5868	0.6203	0.5249	0.5546	0.5282	0.5696
	2		0.7172	0.5932	0.6276	0.5321	0.5635	0.5360	0.5782
(50,20,15)	1		0.7129	0.5986	0.6260	0.5404	0.5658	0.5432	0.5789
	2		0.6592	0.5556	0.5740	0.4979	0.5179	0.4998	0.5296
(15,10,7)	1	8	1.2504	0.8074	0.9766	0.6849	0.7725	0.6881	0.8296
	2		1.3310	0.8242	1.0443	0.6957	0.8122	0.7013	0.8836
(20,10,7)	1		1.0670	0.7379	0.8592	0.6419	0.7242	0.6456	0.7724
	2		1.1498	0.7533	0.9013	0.6364	0.7261	0.6409	0.7790
(25,10,7)	1		1.0922	0.8641	0.9360	0.7972	0.8468	0.7969	0.8808
	2		1.0767	0.8335	0.8939	0.7588	0.7874	0.7520	0.8226
(30,20,15)	1		0.6664	0.5534	0.5855	0.4974	0.5286	0.5011	0.5430
	2		0.6377	0.5351	0.5573	0.4800	0.5033	0.4823	0.5159
(40,20,15)	1		0.6892	0.5682	0.6007	0.5083	0.5371	0.5115	0.5517
	2		0.6946	0.5745	0.6078	0.5153	0.5457	0.5191	0.5600
(50,20,15)	1		0.6904	0.5797	0.6063	0.5234	0.5479	0.5261	0.5606
	2		0.6384	0.5380	0.5559	0.4822	0.5015	0.4840	0.5129

Table 2: The values of EB of ML and Bayesian estimates for α based on the different censoring schemes.

(n, m, k)	$Sch.$	T	$\hat{\alpha}_{ML}$	$\hat{\alpha}_B$					
				$\hat{\alpha}_{BS}$		$\hat{\alpha}_{BL}$		$\hat{\alpha}_{BE}$	
				IP	NIP	IP	NIP	IP	NIP
(15,10,7)	1	5.5	0.6374	0.4429	0.2933	0.2605	0.2608	0.0415	0.0194
	2		0.6830	0.4664	0.3141	0.2823	0.3001	0.0717	0.0151
(20,10,7)	1		0.4033	0.2924	0.1589	0.1213	0.0600	0.1244	0.1955
	2		0.5623	0.4055	0.2615	0.2261	0.1969	0.0080	0.0750
(25,10,7)	1		0.1945	0.1085	0.0130	0.0483	0.1169	0.2803	0.3486
	2		0.3205	0.2146	0.0846	0.0499	0.0090	0.1846	0.2543
(30,20,15)	1		0.2804	0.2677	0.1956	0.1734	0.1284	0.0455	0.0147
	2		0.2789	0.2688	0.1969	0.1745	0.1270	0.0448	0.0134
(40,20,15)	1		0.3224	0.3024	0.2282	0.2067	0.1676	0.0817	0.0518
	2		0.3126	0.2932	0.2194	0.1979	0.1584	0.0730	0.0431
(50,20,15)	1		0.2978	0.2799	0.2067	0.1851	0.1446	0.0602	0.0301
	2		0.2898	0.2777	0.2053	0.1830	0.1371	0.0542	0.0230
(15,10,7)	1	6.5	0.5609	0.3897	0.2581	0.2293	0.2295	0.0365	0.0171
	2		0.6010	0.4104	0.2765	0.2484	0.2640	0.0631	0.0133
(20,10,7)	1		0.3549	0.2573	0.1398	0.1067	0.0528	0.1095	0.1720
	2		0.4949	0.3568	0.2301	0.1989	0.1733	0.0071	0.0660
(25,10,7)	1		0.1711	0.0955	0.0115	0.0425	0.1029	0.2466	0.3068
	2		0.2821	0.1889	0.0745	0.0439	0.0079	0.1625	0.2237
(30,20,15)	1		0.2468	0.2356	0.1721	0.1526	0.1130	0.0400	0.0130
	2		0.2455	0.2365	0.1733	0.1535	0.1118	0.0394	0.0118
(40,20,15)	1		0.2837	0.2661	0.2009	0.1819	0.1475	0.0719	0.0456
	2		0.2750	0.2580	0.1931	0.1741	0.1394	0.0642	0.0379
(50,20,15)	1		0.2621	0.2463	0.1819	0.1629	0.1273	0.0529	0.0265
	2		0.2550	0.2443	0.1806	0.1610	0.1207	0.0477	0.0203
(15,10,7)	1	8	0.5418	0.3765	0.2493	0.2215	0.2217	0.0353	0.0165
	2		0.5805	0.3964	0.2670	0.2400	0.2550	0.0609	0.0128
(20,10,7)	1		0.3428	0.2485	0.1350	0.1031	0.0510	0.1058	0.1662
	2		0.4780	0.3447	0.2223	0.1922	0.1674	0.0068	0.0637
(25,10,7)	1		0.1653	0.0923	0.0111	0.0410	0.0994	0.2382	0.2963
	2		0.2724	0.1824	0.0720	0.0424	0.0077	0.1569	0.2161
(30,20,15)	1		0.2384	0.2276	0.1663	0.1474	0.1091	0.0387	0.0125
	2		0.2371	0.2285	0.1674	0.1483	0.1080	0.0380	0.0114
(40,20,15)	1		0.2740	0.2570	0.1940	0.1757	0.1424	0.0695	0.0440
	2		0.2657	0.2492	0.1865	0.1682	0.1346	0.0621	0.0367
(50,20,15)	1		0.2531	0.2379	0.1757	0.1574	0.1229	0.0511	0.0256
	2		0.2463	0.2360	0.1745	0.1555	0.1166	0.0461	0.0196

