

Point prediction for the proportional hazards family based on progressive Type-II censoring with binomial removals

RAHMAT SADAT MESHKAT*, NAEIMEH DEHQANI
DEPARTMENT OF STATISTICS, YAZD UNIVERSITY, YAZD, IRAN

Received: August 21, 2018/ Revised: September 13, 2018/ Accepted: September 14, 2018

In this paper, some different predictors are presented for failure times of units censored in a progressively censored sample from proportional hazard rate models, where the number of units removed at each failure time follows a binomial distribution. The maximum likelihood predictors, best unbiased predictors and conditional median predictors are derived. Also, the Bayesian point predictors are investigated for the failure times of units with the three common loss function. Finally, a numerical example and a Monte Carlo simulation study are carried out to compare all the prediction methods discussed in this paper.

Keywords: Bayesian point predictor; Best unbiased predictor; Binomial removal; Conditional median predictor; Maximum likelihood predictor; Monte Carlo simulation; Progressive Type-II censoring; Proportional hazard rate model.

Mathematics Subject Classification (2010): 62N01, 62M20.

1 Introduction

In lifetime testing and analysis of reliability data, it is common to use the different censoring scheme in order to reduce costs and time. A censoring scheme, which can balance between (a) total time spent for the experiment; (b) number of units used in the experiment; and (c) the efficiency of statistical inference based on the results of the experiment, is desirable; see Ng et al. (2009). In some life-testing experiments, the experimenter seeks to remove units at different stages in the study for various reasons. This would lead to progressive censoring. In particular, a progressive Type-II censoring scheme is taken account of an important scheme in life-testing experiments.

The progressive Type-II censoring can be described as follows: Suppose n independent units are placed on a life test. Immediately following the first failure, r_1 surviving units are removed from the test at random. Then, immediately following the second

*Corresponding author: r.meshkat@gmail.com

failure, r_2 surviving units are removed from the test at random. This process continues until, at the time of the m -th failure, all the remaining $r_m = n - r_1 - r_2 - \cdots - r_{m-1} - m$ units are removed from the experiment. Suppose $X_1, \dots, X_n$ are the corresponding failure times of these units that are identically distributed with Probability Density Function (PDF) $f(\cdot)$ and Cumulative Distribution Function (CDF) $F(\cdot)$. Therefore, the joint density function $X_{1:m:n}, \dots, X_{m:m:n}$, a progressively Type-II censored order statistics with the censoring scheme $(r_1, \dots, r_m)$ is given by

$$f(x_{1:m:n}, \dots, x_{m:m:n}) = A \prod_{i=1}^m f(x_{i:m:n}) [1 - F(x_{i:m:n})]^{r_i}, \quad (1)$$

where $x_{1:m:n} < \cdots < x_{m:m:n}$ and $A(n, m-1) = n(n-1-r_1)(n-2-r_1-r_2)\cdots(n-m+1-r_1\cdots-r_{m-1})$. Note that $\sum_{i=1}^m r_i = n-m$, $r_i \geq 0$, $i = 1, \dots, m$; see Balakrishnan and Aggarwala (2000). In this scheme, $r_1, r_2, \dots, r_m$ are all pre-fixed. However, these numbers may occur at random in some practical situations. In some reliability experiments, an experimenter may distinguish that it is inappropriate or too dangerous to continue the testing on some of the tested units even though these units have not failed; Yuen and Tse (1996) and Zeinab (2008). In such cases, the pattern of removal at each failure is random. This leads to progressive censoring with random removals (PCR); see Table 1.

Table 1: A schematic representation of the progressive Type-II censoring with binomial removals that each unit leaves with equal probability p .

The numbers in life testing	Binomial removals	Remains
n	$R_1 \sim B(n-m, p)$	$n-1-r_1$
$n-1-r_1$	$R_2 \sim B(n-m-r_1, p)$	$n-2-r_1-r_2$
$\vdots$	$\vdots$	$\vdots$
$n-(m-2)-\sum_{j=1}^{m-2} r_j$	$R_{m-1} \sim B(n-m-\sum_{j=1}^{m-2} r_j, p)$	$n-(m-1)-\sum_{j=1}^{m-1} r_j$
$n-(m-1)-\sum_{j=1}^{m-1} r_j$	$R_m = n-m-\sum_{j=1}^{m-1} r_j$	0

The statistical inference on lifetime distributions under progressive censoring with random removals is discussed by several authors. Tse et al. (2000) presented the maximum likelihood estimations of Weibull distribution under progressive censoring with binomial removals. Wu and Chang (2003) and Wu et al. (2007) discussed the estimation of the Burr Type-XII distribution and Pareto distribution based on progressively censored samples with random removals, where the number of removed units has a discrete uniform removal pattern and binomial distribution. Xiang and Tse (2005) discussed the maximum likelihood estimation of the model parameters and derived the corresponding asymptotic variances based on Type-II progressive interval censoring with binomial removals. Wu et al. (2006) obtained Maximum Likelihood Estimator (MLE) and the

estimated expected test time for the two-parameter Gompertz distribution under progressive censoring with binomial removals. Recently, Weian et al. (2011) presented statistical analysis of generalized exponential distribution under progressive censoring with binomial removals. Prediction of future events on the basis of past and present knowledge is a fundamental problem in statistical inference. Viveros and Balakrishnan (1994) used the conditional method of inference to develop a conditional prediction interval for an observation from an independent future sample based on an observed progressively Type-II right censored sample. Some key references about prediction based on a progressively censored sample include Balakrishnan and Lin (2002), Basak et al. (2006) and Raqab et al. (2010). Recently, Asgharzadeh and Valiollahi (2010) discussed point prediction for the proportional hazards family under progressive Type-II censoring with fixed removal.

Suppose $F_0(\cdot)$ is a CDF with a corresponding hazard rate function $r_0(\cdot)$. The family of random variables with hazard rate function of the form $\{\theta r_0(\cdot) : \theta > 0\}$ is called PHR family and the CDF $F_0(\cdot)$ is called the baseline CDF of that family. Cox (1972) introduced this model and it has been extensively discussed in the statistical literature. This family of distributions includes several well-known lifetime distributions such as exponential, Pareto (Types I and II), Beta, Burr Type XII and so on. Let X be from proportional hazard family with the baseline CDF $F_0(\cdot)$. The distribution function of X is defined by

$$F(x; \theta) = 1 - [\bar{F}_0(x)]^\theta, \quad -\infty \leq c < x < d \leq \infty, \quad \theta > 0. \quad (2)$$

or equivalently

$$\bar{F}(x; \theta) = [\bar{F}_0(x)]^\theta, \quad -\infty \leq c < x < d \leq \infty, \quad \theta > 0,$$

where $\bar{F}_0(\cdot) = 1 - F_0(\cdot)$ is the baseline survival function and $F_0(c) = 0$ and $F_0(d) = 1$. The corresponding probability density function is as following

$$f(x; \theta) = \theta f_0(x) [\bar{F}_0(x)]^{\theta-1}, \quad -\infty \leq c < x < d \leq \infty, \quad \theta > 0, \quad (3)$$

where $f_0(\cdot)$ is the PDF of $F_0(\cdot)$. Lawless (2003), Ahmadi et al. (2009a,b) and Asgharzadeh and Valiollahi (2009, 2010) have extensively discussed the PHR model.

