

Using Bayesian and classical methods to estimation parameter for inverse Pareto distribution under randomly censored data

Mohammed Abdul Ridha Al-Bojmal [†]

Ramin Imany-Nabiyi ^{*}

Department of Statistics

Faculty of Mathematical Sciences

University of Tabriz

Tabriz

Iran

Abstract

In this article, we will use randomly censored data and to discuss both Bayesian and classical methods to the estimating parameters the inverse Pareto distribution (IPD). We derive the Bayes and maximum likelihood methods to obtain the parameters estimators. Asymptotic confidence intervals for the parameters are computed by using the observed Fisher information matrix. Using gamma informative priors, Bayes estimators for the parameters are derived under the (LINEX) loss function. For Bayesian estimates, the method of Markov Chain Monte Carlo (MCMC) is employed. In addition, using (MCMC) techniques, credible intervals with the largest posterior density for the parameters are generated. A simulation study compares how well each estimators performs. Lastly, two real datasets are taken into consideration for demonstration.

Subject Classification: 11J71, 11K06.

Keywords: Bayesian estimation, Random censoring, Markov Chain Monte Carlo Techniques, Asymptotic confidence interval, (HPD) credible interval, Maximum likelihood estimation.

1. Introduction

In survival analysis, the entire lifetime of a person, a plant, a machine or an animal is not always observable. Some lifetimes may be censored, in that case only a part of the lifetime is recorded. Therefore, censoring is a necessary part

This article has been corrected with minor changes. These changes do not impact the academic content of the article.

[†] E-mail: mohammed.abd939@gmail.com

^{*} E-mail: imany@tabrizu.ac.ir (Corresponding Author)

for life testing experiment. The unites in these experiements are removed or lost, resulting in incomplete information. There are different types of censoring schemes available in the literature [1-4]. Type (I), and type (II) censoring schemes are most popular censoring schemes. In these censoring schemes, either the censoring time or the number of censored items is prefixed. These censoring schemes with different lifetime models have been studied extensively by many authors like Mann *et a land* Sinhia [26], etc. Recently, progressive censoring and hybrid censoring schemes have also been studied in detail in the literature [Krishne and Kumir [18], Krishne and Kumir [13], Kumir *et al.* [17], Chaturvedi *et al.* [15] etc. The random censoring is a common occurrence in life testing experiments. For instance, let's say we are studying patients with leukemia who have received treatment. Our goal is to track their lifetimes, but various forms of censoring can occur. This includes situations like a patient being lost to follow-up (as they may choose to move away), dropping out of the study due to negative side effects or not completing the treatment, death from unrelated diseases, or the study ending before complete failure for an item in reliability engineering, such as if it breaks or to save time and money, several authors investigated the usefulness of random censoring in statistical inference for different lifetime models like Ghitany and Al-Awadhe [7], Kumir and Gerg [16], Krishna *et al.* [14], Garig *et al.* [6], Krishne and Guel [12], Kumir [15], [18], Gerg *et al.* [5] etc. Mathematically random censoring can be described as follows: suppose failure times $x_1, x_2, \dots, x_n$ are independent and identically distributed (iid) random variables with probability density function (pdf) $f_x, x > 0$ and survival function $S_x, x > 0$. Associated with these failure times $T_1, T_2, \dots, T_n$ are (iid) censoring times with (pdf) $f_T, t > 0$ and survival function $S_T, t > 0$. Further, let X_i 's and T_i 's be mutually independent all $i = 1, 2, \dots, n$. We observe failure or censored time $Y_i = \min(X_i', T_i')$; ... = $1, 2, \dots, n$ and the corresponding censor indicators [19-28]

$$D_i = \begin{cases} 1 & \text{if failure occurs} \\ 0 & \text{if censoring occurs} \end{cases}$$

This censouring scheme includes entire sample instances as special cases when $T_i = \infty \forall, i = 1, 2, \dots, n$ and Type(I) censoring when $T_i = t_0 \forall i = 1, 2, \dots, n$, where is t_0 is the prefixed studying period. Thus; the joint (pdf) of Y and Q is give [8-10]:

$$f_{Y,Q}(y, q) = \{f_x(y)S_T(y)\}^q \{f_T(y)S_x(y)\}^{1-q} \quad y > 0, q = 0, 1 \quad (1)$$

The marginal distribution of Y and Q can be obtained as:

$f_Y(y) = f_x(y)S_T(y) + f_T(y)S_x(y) \quad y > 0$ and $P[Q = q] = P^q(1 - P)^{1-q}$; $q = 0$ respectively, the probability of witnessing a failure, denoted as (p) , can be calculated using the following equation:

This paper aims to develop the Bayesian and classical methods of parameter estimation for the inverse Pareto distribution(IPD), using data that has been randomly censored.