Table 3: The values of ER of ML and Bayesian estimates for β based on the different censoring schemes.

(n, m, k)	Sch.	T	$\hat{\beta}_{ML}$	$\hat{\beta}_B$					
				$\hat{\beta}_{BS}$		$\hat{\beta}_{BL}$		$\hat{\beta}_{BE}$	
				IP	NIP	IP	NIP	IP	NIP
(15,10,7)	1	5.5	0.2515	0.1796	0.1825	0.1793	0.1850	0.1795	0.2895
	2		0.2625	0.1902	0.1931	0.1898	0.1956	0.1901	0.2963
(20,10,7)	1		0.1903	0.1389	0.1426	0.1390	0.1446	0.1391	0.1439
	2		0.1867	0.1374	0.1414	0.1375	0.1433	0.1376	0.1426
(25,10,7)	1		0.1404	0.1019	0.1062	0.1021	0.1078	0.1021	0.1072
	2		0.1397	0.1032	0.1079	0.1035	0.1096	0.1034	0.1090
(30,20,15)	1		0.1195	0.0854	0.0862	0.0855	0.0865	0.0855	0.0864
	2		0.1203	0.0855	0.0862	0.0855	0.0865	0.0855	0.0864
(40,20,15)	1		0.0848	0.0616	0.0624	0.0617	0.0627	0.0617	0.0626
	2		0.0908	0.0645	0.0649	0.0645	0.0651	0.0645	0.0650
(50,20,15)	1		0.0707	0.0499	0.0503	0.0499	0.0504	0.0499	0.0504
	2		0.0701	0.0495	0.0500	0.0496	0.0501	0.0496	0.0501
(15,10,7)	1	6.5	0.2398	0.1720	0.1771	0.1722	0.1805	0.1723	0.2750
	2		0.2344	0.1653	0.1708	0.1656	0.1744	0.1656	0.2342
(20,10,7)	1		0.1675	0.1192	0.1251	0.1197	0.1279	0.1196	0.1269
	2		0.1818	0.1329	0.1374	0.1332	0.1397	0.1331	0.1389
(25,10,7)	1		0.1334	0.0993	0.1041	0.0995	0.1059	0.0995	0.1052
	2		0.1327	0.0981	0.1025	0.0984	0.1041	0.0983	0.1035
(30,20,15)	1		0.1135	0.0847	0.0852	0.0847	0.0854	0.0847	0.0853
	2		0.1143	0.0812	0.0819	0.0812	0.0822	0.0812	0.0821
(40,20,15)	1		0.0846	0.0608	0.0617	0.0609	0.0620	0.0609	0.0619
	2		0.0870	0.0629	0.0636	0.0630	0.0638	0.0630	0.0637
(50,20,15)	1		0.0671	0.0474	0.0478	0.0474	0.0479	0.0474	0.0479
	2		0.0679	0.0494	0.0475	0.0494	0.0476	0.0494	0.0476
(15,10,7)	1	8	0.2326	0.1668	0.1718	0.1670	0.1751	0.1671	0.2429
	2		0.2274	0.1604	0.1657	0.1606	0.1692	0.1606	0.2272
(20,10,7)	1		0.1625	0.1156	0.1214	0.1161	0.1240	0.1160	0.1231
	2		0.1763	0.1299	0.1365	0.1305	0.1355	0.1303	0.1386
(25,10,7)	1		0.1294	0.0963	0.1009	0.0966	0.1027	0.0965	0.1020
	2		0.1287	0.0951	0.0994	0.0954	0.1010	0.0953	0.1004
(30,20,15)	1		0.1101	0.0822	0.0826	0.0822	0.0828	0.0822	0.0828
	2		0.1108	0.0788	0.0794	0.0788	0.0797	0.0788	0.0796
(40,20,15)	1		0.0821	0.0590	0.0598	0.0591	0.0601	0.0590	0.0600
	2		0.0844	0.0610	0.0617	0.0611	0.0619	0.0611	0.0618
(50,20,15)	1		0.0651	0.0460	0.0464	0.0460	0.0465	0.0460	0.0464
	2		0.0658	0.0479	0.0461	0.0480	0.0462	0.0479	0.0462