Consider $X_{1:m:n}, \dots, X_{m:m:n}$ indicate the progressively Type-II order statistics from the PHR model given in (2) or (3) obtained from a sample of size n with the censoring scheme $(r_1, \dots, r_m)$. To simplify the notation, we will use X_i in place of $X_{i:m:n}$. The purpose of this paper is to obtain the prediction of the life-lengths $Y = X_{i,(s)}$ ($s = 1, 2, \dots, r_i$, $i = 1, 2, \dots, m$) of all censored units in all m stages of censoring based on observed data $\mathbf{x} = (x_1, \dots, x_m)$, where $X_{i,(s)}$ indicates the s -th order statistic out of r_i removed units at stage i ($i = 1, 2, \dots, m$). In Section 2, we derive some different point predictors for failure times of units censored (non-Bayesian and Bayesian approaches). A numerical example and a Monte Carlo simulation study are carried out in Section 3 to compare all the prediction methods mentioned in this paper. Finally, some results is presented in Section 3.2.

2 Different methods for point prediction

Consider $X_{1:m:n}, \dots, X_{m:m:n}$ indicate the progressively Type-II order statistics from the PHR model given in (2) or (3) obtained from a sample of size n with the censoring scheme $(r_1, \dots, r_m)$. To simplify the notation, we will use X_i in place of $X_{i:m:n}$. Our interest is to obtain the prediction of the life-lengths $Y = X_{i,(s)}$ ($s = 1, 2, \dots, r_i$, $i = 1, 2, \dots, m$) of all censored units in all m stages of censoring based on observed data $\mathbf{x} = (x_1, \dots, x_m)$, where $X_{i,(s)}$ indicates the s -th order statistic out of r_i removed units at stage i , $i = 1, 2, \dots, m$.

2.1 Maximum likelihood predictor

The maximum likelihood is one of the most important and widely used methods in statistics. Kaminsky and Rhodin (1985), Basak et al. (2006), and Asgharzadeh and Valiollahi (2010) have used maximum likelihood prediction for predicting future observations. Our interest is to predict the unobservable future value of Y , having observed $\mathbf{x}$. Consider $\mathbf{X} = (X_1, \dots, X_m)$ and Y have the joint PDF $f(\mathbf{x}, y, \mathbf{r}; \theta, p)$ indexed by the parameters θ and p . Thus, we can define

$$\begin{aligned} L(y, \theta, p; \mathbf{x}, \mathbf{r}) &= f(\mathbf{x}, y, \mathbf{r}; \theta, p) \\ &= f_{Y|\mathbf{X}}(y|\mathbf{x}; \theta) f(\mathbf{x}|\mathbf{r}; \theta) f(\mathbf{r}; p), \end{aligned} \quad (4)$$

as the Predictive Likelihood Function (PLF) of y and θ and p , where the joint density function (1) based on PHR of $\mathbf{x} = (x_1, \dots, x_m)$ can be rewritten as

$$f(\mathbf{x}|\mathbf{r}; \theta) = A \left[\prod_{i=1}^m \frac{f_0(x_i)}{\bar{F}_0(x_i)} \right] \theta^m e^{-\theta T(\mathbf{x})}, \quad (5)$$

where $A = n(n-1-r_1)(n-2-r_1-r_2) \cdots (n-m+1-r_1 \cdots -r_{m-1})$ and $T(\mathbf{x}) = -\sum_{i=1}^m (r_i+1) \ln \bar{F}_0(x_i)$ and the joint probability distribution of removals vector $\mathbf{R} = (R_1, \dots, R_m)$ is given by

$$f(\mathbf{r}; p) = P(\mathbf{R} = \mathbf{r}; p) = B p^D (1-p)^E$$

with $B = \frac{(n-m)!}{(n-m-\sum_{j=1}^{m-1} r_j)! \prod_{j=1}^{m-1} r_j!}$, $D = \sum_{j=1}^{m-1} r_j$, $E = (m-1)(n-m) - \sum_{j=1}^{m-1} (m-j)r_j$; see Weian et al. (2011). It is clear that the MLE of θ is

$$\hat{\theta}_{ML} = \frac{m}{T(\mathbf{x})}.$$

Because of the Markovian property of progressively Type-II right censored-order statistics, it can be shown that the conditional distribution of $X_{i,(s)}$ given $\mathbf{X}$ is just the distribution of $X_{i,(s)}$ given $X_i = x_i$; see Balakrishnan and Aggarwala (2000). This implies that the density of $X_{i,(s)}$ given $\mathbf{X} = \mathbf{x}$ is the same as the density of the s -th order statistic out of r_i units from the population with density $f(y)/(1-F(x_i))$, $y \geq x_i$ (left

truncated density at x_i). Thus, the conditional density of $Y = X_{i,(s)}$ given $X_i = x_i$, for $y \geq x_i$, is given by

$$\begin{aligned} f_{Y|X_i}(y|x_i, r_i; \theta) &= s \binom{r_i}{s} f(y; \theta) [F(y; \theta) - F(x_i; \theta)]^{s-1} \\ &\quad \times [1 - F(y; \theta)]^{r_i-s} [1 - F(x_i; \theta)]^{-r_i}. \end{aligned} \quad (6)$$

From (3), we have

$$\begin{aligned} f_{Y|X_i}(y|x_i, r_i; \theta) &= s \binom{r_i}{s} \theta \frac{f_0(y)}{\bar{F}_0(y)} [(\bar{F}_0(x_i))^\theta - (\bar{F}_0(y))^\theta]^{s-1} \\ &\quad \times [(\bar{F}_0(y))^\theta]^{r_i-s+1} [(\bar{F}_0(x_i))^\theta]^{-r_i}, \quad y \geq x_i. \end{aligned} \quad (7)$$

Therefore, PLF of Y , θ and p can be extended as

$$\begin{aligned} L(y, \theta, p; \mathbf{x}, \mathbf{r}) &= f_{Y|X_i}(y|x_i, r_i; \theta) f(x_i; \mathbf{r}; \theta) f(\mathbf{r}; p) \\ &= c f(y; \theta) [F(y; \theta) - F(x_i; \theta)]^{s-1} [1 - F(y; \theta)]^{r_i-s} \\ &\quad \times \prod_{j=1}^m f(x_j; \theta) \prod_{j=1, j \neq i}^m [1 - F(x_j; \theta)]^{r_j} \\ &\quad \times p^D (1-p)^E, \quad y \geq x_i, \end{aligned}$$

where c is a constant factor. Disregarding the constant term, the predictive log-likelihood function for $y \geq x_i$ is

$$\begin{aligned} \ln L(y, \theta, p; \mathbf{x}, \mathbf{r}) &= \ln f(y; \theta) + (s-1) \ln [F(y; \theta) - F(x_i; \theta)] \\ &\quad + (r_i - s) \ln [1 - F(y; \theta)] + \sum_{j=1}^m \ln f(x_j; \theta) \\ &\quad + \sum_{j=1, j \neq i}^m r_j \ln [1 - F(x_j; \theta)] + D \ln p + E \ln (1-p). \end{aligned}$$

