

Estimating Parameters from the Generalized Inverse Lindley Distribution Under Hybrid Censoring Scheme

Mojammel Haque Sarkar^a, Manas Ranjan Tripathy^{a*} and Debasis Kundu^b

^a*Department of Mathematics, National Institute of Technology Rourkela, Rourkela-769008, India*

^b*Department of Mathematics and Statistics, Indian Institute of Technology Kanpur, Kanpur-208 016, India*

Abstract: Estimation of parameters of the generalized inverse Lindley (GIL) distribution is considered under a hybrid censoring scheme. The point estimators, such as the maximum likelihood estimators using the Expectation-Maximization (E-M) algorithm, have been derived. The two approximate Bayes estimators using Tierney and Kadane's method and Gibbs sampling procedure, using the gamma prior and the general entropy loss (GEL) function, have been obtained. Several confidence intervals are proposed, such as the asymptotic confidence intervals (ACIs), bootstrap confidence intervals, and the highest posterior density (HPD) credible intervals. The prediction for future observations has been considered under one and two-sample Bayesian prediction methods using the type-i hybrid censoring scheme. An extensive simulation study has been conducted to numerically evaluate all the estimators' performances. The point estimators are compared through their biases and mean squared errors (MSEs). The performances of confidence intervals are evaluated using coverage probability (CP), average length (AL), and probability coverage density (PCD). Two real-life datasets have been considered for illustrative purposes.

AMS 2001 subject classification: 62C15, 62C20, 62F10

Keywords: Approximate Bayes estimator; Asymptotic confidence interval; Bayesian prediction; Bootstrap confidence interval; Expectation-Maximization algorithm; Gibbs sampling; Maximum likelihood estimator; Numerical comparison.

1 Introduction

The generalized version of a statistical model gives better flexibility for modeling and analyzing real-life data which arise in practice. One of such generalizations is the generalized inverse Lindley (GIL) distribution which has been proposed by authors as an alternative choice for modeling lifetime data. The GIL distribution was proposed by Sharma et al. (2016), which is an extension of the inverse Lindley distribution. The probability density function and the cumulative distribution function of the

*Corresponding author. E.mail: manasmath@yahoo.co.in(Manas Ranjan Tripathy)

GIL distribution are respectively given by

$$f(x; \alpha, \lambda) = \frac{\alpha \lambda^2}{\lambda + 1} (1 + x^{-\alpha}) x^{-\alpha-1} e^{-\lambda x^{-\alpha}}, \quad x > 0, \alpha > 0, \lambda > 0 \quad (1.1)$$

and

$$F(x; \alpha, \lambda) = \left(1 + \frac{\lambda}{\lambda + 1} x^{-\alpha}\right) e^{-\lambda x^{-\alpha}}, \quad x > 0, \alpha > 0, \lambda > 0. \quad (1.2)$$

The parameters α and λ both describe the shape of the distribution; however, α controls the rate of decay of the upper tail. Hence the parameter α is also known as the ‘heavy tailed’ parameter. This article considers both point and interval estimation of the associated model parameters α and λ of the GIL distribution using the type-i hybrid censoring scheme. The Bayes estimators for the associated parameters will be considered with respect to the general entropy loss (GEL) function,

$$L(\delta, \theta) \propto \left(\frac{\delta}{\theta}\right)^w - w \log \left(\frac{\delta}{\theta}\right) - 1, w \neq 0 \quad (1.3)$$

where δ is an estimator for estimating the parameter θ .

Note that the inverse version of a probability distribution shows the upside-down bathtub shapes for their hazard rates which enables to model certain practical situations conveniently. Sharma et al. (2016) discussed some practical problems where the datasets can be conveniently modeled using an inverse probability distribution. The shapes of the probability density function and the hazard rate function of the GIL distribution certainly differentiate from other distributions and make it flexible for modeling wide ranges of lifetime datasets arising in practice.

Much research has been done in estimating the parameters of various probability models using some types of censoring schemes in the past from classical and Bayesian points of view. The most commonly used censoring schemes available are type-i (the number of failures is random and the censoring time is fixed), type-ii (the number of failures is fixed and the censoring time is random) or some mixture of these two schemes. One of the modifications or hybridization of these two types of censoring schemes is the type-i hybrid censoring scheme.