2 The Model

If the random variable (r.v) (X) follows the inverse Pareto distribution (IPD). with the parameter (α), the pdf of (IPD) is denoted by (IPD) (α) and is given by:

$$f_X(y, \alpha) = \frac{\alpha x^{\alpha-1}}{(1+x)^{\alpha+1}} \quad \alpha > 0, x > 0 \quad (2)$$

$$f_X(y, \alpha) = \left(\frac{x}{1+x}\right)^\alpha, \quad \alpha > 0, x > 0 \quad (3)$$

$$S(x, \alpha) = P(X > x) = 1 - \left(\frac{x}{1+x}\right)^\alpha, \quad \alpha > 0, x > 0 \quad (4)$$

and

$$h(x, \alpha) = \frac{\alpha x^{\alpha-1}}{(1+x)^{\alpha+1} \left[1 - \left(\frac{x}{1+x}\right)^\alpha\right]}, \quad \alpha > 0, x > 0 \quad (5)$$

Now, suppose that the failure time is $[X]$ and censoring time is $[T]$ following the IPD (α) and IPD (α), the (pdf) of randomly censored (IPD) can be write

$$f_{Y,Q}(y, q, \alpha, \beta) = \frac{\alpha^q \beta^{1-q} y^{q(\alpha-\beta)+\beta-1}}{(1+y)^{q(\alpha+\beta)+\beta-1}} \left[1 - \left(\frac{y}{1+y}\right)^\beta\right]^q \times \left[1 - \left(\frac{y}{1+y}\right)^\alpha\right]^{1-q};$$

$$y > 0, \alpha > 0, x > 0, q = 0, 1 \quad (6)$$

and it is found that the possibility of experiencing a failure is $P = \int_0^\infty S_T(y) f_X(y) dy = \frac{\beta}{\alpha+\beta}$

3. Estimator Method

3.1 Maximum Likelihood Estimation (MLEs)

In this section, we derive the Maximum Likelihood estimators $\hat{\alpha}$ and $\hat{\beta}$ of α and β , respectively. For the observed sample $(y, q) = (y_1, q_1), (y_2, q_2), \dots, (y_n, q_n)$ of size n . Asymptotic confidence intervals (ACI) for the parameters are also obtained using the observed Fisher information matrix. The likelihood function can be written as:

$$L(y, q, \alpha, \beta) = \prod_{i=1}^n \frac{\alpha^{q_i} \beta^{1-q_i} y_i^{q_i(\alpha-\beta)+\beta-1}}{(1+y_i)^{q_i(\alpha+\beta)+\beta-1}} \times \left[1 - \left(\frac{y_i}{1+y_i}\right)^\beta\right]^{q_i} \left[1 - \left(\frac{y_i}{1+y_i}\right)^\alpha\right]^{1-q_i} \quad (7)$$

Thus, the Log likelihood function becomes:

$$\ln(\alpha, \beta | \text{data}) = m \ln \alpha + (n - m) \ln \beta + (\alpha - \beta) \times \sum_{i=1}^n q_i \ln y_i + (\beta - 1) \sum_{i=1}^n \ln y_i - (\alpha - \beta) \sum_{i=1}^n q_i \ln \left[1 - \left(\frac{y_i}{1+y_i}\right)^\beta\right] + \sum_{i=1}^n (1 - q_i) \ln \left[1 - \left(\frac{y_i}{1+y_i}\right)^\alpha\right] \quad (8)$$

Where, $m = \sum_{i=1}^n q_i$. The following normal equations result from partly differentiating the log likelihood function with regard to the parameters α and β , respectively.

$$\frac{\partial \ln L(\alpha, \beta | data)}{\partial \alpha} = \frac{m}{\alpha} + \sum_{i=1}^n q_i \ln y_i - \sum_{i=1}^n q_i \ln(1 + y_i) - \sum_{i=1}^n (1 - q_i) \frac{\left(\frac{y_i}{1+y_i}\right)^\alpha \ln\left(\frac{y_i}{1+y_i}\right)}{\left[1 + \left(\frac{y_i}{1+y_i}\right)^\alpha\right]} = 0 \quad (9)$$

$$\frac{\partial \ln L(\alpha, \beta | data)}{\partial \beta} = \frac{n-m}{\beta} + \sum_{i=1}^n \ln y_i - \sum_{i=1}^n q_i \ln y_i + \sum_{i=1}^n q_i \ln(1 + y_i) - \sum_{i=1}^n \ln(1 + y_i) - \sum_{i=1}^n q_i \frac{\left(\frac{y_i}{1+y_i}\right)^\beta \ln\left(\frac{y_i}{1+y_i}\right)}{\left[1 - \left(\frac{y_i}{1+y_i}\right)^\beta\right]} = 0 \quad (10)$$

The (ML) estimates $\hat{\alpha}$ and $\hat{\beta}$ of the parameter (α) and (β), respectively; are the solution of the non-line Equations (9) and (10). Equations (9) and (10) do not have closed form solutions, any iterative method can be used to solve these equations for α and β , respectively.