Table 4: The values of EB of ML and Bayesian estimates for β based on the different censoring schemes.

(n, m, k)	$Sch.$	T	$\hat{\beta}_{ML}$	$\hat{\beta}_B$					
				$\hat{\beta}_{BS}$		$\hat{\beta}_{BL}$		$\hat{\beta}_{BE}$	
				IP	NIP	IP	NIP	IP	NIP
(15,10,7)	1	5.5	0.1740	0.0164	0.0113	0.0087	0.0232	0.0117	0.0282
	2		0.1802	0.0208	0.0076	0.0129	0.0199	0.0161	0.0246
(20,10,7)	1		0.1312	0.0102	0.0115	0.0057	0.0186	0.0075	0.0159
	2		0.1285	0.0097	0.0110	0.0053	0.0178	0.0071	0.0151
(25,10,7)	1		0.0988	0.0034	0.0132	0.0005	0.0176	0.0017	0.0159
	2		0.0958	0.0005	0.0176	0.0034	0.0221	0.0023	0.0203
(30,20,15)	1		0.0838	0.0021	0.0049	0.0002	0.0072	0.0009	0.0063
	2		0.0855	0.0039	0.0030	0.0020	0.0053	0.0028	0.0044
(40,20,15)	1		0.0584	0.0023	0.0074	0.0034	0.0086	0.0030	0.0081
	2		0.0648	0.0039	0.0011	0.0029	0.0023	0.0033	0.0018
(50,20,15)	1		0.0502	0.0015	0.0025	0.0008	0.0033	0.0011	0.0030
	2		0.0499	0.0011	0.0029	0.0004	0.0037	0.0007	0.0034
(15,10,7)	1	6.5	0.1674	0.0095	0.0109	0.0018	0.0223	0.0048	0.0271
	2		0.1662	0.0074	0.0073	0.0003	0.0191	0.0027	0.0236
(20,10,7)	1		0.1181	0.0034	0.0111	0.0055	0.0179	0.0061	0.0152
	2		0.1244	0.0030	0.0106	0.0015	0.0170	0.0003	0.0145
(25,10,7)	1		0.0948	0.0017	0.0127	0.0005	0.0169	0.0002	0.0152
	2		0.0919	0.0003	0.0169	0.0034	0.0212	0.0022	0.0195
(30,20,15)	1		0.0837	0.0020	0.0043	0.0002	0.0065	0.0009	0.0056
	2		0.0845	0.0035	0.0029	0.0016	0.0051	0.0024	0.0042
(40,20,15)	1		0.0561	0.0013	0.0063	0.0023	0.0076	0.0019	0.0071
	2		0.0602	0.0007	0.0010	0.0017	0.0022	0.0013	0.0018
(50,20,15)	1		0.0482	0.0014	0.0003	0.0008	0.0012	0.0010	0.0008
	2		0.0475	0.0010	0.0028	0.0004	0.0036	0.0007	0.0033
(15,10,7)	1	8	0.1641	0.0070	0.0107	0.0015	0.0218	0.0018	0.0266
	2		0.1629	0.0011	0.0072	0.0003	0.0187	0.0027	0.0232
(20,10,7)	1		0.1157	0.0004	0.0108	0.0054	0.0175	0.0035	0.0149
	2		0.1219	0.0002	0.0104	0.0015	0.0167	0.0003	0.0142
(25,10,7)	1		0.0929	0.0012	0.0124	0.0005	0.0166	0.0002	0.0149
	2		0.0901	0.0003	0.0165	0.0033	0.0207	0.0021	0.0191
(30,20,15)	1		0.0820	0.0020	0.0025	0.0002	0.0048	0.0009	0.0039
	2		0.0828	0.0034	0.0018	0.0016	0.0040	0.0023	0.0031
(40,20,15)	1		0.0550	0.0013	0.0032	0.0012	0.0045	0.0016	0.0040
	2		0.0590	0.0006	0.0010	0.0005	0.0022	0.0001	0.0017
(50,20,15)	1		0.0472	0.0014	0.0003	0.0007	0.0011	0.0010	0.0008
	2		0.0465	0.0000	0.0027	0.0004	0.0035	0.0004	0.0032

Table 5: The A.W. of 95% confidence intervals and corresponding C.P. for $\hat{\alpha}_{ML}$ and $\hat{\alpha}_B$.