Probably, Epstein (1954) was the first to introduce the type-i hybrid censoring scheme. Using this censoring scheme, the author proposed various point and interval estimators of the parameters of an exponential distribution. Though this censoring scheme had been considered a long back, only during the last two decades more attention has been paid by several authors. We refer to Draper and Guttman (1987), Kundu and Pradhan (2009), Dey and Pradhan (2014), Tripathi and Rastogi (2016), Sultana et al. (2018), Kayal et al. (2018), Abushal (2019), and the references therein, for some results on estimating parameters of certain lifetime statistical models using the type-i hybrid censoring scheme. Authors have also considered the Lindley distribution and its various versions as some alternative for modeling lifetime data arising in practice. We refer to Gupta and Singh (2013), Basu et al. (2019) and Singh et al. (2016) for some results on estimating parameters of Lindley, inverse Lindley and generalized Lindley distributions, respectively, under type-i hybrid censoring scheme. Pandey et al. (2021) considered the point and interval estimation of the parameters of Lindley distribution using general progressive hybrid censored samples under a step-stress partially accelerated life test model. Sen et al. (2019) considered the generalized exponential distribution and estimated the associated parameters when the samples are type-ii progressive hybrid. The authors also predicted the future observation using one and two sample prediction procedures. The estimation of parameters from inverse Weibull distribution using adoptive type-ii progressive hybrid censoring has been considered by Nassar and Abo-Kasem (2017).

The type-i hybrid censoring has certain merits over the two conventional censoring schemes because it saves time and cost. Suppose n identical units are put for a life-testing experiment. The experiment

is terminated when a preassigned number, say r out of n items, has failed or a predetermined time, say T , has been reached. Under type-i hybrid censoring scheme, the samples are observed as follows.

$$\begin{cases} \text{C1: } X_{1:n} \leq X_{2:n} \leq \dots \leq X_{r:n}, \text{ if } X_{r:n} \leq T, \\ \text{C2: } X_{1:n} \leq X_{2:n} \leq \dots \leq X_{m:n}, \text{ if } 0 < m < r, X_{m:n} \leq T < X_{m+1:n}, \end{cases} \quad (1.4)$$

where we denote $X_{1:n} \leq X_{2:n}, \dots$ as the observed failure times of the experimental units. In C2, m number of failures occur up to time T and $(m+1)^{th}$ failure occurs after time T .

The type-i hybrid censoring scheme has certain advantages over the type-ii hybrid censoring scheme. For example, in type-i hybrid censoring scheme the termination time of the experiment is minimum of a prefixed time (T) and the r^{th} failure time ($X_{r:n}$), that is $\min(T, X_{r:n})$. That is, the experiment time will never exceed the prefixed time T , which is a clear advantage. However, this is not the case with the type-ii hybrid censoring scheme; the experiment's termination time is unknown, which may be a disadvantage. There are certain advantages and disadvantages of these two types of hybrid censoring schemes. However, one may prefer one over the other under given situations. A detailed comparison of these two types of hybrid censoring methods has been made by Childs et al. (2003). The results obtained in this paper using the type-i hybrid censoring scheme can also be derived under the type-ii hybrid scheme in a very similar fashion. Moreover, we have also discussed the estimation procedure using the type-ii hybrid censoring scheme in Section 7 (though the details have been omitted for brevity).

The rest of the paper is organized as follows. The MLEs of the associated parameters could not be obtained in closed-form expressions; hence we use the E-M algorithm to obtain it numerically in Section 2. Further, the asymptotic confidence intervals (ACIs) are derived using the observed Fisher information matrix. The Bayesian estimation of the associated parameters is considered using the gamma prior and the GEL function in Section 3. The Bayes estimators are derived using the approximation methods of Tierney and Kadane (1986) and Markov-chain-Monte-Marlo (MCMC). The MCMC posterior samples are utilized to construct the highest probability density (HPD) credible intervals for the parameters. In Section 4, we have compared the performances of all the proposed estimators (both point and interval) using an extensive simulation study. The point estimators are compared through their biases and the MSEs. The performances of interval estimators are evaluated in terms of coverage probability (CP), average length (AL) and a unique measure known as probability coverage density (PCD). Section 5 considers one and two-sample Bayesian prediction procedures to predict future observations. In Section 6, we consider two real-life situations which have been modeled using the GIL distribution, and the estimation methodologies have been explained. Section 7 gives a glimpse of the inferential results using the type-ii hybrid censoring scheme for completeness from the reader's point of view. In Section 8, we conclude the results obtained in this article.

Remark 1.1. *Surprisingly and interestingly, not much study has been done on estimating parameters of the GIL distribution using censored samples, mainly utilizing the hybrid censoring scheme. The reason may be the complicated nature of the probability density function, which prohibits deriving any nice theoretical inferential results. However, this distribution has been proposed as an alternative for lifetime data analysis. It is essential to estimate the associated parameters using complete and censored samples. In this paper, for simplicity and convenience, we have obtained the results under type-i and type-ii hybrid (a glimpse of the idea has been given in Section 7) censoring schemes. The results also can be obtained under some more generalized hybrid censoring schemes, see for example Górný and Cramer (2018), which may not be straightforward.*