3.1.1. Asymptotic Confidence Intervals

It is difficult to drive the precise distributions of the maximum likelihood estimators since they are not in closed form. Therefore, we derive asymptotic confidence intervals for the parameters based on the observed Fisher information matrix using the asymptotic distribution of (MLEs). Let $\hat{\theta} = (\hat{\alpha}, \hat{\beta})$, be the MLE of $\theta = (\alpha, \beta)$. The observed Fisher information matrices are presented as follows:

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

Where

$$\frac{\partial^2 \ln L(\alpha, \beta)}{\partial \alpha^2} = -\frac{m}{\alpha^2} - \sum_{i=1}^n (1 - q_i) \frac{\left(\frac{y_i}{1+y_i}\right)^\alpha \left(\ln\left(\frac{y_i}{1+y_i}\right)\right)^2}{\left[1 + \left(\frac{y_i}{1+y_i}\right)^\alpha\right]^2}, \quad \frac{\partial^2 \ln L(\alpha, \beta)}{\partial \beta^2} = -\frac{n-m}{\beta^2} - \sum_{i=1}^n q_i \frac{\left(\frac{y_i}{1+y_i}\right)^\beta \left(\ln\left(\frac{y_i}{1+y_i}\right)\right)^2}{\left[1 - \left(\frac{y_i}{1+y_i}\right)^\beta\right]^2}, \quad \frac{\partial^2 \ln L(\alpha, \beta)}{\partial \alpha \partial \beta} = \frac{\partial^2 \ln L(\alpha, \beta)}{\partial \beta \partial \alpha} = 0$$

The variance-covariance matrix which is observed thus becomes $I^{-1}(\hat{\theta})$, The MLEs $\hat{\theta}$ have an asymptotic distribution that is a bivariate normal distribution, as $\hat{\theta} \sim N(\theta, I^{-1}(\hat{\theta}))$, [Lawless] [19] Consequently, two sided equal tailed, $100(1 - \xi)\%$ asymptotic confidence intervals for the parameter α and β are given by $\left(\hat{\alpha} \pm z_{\frac{\xi}{2}} \sqrt{\hat{Var}(\hat{\alpha})}\right)$ and $\left(\hat{\beta} \pm z_{\frac{\xi}{2}} \sqrt{\hat{Var}(\hat{\beta})}\right)$; Respectively. Here, $\hat{Var}(\hat{\alpha})$ and $\hat{Var}(\hat{\beta})$ are diagonal elements of the observed variance covariance

matrix $I^{-1}(\hat{\theta})$ and $Z_{\frac{\xi}{2}}$ is the upper $\left(Z_{\frac{\xi}{2}}\right)^{th}$ percentile of the standard normal distribution, the coverage probability (CPs) for the parameters are given by:

$$CP_{\alpha} = \left[\left| \frac{\hat{\alpha} - \alpha}{\sqrt{\hat{V}ar(\hat{\alpha})}} \right| \leq Z_{\frac{\xi}{2}} \right] \text{ And } CP_{\beta} = \left[\left| \frac{\hat{\beta} - \beta}{\sqrt{\hat{V}ar(\hat{\beta})}} \right| \leq Z_{\frac{\xi}{2}} \right]$$

3.2 Bayesian Estimation

This section discusses the Bayes estimators under the linear exponential (LINEX) loss function and using the (MCMC) method for of the model's (6) unknown parameters. The optimal choice in decision theory must be selected with a sufficient loss function provided. Usually, this is the function that is used for squared error loss (SELF). when the outcomes of equally large overestimations and underestimates amount by have the same impacts, The asymmetric loss functions are used to explain the effects of different errors, and the use of the (SELF) is fully warranted. Varian [27] first introduced the (LINEX) loss function. It is defined as this loss function:

$$L(\phi, \hat{\phi}) = e^{k(\phi, \hat{\phi})} - k(\phi, \hat{\phi}) - 1, \quad (11)$$

Where $\phi, k \neq 0$ is the known value and $(\hat{\phi})$ is an estimate of (ϕ) . The direction and degree of asymmetry are reflected in the sign and magnitude of the parameter (k) , respectively; When (k) is positive the situation is reversed and overestimation is more dangerous than underestimation, when (k) is negative. If $|k|$ tends to zero, Since the (SELF) loss function tends to be (LINEX), it is almost symmetric [Zellner][28]. Under the (LINEX) loss function, the Bayes estimate of the (ϕ) is given by:

$$\hat{\phi}_{Bayes} = -\frac{1}{k} \ln E[e^{-k\phi} | data]$$

Where, $E[e^{-k\phi} | data]$ is posterior expectation which exists and is finite? Now, we take into consideration the piecewise independent gamma priors for the parameters α and β as follows:

$$g_1(\alpha) = \frac{b_1^{a_1}}{\Gamma(a_1)} \alpha^{a_1-1} e^{-b_1 \alpha} \quad a_1, \alpha, b_1 > 0, \quad g_2(\beta) = \frac{b_2^{a_2}}{\Gamma(a_2)} \beta^{a_2-1} e^{-b_2 \beta} \quad a_2, \beta, b_2 > 0$$