(n, m, k)	Sch.	T	$\hat{\alpha}_{ML}$		$\hat{\alpha}_B$			
					IP		NIP	
			A.W.	C.P.	A.W.	C.P.	A.W.	C.P.
(15,10,7)	1	5.500	4.7545	0.929	3.8462	0.969	4.4330	0.944
	2		4.8031	0.941	3.8846	0.974	4.4816	0.956
(20,10,7)	1		4.6767	0.935	3.8004	0.971	4.3551	0.950
	2		4.6601	0.934	3.7865	0.977	4.3386	0.949
(25,10,7)	1		4.6450	0.933	3.7726	0.968	4.3234	0.948
	2		4.6234	0.939	3.7636	0.976	4.3018	0.954
(30,20,15)	1		3.1399	0.919	2.6582	0.950	2.8184	0.934
	2		3.1219	0.927	2.6439	0.950	2.8004	0.942
(40,20,15)	1		3.1452	0.931	2.6626	0.953	2.8237	0.946
	2		3.1237	0.932	2.6453	0.960	2.8021	0.947
(50,20,15)	1		3.1885	0.917	2.6979	0.948	2.8669	0.932
	2		3.1349	0.933	2.6543	0.955	2.8134	0.948
(15,10,7)	1	6.500	4.6927	0.936	3.8166	0.974	4.3711	0.951
	2		4.6336	0.943	3.7758	0.974	4.3121	0.947
(20,10,7)	1		4.5025	0.940	3.6908	0.971	4.1810	0.955
	2		4.6001	0.950	3.7630	0.983	4.2786	0.965
(25,10,7)	1		4.3385	0.951	3.5797	0.978	4.0170	0.966
	2		4.4666	0.956	3.6712	0.984	4.1451	0.971
(30,20,15)	1		3.0507	0.943	2.5692	0.969	2.7292	0.958
	2		3.0554	0.936	2.5729	0.965	2.7339	0.951
(40,20,15)	1		3.0903	0.936	2.6109	0.960	2.7810	0.951
	2		3.1157	0.938	2.5730	0.965	2.7341	0.953
(50,20,15)	1		3.1224	0.927	2.6683	0.954	2.7307	0.942
	2		3.1292	0.937	2.5749	0.961	2.7377	0.952
(15,10,7)	1	8.000	4.6880	0.943	3.8041	0.979	4.5587	0.958
	2		4.6721	0.944	3.7534	0.974	4.5996	0.938
(20,10,7)	1		4.4990	0.945	3.6678	0.971	4.3685	0.960
	2		4.5988	0.966	3.7451	0.989	4.4266	0.981
(25,10,7)	1		4.3326	0.969	3.5567	0.988	4.2045	0.984
	2		4.4654	0.973	3.6588	0.992	4.1326	0.988
(30,20,15)	1		3.0438	0.967	2.5457	0.988	2.7167	0.982
	2		3.0443	0.945	2.5560	0.980	2.9721	0.960
(40,20,15)	1		3.0878	0.941	2.5880	0.967	2.8769	0.956
	2		3.1103	0.944	2.5466	0.970	2.8722	0.959
(50,20,15)	1		3.1110	0.937	2.6559	0.960	2.7824	0.952
	2		3.1217	0.941	2.5462	0.967	2.7523	0.956

Table 6: The A.W. of 95% confidence intervals and C.P. percentages for $\hat{\beta}_{ML}$ and $\hat{\beta}_B$.