By substituting (2) and (3), the log PLF of $Y = y$ and θ , for $y \geq x_i$, is given by

$$\begin{aligned} \ln L(y, \theta, p; \mathbf{x}, \mathbf{r}) &= (m+1) \ln \theta + \ln \left[\frac{f_0(y)}{\bar{F}_0(y)} \right] + \sum_{j=1}^m \ln \frac{f_0(x_j)}{\bar{F}_0(x_j)} \\ &\quad + (s-1) \ln \left[1 - \left(\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right)^\theta \right] \\ &\quad + \theta(r_i - s + 1) [\ln \bar{F}_0(y) - \ln \bar{F}_0(x_i)] \\ &\quad + \theta \sum_{j=1}^m (r_j + 1) \ln \bar{F}_0(x_j) + D \ln p + E \ln (1-p). \end{aligned} \quad (8)$$

The maximum likelihood predictor (MLP) of Y and the Predictive Maximum Likelihood Estimator (PMLE) of θ are obtained directly by maximizing $\ln L(y, \theta; \mathbf{x})$, since

$f(\mathbf{r}; p)$ does not involve y and θ . Thus, the predictive likelihood equations (PLEs) for y and θ (for $y \geq x_i$) are given by

$$\begin{aligned} \frac{\partial \ln L(y, \theta; \mathbf{x})}{\partial y} &= \frac{1}{\bar{F}_0(y)} \left[\frac{f'_0(y)\bar{F}_0(y) + f_0^2(y)}{f_0(y)} - \theta(r_i - s + 1)f_0(y) \right. \\ &\quad \left. + \theta(s - 1)f_0(y) \frac{\left[\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right]^\theta}{1 - \left[\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right]^\theta} \right] = 0, \end{aligned}$$

and

$$\begin{aligned} \frac{\partial \ln L(y, \theta; \mathbf{x})}{\partial \theta} &= \frac{m + 1}{\theta} - (s - 1) \ln \left[\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right] \frac{\left[\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right]^\theta}{1 - \left[\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right]^\theta} \\ &\quad + (r_i - s + 1) [\ln \bar{F}_0(y) - \ln \bar{F}_0(x_i)] \\ &\quad + \sum_{j=1}^m (r_j + 1) \ln [\bar{F}_0(x_j)] = 0. \end{aligned}$$

Independently, the MLE of p can be obtained by maximizing $f(\mathbf{r}; p)$ directly,

$$\frac{\partial \ln f(\mathbf{r}; p)}{\partial p} = \frac{D}{p} - \frac{E}{1 - p} = 0,$$

hence,

$$\hat{p} = \frac{D}{D + E}.$$

Example 2.1. For exponential distribution, we have

$$\bar{F}_0(x) = e^{-x} \quad x > 0.$$

Hence, the PLEs reduce to:

$$\frac{\partial \ln L(y, \theta; \mathbf{x})}{\partial y} = -\theta(r_i - s + 1) + \theta(s - 1) \frac{e^{-\theta(y-x_i)}}{1 - e^{-\theta(y-x_i)}} = 0,$$

and

$$\begin{aligned} \frac{\partial \ln L(y, \theta; \mathbf{x})}{\partial \theta} &= \frac{m + 1}{\theta} - \left[\sum_{j=1}^m (r_j + 1)x_j + (r_i - s + 1)(y - x_i) \right] \\ &\quad + (s - 1) \frac{(y - x_i)e^{-\theta(y-x_i)}}{1 - e^{-\theta(y-x_i)}} = 0. \end{aligned}$$

Now, the MLP of Y and PMLE of θ can be obtained as

$$\hat{Y}_{MLP} = x_i + \frac{1}{\hat{\theta}_{PML}} \ln \left[\frac{R_i}{R_i - s + 1} \right],$$

and

$$\hat{\theta}_{PML} = \frac{m+1}{T_0(\mathbf{x})},$$

where

$$T_0(\mathbf{x}) = \sum_{i=1}^m (R_i + 1)x_i.$$

2.2 Best unbiased predictor

In this section, we first define a best unbiased predictor and then we obtain BUP for the interested variables.

Definition 2.2. A statistic $\hat{Y}$ is called a best unbiased predictor (BUP) of $Y = X_{i,(s)}$, if the predictor error $\hat{Y} - Y$ has a mean zero and its prediction error variance $\text{Var}[\hat{Y} - Y]$ is less than or equal to that of any other unbiased predictor of Y .

Nayak (2000) showed that since the conditional distribution of Y given $\mathbf{x} = (x_1, \dots, x_m)$ is just the distribution of Y given X_i , the BUP of Y is

$$\hat{Y}_{BUP} = E[Y|X_i = x_i].$$

By (7), we have

$$\begin{aligned} \hat{Y}_{BUP} &= \int_{x_i}^{\infty} y f(y|x_i; \theta) dy \\ &= \int_0^1 \bar{F}_0^{-1} \left(u^{\frac{1}{\theta}} \bar{F}_0(x_i) \right) \frac{u^{r_i-s}(1-u)^{s-1}}{\text{Beta}(r_i-s+1, s)} du. \end{aligned}$$

Note that R_i is a random variable, so we define

$$\hat{Y}_{EBUP} = E_{R_i} \left[\int_0^1 \bar{F}_0^{-1} \left(u^{\frac{1}{\theta}} \bar{F}_0(x_i) \right) \frac{u^{r_i-s}(1-u)^{s-1}}{\text{Beta}(r_i-s+1, s)} du \right].$$

Example 2.3. For the exponential distribution and after replacing θ by its MLE, we have

$$\begin{aligned} \hat{Y}_{EBUP} &= E_{R_i} \left[\int_0^1 -\ln \left(u^{\frac{T_0(\mathbf{x})}{m}} e^{-x_i} \right) \frac{u^{r_i-s}(1-u)^{s-1}}{\text{Beta}(r_i-s+1, s)} du \right] \\ &= x_i + \frac{T_0(\mathbf{x})}{m} E_{R_i} \left[\int_0^1 (-\ln u) \frac{u^{r_i-s}(1-u)^{s-1}}{\text{Beta}(r_i-s+1, s)} du \right]. \end{aligned}$$

2.3 Empirical conditional median predictor

As another possible predictor, it can be considered Conditional Median Predictor (CMP); see Raqab and Nagaraja (1995).