Respectively. In this manner, the joint prior distribution of α and β can may be expressed as follows:

$$g(\alpha, \beta) \propto \alpha^{a_1-1} e^{-b_1 \alpha} \beta^{a_2-1} e^{-b_2 \beta} \quad a_1, b_1, a_2, b_2 > 0 \quad (12)$$

When the hyper parameters $a_1 = b_1 = a_2 = b_2 = 0$, the non-informative priors have been shown to be specific examples of independent gamma priors. The joint posterior distribution of α and β is supplied by the likelihood function, which is based on the observed randomly censored data. in (7), and joint prior distribution of (α, β) in (12), is given by:

$$\pi \ln L(\alpha, \beta | data) = \frac{L(data | \alpha, \beta) g(\alpha, \beta)}{\int_0^{\infty} \int_0^{\infty} L(data | \alpha, \beta) g(\alpha, \beta) d\alpha d\beta}$$

$$\pi(\alpha, \beta | data) \propto \alpha^{m+a_1-1} e^{-\alpha \left[b_1 - \sum_{i=1}^n q_1 \ln \left(\frac{y_i}{1+y_i} \right) \right]} \prod_{i=1}^n \left[1 - \left(\frac{y_i}{1+y_i} \right)^\alpha \right]^{1-q_i} \times \\ \beta^{n-m+a_1-1} e^{-\beta \left[b_2 - \sum_{i=1}^n (1-q_1) \ln \left(\frac{y_i}{1+y_i} \right) \right]} \prod_{i=1}^n \left[1 - \left(\frac{y_i}{1+y_i} \right)^\beta \right]^{1-q_i} \quad (13)$$

We can see that, according to the combined posterior distribution of α and β given in Eq.(13), the joint posterior distribution of the (α) and (β) are independents . The marginal marginal posterior distribution of the given data (y, q) is calculated as follows:

$$\pi_1(\alpha, \beta | data) \propto \alpha^{m+a_1-1} e^{-\alpha \left[b_1 - \sum_{i=1}^n q_1 \ln \left(\frac{y_i}{1+y_i} \right) \right]} \prod_{i=1}^n \left[1 - \left(\frac{y_i}{1+y_i} \right)^\alpha \right]^{1-q_i}, \alpha > 0$$

Similar to this, the marginal posterior distribution of the (β) supplied data (y, q) may be calculated as follows:

$$\pi_2(\alpha, \beta | data) \propto \beta^{n-m+a_1-1} e^{-\beta \left[b_2 - \sum_{i=1}^n (1-q_1) \ln \left(\frac{y_i}{1+y_i} \right) \right]} \prod_{i=1}^n \left[1 - \left(\frac{y_i}{1+y_i} \right)^\beta \right]^{q_i} B > 0 \quad (15)$$

Thus, the expectations of any function of α say $\phi_1(\alpha)$ and β say $\phi_2(\beta)$, respectively are given by

$$E[\phi_1(\alpha) | data] = \int_0^\infty \phi_1(\alpha) \pi_1(\alpha | data) d\alpha \quad (16)$$

$$\text{and} \quad E[\phi_2(\beta) | data] = \int_0^\infty \phi_2(\beta) \pi_2(\beta | data) d\beta \quad (17)$$

From the above Equations (16) and (17), We note that there are not available solutions. Numerically solving the aforementioned integrals is possible. To get Bayes estimates for the parameters α and B , respectively, we employ (MCMC) methods such the Metropolis-Hastings [21] (M-H) algorithm

3.2.1 Markov Chain Monte Carlo Techniques (MCMC)

Here, to develop and generate sequences of samples from the full conditional distribution's of the parameters, we will use (MCMC) methods. The calculation of the Bayes estimates of the unknown parameters is done using the Metropolis Hastings (M-H) algorithm. For further information on the (MCMC) and (M-H) algorithms, see Robert and Casella [23], Metropolis et al. [21], Hastings [9]. The marginal posterior distributions of the parameters (α) and (β) in Equations (14) and (15), respectively; are not well-known distributions, thus random numbers can be generated from them using the (M-H) algorithm.

Following steps are used for generation of random numbers from the marginal distribution in (14).