(n, m, k)	Sch.	T	$\hat{\beta}_B$					
			$\hat{\beta}_{ML}$		IP		NIP	
			A.W.	C.P.	A.W.	C.P.	A.W.	C.P.
(15,10,7)	1	5.500	0.6030	0.914	0.5404	0.932	0.5815	0.926
	2		0.5933	0.932	0.5330	0.956	0.5718	0.944
(20,10,7)	1		0.4623	0.945	0.4127	0.959	0.4408	0.957
	2		0.4651	0.936	0.4149	0.960	0.4436	0.948
(25,10,7)	1		0.3810	0.940	0.3374	0.964	0.3595	0.952
	2		0.3802	0.933	0.3368	0.957	0.3587	0.945
(30,20,15)	1		0.2910	0.935	0.2617	0.959	0.2694	0.947
	2		0.2917	0.941	0.2623	0.965	0.2702	0.953
(40,20,15)	1		0.2231	0.942	0.1967	0.955	0.2016	0.954
	2		0.2242	0.948	0.1977	0.972	0.2026	0.960
(50,20,15)	1		0.1813	0.936	0.1565	0.952	0.1598	0.948
	2		0.1833	0.944	0.1583	0.968	0.1618	0.956
(15,10,7)	1	6.500	0.6101	0.933	0.5390	0.945	0.5798	0.939
	2		0.5984	0.944	0.5240	0.969	0.5628	0.957
(20,10,7)	1		0.4653	0.959	0.4037	0.972	0.4318	0.968
	2		0.4654	0.935	0.4137	0.957	0.4421	0.945
(25,10,7)	1		0.4572	0.956	0.3284	0.978	0.3505	0.963
	2		0.3792	0.947	0.3278	0.968	0.3497	0.959
(30,20,15)	1		0.2869	0.948	0.2589	0.973	0.2664	0.958
	2		0.2788	0.954	0.2588	0.977	0.2662	0.966
(40,20,15)	1		0.2319	0.955	0.1933	0.968	0.1977	0.967
	2		0.2422	0.935	0.1953	0.953	0.2000	0.946
(50,20,15)	1		0.1872	0.947	0.1475	0.964	0.1508	0.966
	2		0.1882	0.936	0.1572	0.958	0.1606	0.949
(15,10,7)	1	8.000	0.5958	0.935	0.5234	0.953	0.5574	0.947
	2		0.5788	0.956	0.5178	0.975	0.5457	0.969
(20,10,7)	1		0.4478	0.967	0.3898	0.978	0.4026	0.970
	2		0.4582	0.945	0.4078	0.965	0.4237	0.958
(25,10,7)	1		0.3665	0.963	0.3123	0.985	0.3450	0.976
	2		0.3657	0.954	0.3122	0.976	0.3244	0.968
(30,20,15)	1		0.2824	0.956	0.2453	0.978	0.2609	0.967
	2		0.2823	0.963	0.2453	0.986	0.2627	0.974
(40,20,15)	1		0.2137	0.964	0.1788	0.976	0.1922	0.970
	2		0.2160	0.943	0.1790	0.967	0.1945	0.954
(50,20,15)	1		0.1668	0.957	0.1402	0.975	0.1435	0.969
	2		0.1766	0.944	0.1532	0.966	0.1545	0.956

Table 7: Bayesian point predictor and 95% ET and HPD prediction intervals for $X_{\rho:R_j^*}$.