Definition 2.4. A predictor $\hat{Y}$ is called the CMP of Y , if it is the median of the conditional distribution of Y given $X_i = x_i$; that is,

$$P_{\theta}(Y \leq \hat{Y}|X_i = x_i) = P_{\theta}(Y \geq \hat{Y}|X_i = x_i).$$

By using the relation

$$P_\theta(Y \leq \hat{Y} | X_i = x_i, R_i = r_i) = P_\theta \left[\left(\frac{\bar{F}_0(Y)}{\bar{F}_0(X_i)} \right)^\theta \geq \left(\frac{\bar{F}_0(\hat{Y})}{\bar{F}_0(X_i)} \right)^\theta \middle| X_i = x_i, R_i = r_i \right],$$

and this fact that the distribution of $\left(\frac{\bar{F}_0(Y)}{\bar{F}_0(X_i)} \right)^\theta$ given $X_i = x_i$ and $R_i = r_i$ is a $Beta(r_i - s + 1, s)$ distribution, we obtain the CMP of Y as

$$\hat{Y}_{CMP} = \bar{F}_0^{-1} \left(\bar{F}_0(x_i) [Med(U|r_i)]^{\frac{1}{\theta}} \right),$$

where $U|r_i$ has $Beta(r_i - s + 1, s)$ distribution and $Med(U|r_i)$ stands for median of $U|r_i$. Note that R_i is a random variable, so we define Empirical Conditional Median Predictor (ECMP) as below

$$\hat{Y}_{ECMP} = E_{R_i} \left[\bar{F}_0^{-1} \left(\bar{F}_0(x_i) [Med(U|r_i)]^{\frac{1}{\theta}} \right) \right],$$

as the expectation of CMP of Y . Substituting θ with its MLE leads us to ECMP in the following form,

$$\hat{Y}_{ECMP} = E_{R_i} \left[\bar{F}_0^{-1} \left(\bar{F}_0(x_i) [Med(U|r_i)]^{\frac{T_0(\mathbf{x})}{m}} \right) \right].$$

Example 2.5. Taking $\bar{F}_0(x) = e^{-x}$, we obtain

$$\begin{aligned} \hat{Y}_{ECMP} &= E_{R_i} \left[-\ln \left(e^{-x_i} [Med(U|r_i)]^{\frac{T_0(\mathbf{x})}{m}} \right) \right] \\ &= x_i - \frac{T_0(\mathbf{x})}{m} E_{R_i} [\ln[Med(U|r_i)]] . \end{aligned}$$

2.4 Bayesian point predictors

Bayesian predictors are obtained from $f^*(y|\mathbf{x})$, the Bayes predictive density function of Y given $\mathbf{X} = \mathbf{x}$, and given loss function. The loss function $L(y, \hat{y})$ indicates the loss for using $\hat{y}$ as the predicted value of Y when the realized value is y . Here, we consider the three loss functions to discuss Bayesian predictors.

Definition 2.6. The Squared Error Loss (SEL) is a symmetric and most commonly used loss function, defined as

$$L(y, \hat{y}) = (\hat{y} - y)^2.$$

The Bayes point predictor of y , under SEL function, $(\hat{y}_{SEP})$ is $E[Y|\mathbf{x}]$.

Definition 2.7. The Absolute Difference Loss (ADL) is a symmetric loss function, defined as

$$L(y, \hat{y}) = |y - \hat{y}|.$$

The Bayes point predictor of y , under ADL function, $(\hat{y}_{ADL})$ is the median of the Bayes predictive distribution.

Definition 2.8. *The General Entropy Loss (GEL) function is a useful asymmetric loss function, that*

$$L(y, \hat{y}) \propto \left(\frac{\hat{y}}{y}\right)^q - q \ln \left(\frac{\hat{y}}{y}\right) - 1, \quad q \neq 0, \quad (9)$$

whose minimum occurs at $\hat{y} = y$. This loss function is a generalization of the Entropy-loss.

The Bayes point predictor $\hat{y}_{GEP}$ of y under GEL function becomes

$$\hat{y}_{GEP} = (E_Y(Y^{-q}))^{-\frac{1}{q}}, \quad (10)$$

where $E_Y(\cdot)$ denotes the posterior expectation with respect to the Bayes predictive density function of y ; see Soliman (2005). When the parameters p and θ are unknown, if p and θ are independent, all Bayesian predictors are the same as the fixed removal obtained by Asgharzadeh and Valiollahi (2010). Now, consider p and θ are dependent. Let θ given p be a non-informative prior distribution, that is $f(\theta|p) = 0$ for $\theta > 0$ and $p \sim \text{Beta}(a, b)$. In this cases, the posterior density function of p and θ given the data, $\pi(p, \theta|\mathbf{x}, r)$, can be obtained as below

$$\pi(p, \theta|\mathbf{x}, r) = \frac{[T(\mathbf{x})]^{m+1} \theta^m e^{-\theta T(\mathbf{x})} p^{a+D} (1-p)^{b+E-1}}{\Gamma(m+1) \text{Beta}(a+D+1, E+b)} I_{(0, \frac{1}{p})}(\theta). \quad (11)$$

The Bayes predictive density function of Y given $X_i = x_i$ is given by

$$f^*(y|x_i, r) = \int_0^\infty \int_0^1 f(y|x_i, r; \theta, p) \pi(p, \theta|\mathbf{x}, r) dp d\theta. \quad (12)$$

By substituting (7) and (11), for $y \geq x_i$ we get

$$\begin{aligned} f^*(y|x_i) &= \int_0^\infty \int_0^1 s \binom{r_i}{s} \theta \frac{f_0(y)}{\bar{F}_0(y)} [(\bar{F}_0(y))^\theta]^{r_i-s+1} \\ &\quad \times [(\bar{F}_0(x_i))^\theta - (\bar{F}_0(y))^\theta]^{s-1} [(\bar{F}_0(x_i))^\theta]^{-r_i} \\ &\quad \times \frac{[T(\mathbf{x})]^{m+1} \theta^m e^{-\theta T(\mathbf{x})} p^{a+D} (1-p)^{b+E-1}}{\Gamma(m+1) \text{Beta}(a+D+1, E+b)} dp d\theta. \end{aligned} \quad (13)$$

Using bivariate expansion, we have

$$[(\bar{F}_0(x_i))^\theta - (\bar{F}_0(y))^\theta]^{s-1} = \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j [\bar{F}_0(y)]^{\theta j} [\bar{F}_0(x_i)]^{\theta(s-j-1)}. \quad (14)$$

The equation (13) can be rewritten as

$$\begin{aligned} f^*(y|x_i) &= s(m+1) \binom{r_i}{s} [T(\mathbf{x})]^{m+1} \frac{f_0(y)}{\bar{F}_0(y)} \\ &\quad \times \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j \end{aligned}$$

$$\times \left[T(\mathbf{x}) - (r_i - s + j + 1) \ln \left(\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right) \right]^{-(m+2)}. \quad (15)$$

By using (15), under ADL function, the Bayes point predictor of Y , $\hat{Y}_{ADP}$, can be obtained as

$$\begin{aligned} \frac{1}{2} &= \int_{x_i}^{\hat{Y}_{ADP}} f^*(y|x_i, r) dy \\ &= s \binom{r_i}{s} \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j \frac{1}{r_i + j - s + 1} \\ &\quad \times \left[1 - \frac{[T(\mathbf{x})]^{m+1}}{\left[T(\mathbf{x}) - (r_i + j - s + 1) \ln \left(\frac{\bar{F}_0(\hat{Y}_{ADP})}{\bar{F}_0(x_i)} \right) \right]^{m+1}} \right]. \end{aligned} \quad (16)$$