Step 1: You must start with an initial guess $\alpha^{(0)}$

Step 2: Generate a candidate point $\alpha_c^{(j)}$ from the proposal density Proposal density $\delta(\alpha^{(j)} | \alpha^{(j-1)})$,

Step 3: Generate u from Uniform $U(0,1)$ distribution

Step 4: Calculate $z(\alpha_c^{(j)} | \alpha^{(j-1)}) = \min \left\{ \frac{\pi_1((\alpha_c^{(j)} | data) \delta((\alpha^{(j)} | \alpha^{(j-1)}))}{\pi_1((\alpha^{(j-1)} | data) \delta((\alpha_c^{(j)} | \alpha^{(j-1)}))}, 1 \right\}$

Step 5: If $u \leq z$ set $\alpha^{(j)} = \alpha_c^{(j)}$ with acceptance probability z otherwise $\alpha^{(j)} = \alpha^{(j-1)}$

Step 6: Steps should be repeated (2-5), for $j = 1, 2, 3, \dots, M$ in order to acquire the parameter's sequence Aas $(\alpha^{(1)}, \alpha^{(2)}, \dots, \alpha^{(M)})$.

In this case, we treat proposal density as a normal distribution (ND). The (MLE) and variance of (MLE) from posterior distribution of (α) are considered as mean and variance of the proposal normal distribution [N tzoufres][22]. We discard first $M_0, \alpha^{(j)}$; $j = 1, 2, \dots, M_0$ where, $M_0 \leq M$ is the burn-in period to get a sample independently from the stationary Markov chain distribution, which is usually the posterior distribution. Now, the approximate posterior mean of $\phi_1(\alpha)$ use (M-H) algorithm is obtained as: $\phi_{1MH}(\alpha) = \frac{1}{M-M_0} \sum_{j=M_0+1}^M \phi_1(\alpha^{(j)})$.

Similarly, the approximate posterior mean of $\phi_2(\beta)$ using Metropolis-Hastings (M-H) algorithm is obtained as $\phi_{2MH}(\beta) = \frac{1}{M-M_0} \sum_{j=M_0+1}^M \phi_1(\beta^{(j)})$. Therefore, the Bayes estimates of the parameters (α) and (β) under (LINEX) loss function using (M-H) algorithm are; respectively give

$$\hat{\alpha}_{MH} = -\frac{1}{k} \ln \left(\hat{\phi}_{1MH}(\alpha) \right) \text{ and } \hat{\beta}_{MH} = -\frac{1}{k} \ln \left(\hat{\phi}_{2MH}(\beta) \right)$$

3.2.2 Highest Posterior Density (HPD) Credible Intervals

Here, Using the generated (MCMC) samples, we build the (HPD) credible intervals of the parameters (α) and (β) . Let $\alpha_{(1)} < \alpha_{(2)} < \dots < \alpha_{(M-M_0)}$ show the values in order of $\alpha^{(M_0+1)}, \alpha^{(M_0+2)}, \dots, \alpha^{(M)}$. Then, the algorithm proposed using by The two researchers (Chen and Shao)[4], the $100(1 - \xi)\%$, where $0 < \xi < 1$, Highest Posterior Density (HPD) credible interval for (α) is given by: $(\alpha_{(j)}, \alpha_{(j)+[(1-\xi)(\frac{1}{M-M_0})]})$, where j is chosen such that

$$\alpha_{(j)+[(1-\xi)(\frac{1}{M-M_0})]} = \text{Min}_{- \leq i \leq (M-M_0)} (\alpha_{(j)+[(1-\xi)(\frac{1}{M-M_0})]} - \alpha_{(j)}),$$

$$j = 1, 2, \dots, M - M_0,$$

where $[x]$ is the largest integer less than or equal to the $[X]$. Similarly. We can construct the $[100(1 - \xi)\%]$, Highest Posterior Density (HPD) credible interval for (β) .

4. Monte Carlo Simulation Study

This part of the article focuses on a study involving the use of Monte Carlo simulation. Here, we evaluate different and compare estimators that were developed in previous sections by using a Monte Carlo simulation studies. Six different sample sizes $n = 30, 40, 50, 60, 70$ and 80 are considered in the simulation study. Following combinations of the true values of the parameters $(\alpha, \beta) = (0.75, 1.5)$ and $(1.5, 0.75)$ are taken. The (ML) and Bayes estimates for the unknown parameters are calculated in each case, In Bayesian computations, hyper parameters $(a_1, b_1, a_2, b_2) = (3, 2, 3, 4)$ and $(3, 4, 3, 2)$ are considered in gamma informative priors [11, 19, and 20]. Two different values of $(k =$

-1 and 1) are taken for (LINEX) loss function. For (MCMC) method ($M = 10,000$) with burn-in-period $M_0 = 1000$ is used. We calculate the 95% confidence intervals using the observed Fisher information matrix, and the (MCMC) techniques are used to determine the (HPD) credible intervals. This process is repeated (1000) times as shown in the following tables, Table 1 to Table 3.