Sch.	R_j^*	ρ	IP			NIP		
			$\hat{X}_{\rho:R_j^*}$	ET interval	HPD interval	$\hat{X}_{\rho:R_j^*}$	ET interval	HPD interval
1	3	1	6.576	(5.381,10.931)	(5.358,9.410)	6.712	(5.281,11.665)	(5.258,9.826)
		2	11.236	(5.113,30.314)	(5.366,22.016)	16.422	(5.112,36.228)	(5.364,24.904)
	6	1	6.740	(5.515,11.204)	(5.492,9.645)	6.880	(5.316,11.956)	(5.492,10.072)
		2	11.515	(5.240,31.073)	(5.500,22.567)	16.792	(5.340,37.134)	(5.498,25.526)
	10	1	6.118	(5.421,6.735)	(5.516,6.554)	6.131	(5.321,6.815)	(5.916,6.606)
		2	6.350	(5.579,7.318)	(5.621,7.051)	6.379	(5.578,7.485)	(5.920,7.162)
		3	6.619	(5.634,7.969)	(5.884,7.588)	6.668	(5.634,8.250)	(6.088,7.777)
		4	6.938	(5.779,8.748)	(6.016,8.275)	7.014	(5.783,9.181)	(6.106,8.562)
		5	7.324	(6.260,9.720)	(6.098,9.093)	7.437	(6.256,10.363)	(6.179,9.514)
		6	7.809	(6.366,10.991)	(6.199,10.141)	7.972	(6.681,11.936)	(6.170,10.755)
		7	8.442	(6.509,12.753)	(6.322,11.563)	8.684	(6.596,14.160)	(6.278,12.460)
		8	9.326	(6.736,15.405)	(6.469,13.643)	9.703	(6.836,17.587)	(6.407,15.003)
		9	10.701	(7.191,19.969)	(6.594,16.942)	11.385	(7.159,23.654)	(6.560,19.302)
		10	13.359	(7.683,30.119)	(6.862,24.283)	15.316	(7.636,37.687)	(6.741,28.552)
		11	25.139	(8.581,76.101)	(7.464,52.349)	52.847	(8.508,106.033)	(6.915,66.841)
2	3	1	6.887	(5.386,12.485)	(5.358,10.479)	7.160	(5.388,13.923)	(5.358,11.346)
		2	13.969	(5.735,41.606)	(5.365,28.698)	23.766	(5.762,55.147)	(5.364,35.469)
	6	1	7.059	(5.520,12.798)	(5.492,10.741)	7.339	(5.522,14.272)	(5.492,11.630)
		2	14.314	(5.878,42.646)	(5.500,29.416)	24.294	(5.906,56.526)	(5.498,36.356)
	9	1	7.540	(5.896,13.669)	(5.866,11.472)	7.839	(5.898,15.244)	(5.866,12.422)
		2	15.277	(6.279,45.550)	(5.874,31.419)	25.762	(6.308,60.375)	(5.872,38.831)
	12	1	8.084	(6.321,14.655)	(6.289,12.299)	8.404	(6.324,16.343)	(6.289,13.318)
		2	16.365	(6.731,48.835)	(6.298,33.685)	27.414	(6.763,64.729)	(6.296,41.632)
	13	1	9.166	(8.016,11.221)	(8.000,10.462)	9.591	(8.018,11.722)	(8.000,10.800)
		2	10.563	(8.173,14.494)	(8.013,13.080)	11.239	(8.185,15.750)	(8.012,13.927)
		3	13.806	(8.507,20.100)	(8.496,17.331)	15.550	(8.539,23.006)	(8.117,19.299)
		4	17.922	(9.087,33.267)	(8.839,26.859)	22.593	(9.152,41.204)	(8.315,31.661)
		5	28.908	(10.232,102.241)	(9.601,68.473)	55.678	(10.369,149.336)	(8.537,91.966)
3	3	1	7.068	(5.388,13.385)	(5.358,11.106)	7.342	(5.391,14.878)	(5.358,12.024)
		2	15.410	(5.778,48.861)	(5.365,33.000)	23.146	(5.811,63.859)	(5.364,40.596)
	6	1	7.245	(5.523,13.719)	(5.492,11.384)	7.526	(5.526,15.250)	(5.492,12.325)
		2	15.792	(5.922,50.083)	(5.500,33.825)	23.687	(5.956,65.456)	(5.498,41.611)
	9	1	7.738	(5.899,14.654)	(5.866,12.159)	8.038	(5.902,16.288)	(5.866,13.164)
		2	16.856	(6.326,53.494)	(5.874,36.129)	25.193	(6.362,69.914)	(5.873,44.444)
	12	1	8.296	(6.325,15.710)	(6.289,13.036)	8.618	(6.328,17.463)	(6.289,14.113)
		2	18.060	(6.782,57.351)	(6.298,38.734)	26.890	(6.821,74.955)	(6.296,47.649)
	15	1	17.386	(13.255,32.925)	(13.180,27.319)	18.059	(13.261,36.597)	(13.180,29.577)
		2	37.597	(13.369,120.192)	(13.198,81.175)	53.975	(13.312,157.085)	(13.195,99.860)

Table 8: Bayesian point predictor and 95% ET and HPD prediction intervals for $Y_{s:\ell:N}$.