Using (10), the Bayes point predictor of Y under GEL function is given by

$$\begin{aligned} \hat{Y}_{GEP} &= \left[\int_{x_i}^{\infty} y^{-q} f^*(y|x_i, r) dy \right]^{-\frac{1}{q}} \\ &= \left[s(m+1) \binom{r_i}{s} [T(\mathbf{x})]^{m+1} \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j I(j) \right]^{-\frac{1}{q}}, \end{aligned} \quad (17)$$

where

$$I(j) = \int_{x_i}^{\infty} y^{-q} \frac{f_0(y)}{\bar{F}_0(y)} \left[T(\mathbf{x}) - (r_i - s + j + 1) \ln \left(\frac{\bar{F}_0(y)}{\bar{F}_0(x_i)} \right) \right]^{-(m+2)} dy. \quad (18)$$

Remark 2.9. The Bayes point predictor under SEL function ($\hat{Y}_{SEP}$) can be obtained by setting $q = -1$ in the Bayes point predictor under GEL function ($\hat{Y}_{GEP}$).

Example 2.10. In the exponential distribution, the Bayes point predictors $\hat{Y}_{ADP}$ and $\hat{Y}_{GEP}$ of Y can be obtained as below

$$\begin{aligned} \frac{1}{2} &= s \binom{r_i}{s} \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j \frac{1}{r_i + j - s + 1} \\ &\quad \times \left[1 - \frac{[T_0(\mathbf{x})]^{m+1}}{[T_0(\mathbf{x}) - (r_i + j - s + 1)(x_i - \hat{Y}_{ADP})]^{m+1}} \right]. \end{aligned} \quad (19)$$

and

$$\hat{Y}_{GEP} = \left(s \binom{r_i}{s} \frac{m + \alpha}{\beta + T_0(\mathbf{x})} \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j I_0(j) \right)^{-\frac{1}{q}}, \quad (20)$$

where

$$I_0(j) = \int_{x_i}^{\infty} y^{-q} \left[1 + \frac{(r_i - s + j + 1)(y - x_i)}{\beta + T_0(\mathbf{x})} \right]^{-(m+\alpha+1)} dy.$$

3 Numerical example and simulation

In this section, we present the results of numerical experiments of the different point predictors investigated in the previous section. We present the analysis of a data set and a Monte Carlo simulation to compare the performance of the point predictors with respect to their biases and Mean Squared Prediction Errors (MSPEs). Consider the exponential distribution $E(\theta)$ with CDF

$$F(x, \theta) = 1 - e^{-\theta x}, \quad x > 0, \theta > 0,$$

as a special case from the model (2). Here, we have $\bar{F}_0(x) = e^{-x}$ and $T_0(\mathbf{x}) = \sum_{i=1}^m (r_i + 1)x_i$.

3.1 Numerical example

First, we generate the $B = 1000$ of progressive censoring sample with binomial removals $r_i (i = 1, 2, \dots, m)$ and then, the point predictors are obtained by Monte-Carlo method. Steps are as follows:

Step 1. Generate the values r_i for the given values $n = 20$, $m = 9$ and $p = 0.25$, from

$$R_i \sim B\left(n - m - \sum_{j=1}^{i-1} r_j, p\right),$$

where $r_m = n - m - \sum_{i=1}^{m-1} r_i$, $i = 1, 2, \dots, m - 1$. Note that $\sum_{j=1}^0 r_j = 0$.

Step 2. We generate $\theta = 1$ from the prior PDF (11).

Step 3. Using the sample from step 2, we generate a progressively Type-II censored sample of size $m = 9$ with binomial removals $r_i (i = 1, 2, \dots, m)$ from step 1, from the exponential distribution. The generated sample is

0.0323, 0.0462, 0.0484, 0.1224, 0.2968, 0.3772, 1.4380, 1.5063, 1.5212.

Step 4. Using this sample, the non-Bayesian point predictors MLP, EBUP and ECMP and Bayesian point predictors ADP, SEP and GEP are obtained for $X_{i,(s)} (s = 1, 2, \dots, r_i; i = 1, \dots, m)$ by using the previous step's sample. The results are presented in Table 2.

Table 2: Different Point Predictors for $n = 20$, $m = 9$ and $\theta = 1$ with random removals $\mathbf{R} = (1, 1, 3, 2, 0, 2, 0, 0, 2)$.

$X_{i,(s)}$	MLP	EBUP	ECMP	ADP	SEP	GEP($q = -0.5$)	GEP($q = 0.5$)
$X_{1,(1)}$	0.051	0.512	0.371	0.349	0.510	0.264	0.416
$X_{2,(1)}$	0.399	0.934	0.771	0.723	0.923	0.594	0.802
$X_{3,(1)}$	0.563	1.149	0.973	0.914	1.133	0.790	1.004
$X_{3,(2)}$	0.566	1.167	0.987	0.929	1.151	0.821	1.025
$X_{3,(3)}$	0.600	1.266	1.065	1.004	1.243	0.897	1.108
$X_{4,(1)}$	0.557	1.253	1.041	0.985	1.234	0.893	1.101
$X_{4,(2)}$	0.581	1.303	1.083	1.028	1.283	0.956	1.153
$X_{6,(1)}$	0.655	1.423	1.188	1.133	1.401	1.072	1.271
$X_{6,(2)}$	0.784	1.584	1.339	1.283	1.560	1.262	1.437
$X_{9,(1)}$	0.997	1.760	1.529	1.479	1.740	1.502	1.641
$X_{9,(2)}$	1.529	2.471	2.182	2.117	2.420	2.193	2.303

3.2 Simulation study

A Monte Carlo simulation study is used to evaluate the biases and MSPEs for the predictors MLP, EBUP, ECMP, ADP, SEP and GEP. We randomly generated 1000 progressively censored sample from exponential distribution with $\theta = 0.5$. Note that we worked with the binomial removal corresponding to $p = 0.25, 0.5, 0.75$. To compute Bayesian predictors, since we do not have any prior information, we can use arbitrary a and b , for example, we assume that $a = 2$ and $b = 3$ in our simulation.

Tables 3, 4 and 5 display the biases and MSPEs of different predictors obtained from this simulation study. All the computations are performed using R software.

Conclusion

From Table 3, 4 and 5, for all n , m and p , we have the best results in terms of Biases for the EBUPs and in terms of MSPEs for the SEPs. As expected, the MLP does not work well, See Asgharzadeh and Valiollahi (2010). Also, we observe that for $n = 10$ most of predictors usually under-predict (because of negative sign of bias) the life-lengths $X_{i,(s)}$ ($s = 1, \dots, r_i; i = 1, 2, \dots, m$), except the EBUP, which over-predict (because of positive sign of bias) most of the time; and for $n = 15$ most of predictors usually under-predict the life-lengths $X_{i,(s)}$ ($s = 1, \dots, r_i; i = 1, 2, \dots, m$), except the EBUP, which over-predict in special cases $m = 13$ as $p = 0.5, 0.75$ and $m = 9$ as $p = 0.5$. This table shows that the Bayes predictors relative to GEL function are sensitive to the value of the shape parameter q .