Table 1

The Bayes & MLs of α and β when, $\alpha = 1.5$, $\beta = 0.75$, $a = 3$, $b = 2$, $c = 3$, $q = 4$

n	m	$\hat{\alpha}_{MLE}$		$\hat{\beta}_{MLE}$		$\hat{\alpha}_{Bayes}$				$\hat{\beta}_{Bayes}$			
						$(k = -1)$		$(k = 1)$		$(k = -1)$		$(k = 1)$	
		AE	MSE	AE	MSE	AE	MSE	AE	MSE	AE	MSE	AE	MSE
30	8	1.6087	0.1760	0.7832	0.0333	1.5626	0.1163	1.5135	0.0980	0.7665	0.0275	0.7574	0.0266
40	12	1.5850	0.1278	0.7791	0.0258	1.5548	0.0951	1.5174	0.0832	0.7673	0.0229	0.7604	0.0223
50	18	1.5600	0.0898	0.7794	0.0231	1.5396	0.0731	1.5101	0.0663	0.7700	0.0211	0.7643	0.0207
60	24	1.5449	0.0716	0.7746	0.0204	1.5298	0.0616	1.5055	0.0572	0.7672	0.0191	0.7625	0.0188
70	26	1.5439	0.0603	0.7722	0.0187	1.5315	0.0534	1.5105	0.0501	0.7661	0.0178	0.7621	0.0176
80	25	1.5379	0.0502	0.7716	0.0178	1.5278	0.0455	1.5095	0.0431	0.7663	0.0171	0.7628	0.0169

Table 2

The coverage probability (CP) and average length (AL) of the 95 asymptotic confidence and credible intervals (HPD), For the parameters α and β , $\alpha = 1.5$, $B = 0.75$

n	m	$\hat{\alpha}_{MLE}$		$\hat{\beta}_{MLE}$		$\hat{\alpha}_{Bayes}$		$\hat{\beta}_{Bayes}$	
		AL	CP	AL	CP	AL	CP	AL	CP
30	8	1.3731	0.9691	0.8464	0.8431	0.5719	0.9611	0.3766	0.8711
40	12	1.1638	0.9621	0.7443	0.8411	0.4939	0.9691	0.3323	0.8621
50	18	1.0184	0.9621	0.6660	0.8531	0.4429	0.9611	0.3009	0.8511
60	24	0.9178	0.9641	0.6083	0.8461	0.4026	0.9661	0.2757	0.8401
70	26	0.8488	0.9611	0.5671	0.8501	0.3722	0.9641	0.2566	0.8511
80	25	0.7901	0.9721	0.5310	0.8581	0.3485	0.9651	0.2416	0.8411

Table 3

Summary fit of the Colon cancer patients data (Data I)

Models	MLE	-LogL	AIC	BIC	KS-TEST	
					KS-Statistic	P=value
$X \sim IED(\alpha)$ $T \sim IED(\beta)$	$\hat{\alpha} = 6.5543$ $\hat{\beta} = 79.4764$	141.8648	285.7200	288.5219	0.1856	0.3238
$X \sim IRD(\alpha)$ $T \sim IRD(\beta)$	$\hat{\alpha} = 12.8202$ $\hat{\beta} = 651.1212$	257.5126	517.0151	519.8175	0.5656	1.8×10^{-8}
$X \sim IPD(\alpha)$ $T \sim IPD(\beta)$	$\hat{\alpha} = 9.466$ $\hat{\beta} = 79.3547$	139.7126	281.4150	284.2174	0.1472	0.6357
$X \sim IED(\alpha)$	$\hat{\alpha} = 0.81474$	139.7452	283.4802	287.6838	0.1381	0.7198

$T \sim IED(\beta)$	$\hat{\beta} = 4.9351$ $\hat{\nu} = 35.3623$					
$X \sim IWD(\alpha)$ $T \sim IWD(\beta)$	$\hat{\alpha} = 0.5647$ $\hat{\beta} = 4.8135$ $\hat{\nu} = 64.2851$	140.2895	284.5688	288.7724	0.1700	0.4372

The mean squared error (MSE) and average estimates (AE) for various estimators are calculated. Calculations are also made for the average length (AL) and coverage probabilities (CP) of 95% intervals. Table 4, and Table 5 list the findings of the Monte Carlo simulation investigation.

Table 4
Summary fit of the Hodgkin's disease patient data (Data II)

Models	MLE	-LogL	AIC	BIC	KS-TEST	
					KS-Statistic	P=value
$X \sim IED(\alpha)$ $T \sim IED(\beta)$	$\hat{\alpha} = 6.4343$ $\hat{\beta} = 76.3664$	62.0777	126.1453	127.5614	0.1404	0.9424
$X \sim IRD(\alpha)$ $T \sim IRD(\beta)$	$\hat{\alpha} = 11.7202$ $\hat{\beta} = 631.0712$	123.4034	248.7967	250.2128	0.5684	0.0102
$X \sim IPD(\alpha)$ $T \sim IPD(\beta)$	$\hat{\alpha} = 7.863$ $\hat{\beta} = 77.3696$	61.7592	125.5083	126.9244	0.1067	1.0066
$X \sim IED(\alpha)$ $T \sim IED(\beta)$	$\alpha = 0.7774$ $\hat{\beta} = 4.9231$ $\hat{\nu} = 33.3523$	62.0591	128.1081	130.2323	0.1287	0.6871
$X \sim IWD(\alpha)$ $T \sim IWD(\beta)$	$\hat{\alpha} = 0.6619$ $\hat{\beta} = 4.7952$ $\hat{\nu} = 63.1718$	62.0501	128.0901	130.2142	0.1258	0.9841