Scheme	s	IP			NIP		
		$\hat{Y}_{s:N}$	ET interval	HPD interval	$\hat{Y}_{s:N}$	ET interval	HPD interval
1	1	5.026	(4.797,5.308)	(4.783,5.292)	5.027	(4.780,5.332)	(4.764,5.314)
	2	5.124	(4.862,5.532)	(4.832,5.488)	5.131	(4.852,5.585)	(4.815,5.530)
	3	5.243	(4.932,5.802)	(4.885,5.719)	5.258	(4.927,5.894)	(4.871,5.788)
	4	5.374	(5.000,6.089)	(4.940,5.965)	5.398	(4.998,6.229)	(4.927,6.065)
	5	5.559	(5.077,6.528)	(5.001,6.333)	5.596	(5.076,6.744)	(4.989,6.483)
	6	5.771	(5.165,7.030)	(5.064,6.751)	5.825	(5.094,7.341)	(5.050,6.964)
	7	6.017	(5.254,7.621)	(5.127,7.241)	6.093	(5.250,8.056)	(5.111,7.535)
	8	6.414	(5.399,9.690)	(5.196,9.101)	6.529	(5.376,9.355)	(5.174,8.541)
	9	6.932	(5.941,10.144)	(5.697,9.144)	7.107	(5.530,11.165)	(5.237,9.916)
	10	11.817	(6.987,18.868)	(6.320,15.139)	14.327	(5.969,19.084)	(5.282,16.749)
2	1	5.037	(4.698,5.462)	(4.675,5.434)	5.042	(4.578,5.662)	(4.618,5.504)
	2	5.186	(4.792,5.803)	(4.745,5.733)	5.214	(4.602,5.803)	(4.697,5.762)
	3	5.369	(4.894,6.220)	(4.822,6.092)	5.428	(4.707,6.469)	(4.784,6.296)
	4	5.573	(4.996,6.671)	(4.904,6.479)	5.667	(4.899,7.040)	(4.875,6.773)
	5	5.866	(5.115,7.382)	(4.996,7.072)	6.015	(5.106,7.953)	(4.976,7.514)
	6	6.209	(5.235,8.217)	(5.092,7.765)	6.291	(5.230,9.053)	(5.080,8.398)
	7	6.617	(5.394,9.236)	(5.191,8.601)	7.585	(5.333,10.431)	(5.186,9.490)
	8	7.297	(5.552,11.176)	(5.300,10.136)	8.879	(5.562,13.121)	(5.300,11.537)
	9	8.226	(5.860,13.989)	(5.414,12.318)	10.980	(5.748,16.202)	(5.416,14.555)
	10	10.789	(6.604,16.191)	(5.470,13.126)	12.416	(6.356,19.407)	(5.459,17.213)
3	1	5.033	(4.737,5.400)	(4.718,5.377)	5.035	(4.506,5.445)	(4.684,5.318)
	2	5.162	(4.819,5.689)	(4.780,5.633)	5.178	(4.698,5.775)	(4.752,5.707)
	3	5.320	(4.908,6.038)	(4.849,5.936)	5.353	(4.897,6.179)	(4.826,6.053)
	4	5.495	(4.997,6.413)	(4.922,6.261)	5.548	(4.985,6.619)	(4.904,6.427)
	5	5.744	(5.102,6.995)	(5.003,6.752)	5.848	(5.095,7.021)	(5.104,6.954)
	6	5.744	(5.157,7.354)	(5.104,7.009)	6.154	(5.125,8.117)	(5.182,7.665)
	7	6.374	(5.347,8.475)	(5.179,7.993)	6.542	(6.015,9.104)	(6.175,8.470)
	8	6.933	(5.710,9.981)	(5.277,9.208)	7.187	(6.154,12.971)	(6.276,10.938)
	9	7.682	(5.756,12.101)	(5.381,10.891)	7.832	(6.254,16.669)	(6.380,15.016)
	10	8.304	(5.879,15.812)	(5.482,13.263)	8.349	(6.452,21.096)	(6.477,18.953)

7 Conclusions and discussion

From Tables 1 and 2, it can be seen that the performance of the ML estimators is quite close to that of the Bayesian estimators based noninformative priors, as expected. Thus, if we have no prior information on the unknown parameters, then it is always better to use the ML rather than the Bayesian estimators, because the Bayesian estimators are computationally more expensive. Also, the Bayesian method with informative priors is the best method for estimation under all different censoring schemes. Moreover, mean-squared error decreases when n and m increase.

From Tables 3 and 4, it can be seen that the average width of 95% confidence intervals and corresponding coverage probability are short than the average width of 99% confidence intervals and corresponding coverage probability. Also, a comparison of the results for the informative priors with the corresponding ones for non-informative priors reveals that the former produce more precise results, as we would expect. Moreover, the average width decreases when n and m increase.

From the results in Tables 5 and 6, we notice that, the point predictor of mean is between the upper and lower bounds of the prediction intervals. Also, the lower bounds are relatively insensitive while the upper bounds are more sensitive. Moreover, a comparison of the results for the informative priors with the corresponding ones for non-informative priors reveals that the former produce more precise results, as we would expect. Finally, the HPD prediction intervals seem to be more precise than the ET prediction intervals.

Acknowledgements

The authors would like to thank the editor, associate editor and two anonymous referees for their thorough review of the paper and their valuable suggestions that improved the original version of the manuscript.