Table 3: Biases and MSPE's of point predictors for $\theta = 0.5$ and $p = 0.25$ with random removals.

n	m	$X_{i,(s)}$		MLP	EBUP	ECMP	ADP	SEP	GEP($q = -0.5$)	GEP($q = 0.5$)
10	4	$X_{1,(1)}$	Bias	-1.369	-0.026	-0.438	-0.570	-0.223	-0.721	-0.464
			MSPE	4.168	2.518	2.457	2.544	2.255	2.764	2.376
		$X_{1,(2)}$	Bias	-1.459	-0.029	-0.450	-0.662	-0.387	-0.767	-0.662
			MSPE	5.219	3.763	3.600	3.691	3.230	3.998	3.471
		$X_{4,(1)}$	Bias	-1.336	0.038	-0.361	-0.573	-0.385	-0.529	-0.643
			MSPE	5.181	4.092	3.846	3.849	3.332	5.060	3.615
		$X_{4,(2)}$	Bias	-1.251	-0.058	-0.404	-0.586	-0.380	-0.494	-0.586
			MSPE	4.157	2.864	2.790	2.899	2.623	3.196	2.917
		$X_{4,(3)}$	Bias	-1.295	-0.120	-0.451	-0.652	-0.436	-0.474	-0.646
			MSPE	4.083	2.876	2.851	2.973	2.600	3.406	2.830
		$X_{4,(4)}$	Bias	-1.991	-0.088	-0.615	-0.992	-1.341	0.070	-1.822
			MSPE	9.548	7.675	7.325	7.409	7.306	23.678	9.121
	6	$X_{2,(1)}$	Bias	-1.599	0.057	-0.451	-0.565	-0.171	-0.761	-0.459
			MSPE	5.766	3.360	3.352	3.441	3.116	3.747	3.268
		$X_{2,(2)}$	Bias	-1.660	0.017	-0.493	-0.645	-0.351	-0.704	-0.643
			MSPE	6.643	4.467	4.380	4.475	4.023	4.782	4.279
		$X_{6,(1)}$	Bias	-1.604	0.009	-0.480	-0.625	-0.382	-0.504	-0.645
			MSPE	6.299	4.152	4.119	4.208	3.805	4.475	4.092
		$X_{6,(2)}$	Bias	-1.796	0.166	-0.429	-0.613	-0.786	0.360	-1.221
			MSPE	7.377	5.075	4.830	4.902	5.137	12.220	6.482
10	8	$X_{2,(1)}$	Bias	-1.822	-0.067	-0.606	-0.698	-0.279	-0.854	-0.569
			MSPE	7.159	4.002	4.214	4.317	3.893	4.620	4.11
		$X_{5,(1)}$	Bias	-2.096	-0.208	-0.788	-0.903	-0.734	-0.501	-1.071
			MSPE	9.206	5.278	5.636	5.791	5.672	8.465	6.52
	9	$X_{1,(1)}$	Bias	-1.408	0.004	-0.429	-0.495	-0.114	-0.737	-0.373
			MSPE	4.588	2.490	2.586	2.644	2.396	2.972	2.508
		$X_{4,(1)}$	Bias	-1.426	0.036	-0.414	-0.520	-0.203	-0.739	-0.486
			MSPE	5.116	3.302	3.293	3.354	3.032	3.631	3.198
		$X_{5,(1)}$	Bias	-1.544	0.017	-0.460	-0.572	-0.264	-0.709	-0.551
			MSPE	5.650	3.435	3.487	3.574	3.279	3.804	3.521
		$X_{5,(2)}$	Bias	-1.713	-0.030	-0.546	-0.658	-0.355	-0.701	-0.643
			MSPE	6.725	4.050	4.160	4.256	3.860	4.475	4.146
		$X_{8,(1)}$	Bias	-1.587	0.118	-0.405	-0.507	-0.225	-0.417	-0.494
			MSPE	5.594	3.488	3.415	3.469	3.124	3.657	3.290
		$X_{9,(1)}$	Bias	-1.899	0.043	-0.554	-0.668	-0.716	0.164	-1.109
			MSPE	7.840	4.746	4.807	4.903	5.203	11.135	6.399
15	11	$X_{1,(1)}$	Bias	-1.617	-0.035	-0.521	-0.582	-0.165	-0.855	-0.451
			MSPE	5.777	3.177	3.337	3.398	3.078	3.801	3.225
		$X_{1,(2)}$	Bias	-1.685	-0.036	-0.546	-0.635	-0.271	-0.844	-0.569
			MSPE	6.706	4.040	4.197	4.283	3.895	4.635	4.124
		$X_{4,(1)}$	bias	-1.769	0.018	-0.532	-0.622	-0.276	-0.693	-0.570
			MSPE	6.937	4.153	4.247	4.330	3.972	4.567	4.193
		$X_{5,(1)}$	Bias	-1.972	-0.016	-0.617	-0.704	-0.534	-0.310	-0.873
			MSPE	8.024	4.454	4.666	4.761	4.601	6.759	5.271
	13	$X_{1,(1)}$	Bias	-1.768	0.098	-0.474	-0.537	-0.084	-0.810	-0.404
			MSPE	6.718	3.779	3.857	3.909	3.613	4.267	3.729
		$X_{2,(1)}$	Bias	-1.885	0.082	-0.525	-0.598	-0.270	-0.538	-0.590
			MSPE	7.197	3.998	4.092	4.154	3.880	5.842	4.169

Table 4: Biases and MSPE's of point predictors for $\theta = 0.5$ and $p = 0.5$ with random removals.