Table 5
The ML, Bayes estimates and (95%) Asymptotic and (HPD) credible intervals of the unknown parameters corresponding to Data(I) and Data (II), respectively

Datasets	Parameters	ME	95% CI	Bayes estimates	95% HPD CI
Data I	α β	7.863 773696	(5.0484, 10.6775) (31.6325, 123.1067)	7.6672 72.6036	(5.7135, 9.598) (40.328, 106.8078)
Data II	α β	6.6481 28.1105	(3.2792, 1(1017) (8.3706, 47.8505)	6.3191 25.6675	(4.0975, 8.6322) (12.6852, 403103)

From these results the following conclusions are made:

In almost all cases, average estimates get closer to the true value of the parameters and MSEs decreases as sample size increases. The average length of interval estimates shrink down one sample size. The coverage probabilities almost always reach the specified confidence levels. Due to the inclusion of previous knowledge, Bayes estimates perform better than (MLEs) in terms of

average estimates, MSEs, and average length. The asymptotic confidence intervals are often longer than the (HPD) credible intervals. Therefore, when some prior parameter information is available, we advise using Bayes estimators. For speedy findings in other cases, ML estimators may be employed.

5. Analysis Real Data

In this section, we will use two real datasets to demonstrate the estimating processes covered in previous sections. The datasets we will be using are the Colon Cancer patient's data, which we will refer to as Data I, and the Hodgkin's patient's data, which we will refer to as Data II. Data (I) includes the times of remission, measured in weeks among of a group of thirty individuals with colon cancer who all received a similar treatment. Data II consists of the survival times, measured in months, of (15) patients with Hodgkin's disease who were treated with nitrogen mustards and had previously undergone intensive therapy. The Data I and Data II are presented below.

Data I : 1, 1, 2, 4, 4, 6, 6, 6, 7, 8, 9, 9, 10, 12, 13, 14, 18, 19, 24, 26, 29, 31+, 42, 45+, 50+, 57, 60, 71+, 85+, 91 **Data II :** 1.05, 2.92, 3.61, 4.20, 4.49, 6.72, 7.31, 9.08, 9.11, 14.49+, 16.85, 18.82+, 26.59+, 30.26+, 41.34+

The observations with + sign are censored times. Before going further, we fit Data I and Data II to randomly censored (IPD) and compare its fitting with some well known lifetime models, namely, inverse exponential (IE), inverse Rayleigh (IR), inverse Weibull (IW) and generalized inverted exponential (GI)

We calculate the MLEs of the unknown parameters for both datasets tests some useful goodness-of-fit measure, namely, the negative log likelihood function $-(\ln L)$, the Akaike[1] information criterion denoted by:

$AIC = 2 \times k - 2 \times \ln L$, proposed by Akaike and Bayesian [2] information criterion denoted by $BIC = k \times \ln(n) - 2 \times \ln L$, proposed by Schwarz, where (k) is the number of parameters in the model, (n) is the number of observations in the given datasets, (L) is the maximized value of the likelihood function for the estimated model and Kolmogorov-Smirnov (K-S) statistic with its p-value. The best distribution corresponds to the lowest $-(\ln L)$, (AIC), (BIC) and (K-S) statistic and the highest (p) values. The (K-S) statistic with its (p-value) are obtained using (k-s). The results of the (MLEs) and measures of goodness-of-fit tests are reported in. These results show that randomly censored (IP) distribution is the best choice for the considered datasets. However, for data (I), according to (K-S) test randomly censored (IWD) is slightly better than the randomly censored (IPD). The randomly censored (IWD) has three parameters, while randomly censored (IPD) has only two parameters. Therefore, randomly censored (IPD) is the best choice from computational point of view. Further, ML and Bayes estimates using no informative priors and their corresponding asymptotic (HPD) 95% CIs of the unknown parameters of considered randomly censored (IPD) model corresponding to the above real datasets (Data I and Data II) are reported. From these results, we see that MLEs and Bayes estimates of parameters based on (MCMC) techniques are quite closed.

6. Conclusion

In this article using randomly censored data to estimate inverse Pareto distribution parameters using classical and Bayesian methods. For the unknown parameters, the maximum likelihood estimators and corresponding asymptotic confidence intervals were developed using the observed Fisher information matrix. Using (MCMC) methods, the Bayes estimates of the parameters under the (LINEX) loss function were estimated. The effectiveness of these estimators was investigated by extensive Monte Carlo simulation analysis. It shown that MLEs could be acquired rapidly and readily with accurate estimations. The Bayes estimating method with known prior information or convenient non-informative priors in the absence of prior information is advised for more efficient estimators.