References

- [1] R. Aggarwala, N. Balakrishnan, Some properties of progressive censored order statistics from arbitrary and uniform distributions with applications to inference and simulation. *J. Stat. Plan. and Infer.* **70** (1998) 35-49.
- [2] E. Cramer, G. Iliopoulos, Adaptive progressive Type-II censoring, *Test* **19** (2010) 342-358.
- [3] M. Z. Raqab, A. Asgharzadeh, R. Valiollahi, Prediction for Pareto distribution based on progressively Type-II censored samples, *Comput. Stat. Data Anal.* **54** (2010) 1732-1743.
- [4] M. M. Mohie El-Din, A. R. Shafay, One-and two-sample Bayesian prediction intervals based on progressively Type-II censored data, *Stat. Pap.* **54** (2013) 287-307.
- [5] N. Balakrishnan, A.C. Cohen, *Order statistics – inference: estimation methods*, Elsevier 2014.
- [6] D. Kundu, A. Joarder, Analysis of Type-II progressively hybrid censored data, *Comput. Stat. Data Anal.* **50** (2006) 2509-2528.
- [7] A. Childs, B. Chandrasekar, N. Balakrishnan, Exact likelihood inference for an exponential parameter under progressive hybrid censoring schemes, *In Stat. models and methods for biomedical and tech. Sys.* (2008) (319-330).
- [8] C.T. Lin, C.C. Chou, Y.L. Huang, Inference for the Weibull distribution with progressive hybrid censoring, *Comput. Stat. Data Anal.* **56** (2012) 451-467.
- [9] C.T. Lin, Y.L. Huang, On progressive hybrid censored exponential distribution, *J. Stat. Comput. Simul.* **82** (2012) 689-709.
- [10] F. Hemmati, E. Khorram, Statistical analysis of the log-normal distribution under Type-II progressive hybrid censoring schemes, *Commun. Stat. Simul. Comput.* **42** (2013) 52-75.
- [11] Y. Cho, H. Sun, K. Lee, Exact likelihood inference for an exponential parameter under generalized progressive hybrid censoring scheme, *Stat. Methodol.* **23** (2015) 18-34.
- [12] V. Pareto, *Cours d'Economie Politique*, Rouge et Cie, Paris 1897.
- [13] H.T. Davis, M.L. Feldstein, The generalized Pareto law as a model for progressively censored survival data, *Biometrika* **66** (1979) 299-306.
- [14] A.C. Cohen, B.J. Whitten, *Parameter Estimation in Reliability and Life Span Models*, Marcel Dekker, Inc, New York 1988.
- [15] S.D. Grimshaw, Computing maximum likelihood estimates for the generalized Pareto distribution, *Technometrics* **35** (1993) 185-191.
- [16] T. Lwin, Estimation of the tail of the Paretian law, *Scandinavian Actuarial J.* **55** (1972) 170-178.
- [17] B.C. Arnold, S.J. Press, Bayesian estimation and prediction for Pareto data. *J. Amer. Stat. Assoc.* **84** (1989) 1079-1084.
- [18] I. Basak, P. Basak, N. Balakrishnan, On some predictors of times to failure of censored items in progressively censored samples. *Comput. stat. Data Anal.* **50** (2006) 1313-1337.
- [19] N. Balakrishnan, A. Childs, B. Chandrasekar An efficient computational method for moments of order statistics under progressive censoring. *Stat. Probabil. Lett.* **60** (2002) 359-365.



Mostafa M. Mohie El-Din is Professor of Mathematical Statistics at Faculty of Science, Al-Azhar University, Egypt. His supervisor for graduated student of statistics. His research interests are in the areas of Statistical Inference, Order Statistics and Censored Data. He has published research articles in reputed international journals of Statistics. He is referee of several international journals in the frame of Statistics.



Magdy Nagy received the PhD degree in Mathematical Statistics at Al-Azhar University, Egypt. He is working as Assistant Professor in Department of Statistics and Operation Research, Faculty of Science, King Saud University, Riyadh, Saudi Arabia. Also, working as specialist Statistics and Mathematics in Department of Mathematics, Faculty of Science, Fayoum University, Fayoum, Egypt. His research interests are in the areas of Statistical Inference, Bayesian Estimation, Bayesian Prediction, Order Statistics and Censored Data. He has published research articles in reputed international journals of mathematical and statistics.



Ahmed R. Shafay received the PhD degree in Mathematical Statistics at Fayoum University, Egypt. His research interests are in the areas of Statistical Inference, Bayesian Estimation, Bayesian Prediction, Order Statistics and Censored Data. He has published research articles in reputed international journals of Statistics. He is referee of Statistical journals.