n	m	$X_{i,(s)}$		MLP	EBUP	ECMP	ADP	SEP	GEP($q = -0.5$)	GEP($q = 0.5$)
10	4	$X_{1,(1)}$	Bias	-0.788	0.015	-0.231	-0.310	-0.052	-0.382	-0.200
			MSPE	1.576	0.946	0.904	0.935	0.890	0.988	0.910
		$X_{1,(2)}$	Bias	-1.064	0.003	-0.298	-0.492	-0.217	-0.627	-0.447
			MSPE	3.316	2.552	2.429	2.486	2.289	2.654	2.417
		$X_{1,(3)}$	Bias	-1.237	0.088	-0.271	-0.544	-0.375	-0.607	-0.669
			MSPE	4.796	4.410	4.024	3.956	3.497	4.361	3.752
		$X_{2,(1)}$	Bias	-1.515	-0.021	-0.434	-0.714	-0.592	-0.669	-0.907
			MSPE	6.254	5.235	4.918	4.910	4.154	5.965	4.600
	6	$X_{4,(1)}$	Bias	-1.644	-0.062	-0.516	-0.768	-0.623	-0.636	-0.926
			MSPE	7.211	5.678	5.465	5.496	4.733	6.450	5.194
		$X_{4,(2)}$	Bias	-1.987	-0.031	-0.610	-0.879	-0.915	-0.274	-1.305
			MSPE	8.847	6.438	6.172	6.293	5.889	10.393	6.915
		$X_{1,(1)}$	Bias	-1.148	-0.036	-0.377	-0.453	-0.125	-0.604	-0.328
			MSPE	3.313	1.909	1.963	2.020	1.842	2.242	1.931
		$X_{1,(2)}$	Bias	-1.431	-0.068	-0.475	-0.637	-0.351	-0.822	-0.631
			MSPE	5.237	3.787	3.715	3.805	3.359	4.158	3.576
		$X_{1,(3)}$	Bias	-1.698	-0.083	-0.563	-0.749	-0.535	-0.794	-0.852
			MSPE	7.459	5.181	5.182	5.309	4.772	6.322	5.236
		$X_{2,(1)}$	Bias	-1.963	-0.041	-0.623	-0.797	-0.608	-0.601	-0.952
			MSPE	8.032	5.063	5.049	5.188	4.650	6.346	5.187
10	8	$X_{2,(1)}$	Bias	-1.700	-0.023	-0.538	-0.626	-0.211	-0.866	-0.507
			MSPE	6.055	3.159	3.322	3.415	3.060	3.801	3.256
		$X_{8,(1)}$	Bias	-1.973	-0.047	-0.639	-0.769	-0.454	-0.854	-0.791
			MSPE	8.524	5.410	5.473	5.583	5.012	5.988	5.338
	9	$X_{1,(1)}$	Bias	-0.796	0.016	-0.233	-0.271	-0.009	-0.396	-0.161
			MSPE	1.632	0.936	0.951	0.969	0.906	1.078	0.919
		$X_{1,(2)}$	Bias	-1.019	0.018	-0.297	-0.393	-0.120	-0.609	-0.345
			MSPE	3.261	2.328	2.306	2.351	2.198	2.583	2.293
		$X_{1,(3)}$	Bias	-1.206	0.009	-0.354	-0.483	-0.270	-0.652	-0.542
			MSPE	4.558	3.419	3.362	3.420	3.108	3.803	3.339
		$X_{3,(1)}$	Bias	-1.370	0.033	-0.388	-0.521	-0.305	-0.642	-0.602
			MSPE	5.641	4.179	4.108	4.162	3.745	4.581	4.002
		$X_{3,(2)}$	Bias	-1.561	0.052	-0.438	-0.563	-0.307	-0.650	-0.615
			MSPE	5.967	3.928	3.893	3.963	3.615	4.388	3.900
		$X_{3,(3)}$	Bias	-1.926	0.001	-0.590	-0.711	-0.453	-0.619	-0.787
			MSPE	8.194	5.081	5.155	5.265	4.901	7.920	5.341
	11	$X_{1,(1)}$	Bias	-1.079	0.086	-0.271	-0.317	0.019	-0.521	-0.197
			MSPE	2.779	1.543	1.534	1.558	1.470	1.752	1.493
		$X_{1,(2)}$	Bias	-1.301	0.056	-0.365	-0.459	-0.145	-0.725	-0.425
			MSPE	4.604	3.140	3.115	3.167	2.918	3.482	3.051
		$X_{1,(3)}$	Bias	-1.591	0.005	-0.486	-0.592	-0.297	-0.797	-0.612
			MSPE	6.504	4.251	4.312	4.395	3.989	4.750	4.249
		$X_{3,(1)}$	Bias	-1.779	0.150	-0.444	-0.544	-0.206	-0.647	-0.538
			MSPE	6.809	3.940	3.956	4.028	3.724	4.229	3.952
15	13	$X_{1,(1)}$	Bias	-1.63	0.078	-0.446	-0.503	-0.073	-0.812	-0.382
			MSPE	5.620	3.048	3.110	3.156	2.891	3.620	3.012
		$X_{1,(2)}$	Bias	-1.96	-0.018	-0.620	-0.704	-0.312	-0.944	-0.655
			MSPE	8.041	4.635	4.816	4.910	4.438	5.387	4.716

Table 5: Biases and MSPE's of point predictors for $\theta = 0.5$ and $p = 0.75$ with random removals.

n	m	$X_{i,(s)}$		MLP	EBUP	ECMP	ADP	SEP	GEP($q = -0.5$)	GEP($q = 0.5$)
10	4	$X_{1,(1)}$	Bias	-0.469	0.012	-0.135	-0.183	0.004	-0.205	-0.083
			MSPE	0.543	0.364	0.350	0.358	0.355	0.376	0.347
		$X_{1,(2)}$	Bias	-0.647	0.042	-0.147	-0.281	-0.031	-0.370	-0.184
			MSPE	1.227	1.104	0.995	0.984	0.944	1.055	0.923
		$X_{1,(3)}$	Bias	-0.942	0.021	-0.221	-0.459	-0.234	-0.550	-0.463
			MSPE	3.049	2.781	2.590	2.570	2.331	2.886	2.447
		$X_{1,(4)}$	Bias	-1.261	-0.002	-0.313	-0.635	-0.543	-0.614	-0.852
			MSPE	5.056	4.759	4.418	4.330	3.700	5.726	4.107
		$X_{1,(5)}$	Bias	-1.510	0.044	-0.361	-0.705	-0.724	-0.448	-1.087
			MSPE	6.635	6.228	5.735	5.556	4.709	14.698	5.378
		$X_{1,(6)}$	Bias	-2.062	-0.096	-0.657	-0.975	-1.031	-0.516	-1.441
			MSPE	10.057	7.798	7.495	7.530	6.895	14.985	8.009
	6	$X_{1,(1)}$	Bias	-0.768	-0.009	-0.242	-0.294	-0.033	-0.371	-0.171
			MSPE	1.485	0.862	0.884	0.910	0.847	0.974	0.863
		$X_{1,(2)}$	Bias	-1.012	0.044	-0.266	-0.408	-0.127	-0.584	-0.356
			MSPE	2.689	1.917	1.831	1.878	1.735	2.046	1.819
		$X_{1,(3)}$	Bias	-1.413	0.014	-0.395	-0.612	-0.424	-0.738	-0.745
			MSPE	5.739	4.453	4.298	4.387	3.959	4.903	4.326
		$X_{2,(1)}$	Bias	-1.849	0.065	-0.500	-0.726	-0.591	-0.718	-0.981
			MSPE	8.170	5.906	5.701	5.784	5.141	6.626	5.688
	8	$X_{1,(1)}$	Bias	-1.430	0.001	-0.438	-0.514	-0.123	-0.729	-0.382
			MSPE	4.474	2.371	2.468	2.533	2.307	2.847	2.429
		$X_{2,(1)}$	Bias	-1.956	-0.080	-0.656	-0.808	-0.486	-1.049	-0.848
			MSPE	8.712	5.673	5.769	5.913	5.312	6.397	5.691
15	9	$X_{1,(1)}$	Bias	-0.492	-0.008	-0.156	-0.179	-0.009	-0.232	-0.096
			MSPE	0.538	0.302	0.312	0.319	0.299	0.340	0.300
		$X_{1,(2)}$	Bias	-0.654	-0.023	-0.211	-0.275	-0.047	-0.412	-0.186
			MSPE	1.336	0.930	0.944	0.968	0.934	1.059	0.955
		$X_{1,(3)}$	Bias	-0.846	-0.004	-0.246	-0.360	-0.121	-0.530	-0.321
			MSPE	2.606	2.030	2.007	2.046	1.920	2.219	2.005
		$X_{1,(4)}$	Bias	-1.186	-0.059	-0.381	-0.540	-0.380	-0.668	-0.659
			MSPE	4.908	3.832	3.812	3.896	3.710	4.271	4.045
		$X_{2,(1)}$	Bias	-1.433	-0.009	-0.426	-0.592	-0.500	-0.637	-0.846
			MSPE	6.109	4.578	4.529	4.594	4.214	5.286	4.715
		$X_{2,(2)}$	Bias	-2.011	-0.126	-0.697	-0.845	-0.639	-0.891	-1.017
			MSPE	9.541	6.208	6.388	6.514	5.876	7.866	6.482
	11	$X_{1,(1)}$	Bias	-0.748	-0.003	-0.231	-0.261	-0.017	-0.372	-0.155
			MSPE	1.364	0.719	0.761	0.776	0.710	0.870	0.732
		$X_{1,(2)}$	Bias	-1.023	-0.027	-0.333	-0.414	-0.123	-0.646	-0.341
			MSPE	3.037	2.088	2.115	2.159	1.999	2.404	2.080
		$X_{2,(1)}$	Bias	-1.431	-0.094	-0.501	-0.623	-0.391	-0.839	-0.694
			MSPE	6.147	4.499	4.551	4.636	4.216	5.109	4.517
		$X_{2,(2)}$	Bias	-1.930	-0.103	-0.663	-0.788	-0.544	-0.958	-0.917
			MSPE	8.626	5.388	5.600	5.720	5.199	6.296	5.705
	13	$X_{1,(1)}$	Bias	-1.322	0.089	-0.344	-0.391	0.008	-0.658	-0.255
			MSPE	3.769	2.022	2.053	2.084	1.935	2.389	1.981
		$X_{1,(2)}$	Bias	-1.808	0.035	-0.542	-0.639	-0.277	-0.982	-0.642
			MSPE	7.597	4.734	4.844	4.926	4.499	5.454	4.768