References

- [1] Akaike, H. A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19(6), 716-723 (1974).
- [2] Breslow, N. and J. Crowley. A large sample study of the life table and product limit estimates under random censorship. *The Annals of Statistics*, 2(3), 437-453. Chaturvedi, A., N (1974).
- [3] Kumar and K. Kumar. Statistical inference for the reliability functions of a family of lifetime distributions based on progressive type II right censoring. *Statistical*, 78(1), 81-101 (2018).
- [4] Chen, M.H. and Q. Shao. Monte Carlo estimation of Bayesian credible and HPD intervals. *Journal of Computational and Graphical Statistics*, 8(1), 69-92 (1999).
- [5] Garg, Renu, D. Madhulika and H. Krishna. Estimation of parameters and reliability characteristics in Lindley distribution using randomly censored data. To appear in *Statistics, Optimization and Information* (2019).
- [6] Garg, R., M. Dube, K.Kumar and H. Krishna. On randomly censored generalized inverted exponential distribution. *American Journal of Mathematical and Management Sciences*, 35(4), 361-379 (2016).
- [7] Ghitany, M. and S. Al-Awadhi. Maximum likelihood estimation of Burr XII distribution parameters under random censoring. *J. Applied Statistics*, 29(7), 955-965 (2002).
- [8] Gilbert, J.P. Random Censorship. Ph.D Thesis, University of Chicago, Department of Statistics (1962).

- [9] Hastings, W.K. Monte Carlo sampling methods using Markov chains and their applications. 57, 97-109 (1970).
- [10] Koziol, J.A. and S.B. Green. A Cram–er-Von Mises statistic for randomly censored data. *Biometrika*, 63(3), 465-474 (1976).
- [11] Krishna, H. and N. Goel. Randomly censored geometric distribution under Koziol-Green model. *Int. J. Agricult. Stat. Sci.*, 13(1), 85-95 (2017).
- [12] Krishna, H. and K. Kumar. Reliability estimation in generalized gamma distribution with progressively censored data. *Int. J. Agricult. Stat* (2017).
- [13] Krishna, H. and K. Kumar. Reliability estimation in generalized inverted exponential distribution with progressively type II censored sample. *J. Statistical Computation and Simulation*, 83(6), 1007-1019 (2013).
- [14] Krishna, H., Vivekanand and K. Kumar. Estimation in Maxwell distribution with randomly censored data. *J. Statistical Computation and Simulation*, 85(17), 3560- 3578 (2015).
- [15] Kumar, K. Classical and Bayesian estimation in loglogistic distribution under random censoring. *International Journal of System Assurance Engineering and Management*, 9(2), 440-451 (2018).
- [16] Kumar, K. and R. Garg. Estimation of the parameters of randomly censored generalized inverted Rayleigh distribution. *Sci.*, 10(1), 147-155 (2014).
- [17] Kumar, K. R. Garg and H. Krishna. Nakagami distribution as a reliability model under progressive censoring. *Int. J. System Assurance Engineering and Management*, 8(1), 109-122 (2017).
- [18] Kumar, K. and I. Kumar. Estimation in inverse Weibull distribution based on randomly censored data. *Statistica*, 79(1), 47-74 (2019).
- [19] Lawless, J. F. *Statistical Models and Methods for Lifetime Data*. Flo-boken, New Jersey. Mann, N.R., N.D. (2003).
- [20] Singpurwalla and R.E. Schafer. *Methods for Statistical Analysis of Reliability and Life Data*. Wiley, New York (1974).
- [21] Metropolis, N., A.W. Rosenbluth, M.N. Rosenbluth, A. H. Teller and E. Teller. Equation of state calculations by fast computing machines. *The Journal of Chemical Physics*, 21(6), 1087-1092 (1953).
- [22] Ntzoufras, I. *Bayesian Modeling Using WinBugs*. Wiley, New (2009).

- [23] R Core Team. R: A Language and Environment for Statistical Computing. R Foundation for Statistical Computing, Vienna, Austria (2018).
- [24] Robert, C. and G. Casella. Monte Carlo Statistical Methods. Verlag NY (2004).
- [25] Schwarz, G. Estimating the dimension of a model. *The Annals of Statistics*, 6(2), 461-464 (1978).
- [26] Sinha, S.K. Reliability and Life Testing. Wiley Eastern Limited, New Delhi (1986).
- [27] Varian, H.R. A Bayesian approach to real estate assessment. Studies in Bayesian Econometric and Statistics in Honor of Leonard J. Savage Savage, eds. Stephen E. Fienberg and Arnold (1975)
- [28] Zellner, Amsterdam: North- Hollan. Zellner, A. Bayesian estimation and prediction using asymmetric loss functions. *Journal of the American Statistical Association*, 81(394), 446-451 (1986).

Received January, 2024