References

- Asgharzadeh, A. and Valiollahi, R. (2009). Inference for the proportional hazards family under progressive Type-II censoring. *JIRSS*, **1-2**, 35–53.
- Asgharzadeh, A. and Valiollahi, R. (2010). Point prediction for the proportional hazards family under progressively Type-II censoring. *JIRSS*, **2**, 127–148.
- Ahmadi, J., Jafari Jozani, M., Marchand, E., and Parsian, A. (2009). Prediction of k -records from a general class of distributions under balanced type loss functions. *Metrika*, **70**, 19–33.
- Ahmadi, J., Jafari Jozani, M., Marchand, E., and Parsian, A. (2009). Bayesian estimation based on k -record data from a general class of distributions under balanced type loss functions. *Journal of Statistical Planning Inference*, **139**, 1180–1189.
- Balakrishnan, N. and Aggarwala, R. (2000). *Progressive Censoring: Theory, Methods and Applications*. Boston, Birkhauser.
- Balakrishnan, N. and Lin C.T. (2002). Exact linear inference and prediction for exponential distributions based on general progressively Type-II censored samples. *Journal of Statistics Computation and Simulation*, **72**, 677–686.
- Basak, I., Basak, P., and Balakrishnan, N. (2006). On some predictors of times to failure of censored items in progressively censored samples. *Computational Statistics and Data Analysis*, **50**, 1313–1337.
- Cox, D.R. (1972). Regression models and life tables (with discussion). *Journal of the Royal Statistical Society, Series B*, **34**, 187–220.
- Kaminsky, K.S and Rhodin, L.S. (1985). Maximum likelihood prediction. *Annals of the Institute of Statistical Mathematics*, **37**, 507–517.
- Lawless, J. F. (2003). *Statistical Models and Methods for Life Time Data*. New York, John Wiley and Sons.
- Nayak, T.K. (2000). On best unbiased prediction and its relationships to unbiased estimation. *Journal of Statistical Planning Inference*, **84**, 171–189.
- Ng, H.K.T., Kundu, D. and Chan, P.S. (2009). Statistical analysis of exponential lifetimes under an adaptive Type-II progressive censoring scheme. *Naval Research Logistics*, **56**(8), 687–698.
- Raqab, M.Z., Asgharzadeh, A., and Valiollahi, R. (2010). Prediction for Pareto distribution based on progressively Type-II censored samples. *Computational Statistics and Data Analysis*, **54**, 1732–1743.
- Raqab, M.Z. and Nagaraja, H.N. (1995). On some predictors of future order statistics. *Metron*, **53**(12), 185–204.
- Soliman, A.A. (2005). Estimation of parameters of life from progressively censored data using Burr-XII model. *IEEE Transactions on Reliability*, **54**(1), 34–42.

- Tse, S.K., Yang, C.Y. and Yuen, H.K. (2000). Statistical analysis of Weibull distributed lifetime data under Type-II progressive censoring with binomial removals. *Journal of Applied Statistics*, **27**(8), 1033–1043.
- Viveros, R. and Balakrishnan, N. (1994). Interval estimation of life characteristics from progressively censored data. *Technometrics*, **36**, 84–91.
- Weian, Y., Yimin, S., Baowei, S. and Zhaoyong, M. (2011). Statistical analysis of generalized exponential distribution under progressive censoring with binomial removals. *Journal of Systems Engineering and Electronics*, **22**(4), 707–714.
- Wu, S.J. and Chang, C.T. (2003). censoring with random removals. *Journal of Applied Statistics*, **30**(2), 163–172.
- Wu, C.C., Wu, S.F. and Chan, H.Y. (2006). MLE and the estimated expected test time for the two-parameter Gompertz distribution under progressive censoring with binomial removals. *Applied Mathematics and Computation*, **181**(1), 1657–1670.
- Wu, S.J., Chen, Y.J. and Chang, C.T. (2007). Statistical inference based on progressively censored samples with random removals from the Burr type XII distribution. *Journal of Statistical Computation and Simulation*, **77**(1), 19–27.
- Xiang, L. M. and Tse, S. K. (2005). Maximum likelihood estimation in survival studies under progressive interval censoring with random removals. *Journal of Biopharmaceutical Statistics*, **15**(6), 981–991.
- Yuen, H.K. and Tse, S.K. (1996). Parameters estimation for Weibull distributed lifetimes under progressive censoring with random removals. *Journal of Statistical Computation and Simulation*, **55**(1), 57–71.
- Zeinab, H.A. (2008). Bayesian inference for the Pareto lifetime model under progressive censoring with binomial removals. *Journal of Applied Statistics*, **35**(11), 1203–1217.