2 The MLE & Asymptotic Confidence Interval

Let $X_1, X_2, \dots, X_n$ be a random sample of size n from the GIL distribution with unknown parameters α and λ . Then a type-i hybrid censored sample is observed as given in (1.4) satisfying the condition C1 or C2. In the first case (C1), we observe the lifetimes of r units up to the time T , whereas in the second case (C2), m number of failures have been observed up to the time T and $(m+1)^{th}$ failure occurs after the time T . Therefore, the likelihood function of the parameters α and λ can be written as

$$l(\alpha, \lambda) \propto \begin{cases} \prod_{i=1}^r f(x_{i:n}) [1 - F(x_{r:n})]^{(n-r)}, & \text{if condition C1 holds,} \\ \prod_{i=1}^m f(x_{i:n}) [1 - F(T)]^{(n-m)}, & \text{if condition C2 holds} \end{cases}$$

Combining the above two conditions, we obtain the general form of the log-likelihood function as

$$\begin{aligned} L(\alpha, \lambda) \propto & d(\log \alpha + 2 \log \lambda - \log \lambda^*) + \sum_{i=1}^d \log(1 + x_{i:n}^{-\alpha}) - \lambda \sum_{i=1}^d x_{i:n}^{-\alpha} \\ & - \alpha^* \sum_{i=1}^d \log x_{i:n} + (n-d) \log \left\{ 1 - \left(1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right) e^{-\lambda c^{-\alpha}} \right\}. \end{aligned} \quad (2.1)$$

where $\alpha^* = 1 + \alpha$, $\lambda^* = 1 + \lambda$ and

$$c = \begin{cases} x_{r:n}, & \text{if C1 holds} \\ T, & \text{if C2 holds} \end{cases} \quad \text{and} \quad d = \begin{cases} r, & \text{if C1 is true} \\ m, & \text{if C2 is true.} \end{cases}$$

Note that, to obtain the MLEs of α and λ , we must have $d > 0$, as for $d = 0$, the MLEs do not exist. Differentiating the log-likelihood function $L(\alpha, \lambda)$ with respect to α and λ and equating to zero, we have the following system of two non-linear equations in α and λ

$$\frac{d}{\alpha} - \sum_{i=1}^d \frac{x_{i:n}^{-\alpha} \log x_{i:n}}{1 + x_{i:n}^{-\alpha}} - \sum_{i=1}^d \log x_{i:n} + \lambda \sum_{i=1}^d x_{i:n}^{-\alpha} \log x_{i:n} + \lambda^2 (n-d) \frac{(1 + c^{-\alpha}) \log c}{h(c)} = 0 \quad (2.2)$$

$$\frac{2d}{\lambda} - \frac{d}{\lambda^*} - \sum_{i=1}^d x_{i:n}^{-\alpha} - (n-d) \frac{\lambda}{\lambda^*} \frac{2 + \lambda + \lambda^* c^{-\alpha}}{h(c)} = 0, \quad (2.3)$$

where $h(c) = \lambda - c^\alpha \lambda^* (e^{\lambda c^{-\alpha}} - 1)$. The above system of non-linear equations ((2.2)-(2.3)) can be solved using certain numerical methods, such as Newton's method. However, this iteration scheme may not converge always (see Pradhan and Kundu (2009)). Hence, we propose to employ the Expectation and Maximization (E-M) algorithm to obtain the MLEs of the parameters. The E-M algorithm was proposed by Dempster et al. (1977) and has been extensively used in the literature to obtain the MLEs of the parameters in the case when the data are incomplete. This method guarantees the convergence though its rate of convergence sometimes is very slow (see Little and Rubin (1983)).

Suppose $X = (X_{1:n}, X_{2:n}, \dots, X_{d:n})$ is the observed data and $Z = (Z_1, Z_2, \dots, Z_{n-d})$ is the censored data. The unobserved data Z can be considered as the missing data, and thus $W = (X, Z)$ forms the complete data set of the experiment. Hence, the associated log-likelihood function can be obtained as

$$\begin{aligned} L_c(\alpha, \lambda) = & n \log \alpha + 2n \log \lambda - n \log \lambda^* + \sum_{i=1}^d \log(1 + x_{i:n}^{-\alpha}) - \lambda \sum_{i=1}^d x_{i:n}^{-\alpha} \\ & - \alpha^* \sum_{i=1}^d \log x_{i:n} + \sum_{i=1}^{n-d} \log(1 + z_i^{-\alpha}) - \alpha^* \sum_{i=1}^{n-d} \log z_i - \lambda \sum_{i=1}^{n-d} z_i^{-\alpha}. \end{aligned} \quad (2.4)$$

Taking the derivative of $L_c(\alpha, \lambda)$ with respect to α and λ , and equating to zero, we get

$$\alpha = \left[\frac{1}{n} \sum_{i=1}^d \frac{\log x_{i:n}}{1 + x_{i:n}^\alpha} + \frac{1}{n} \sum_{i=1}^d \log x_{i:n} - \frac{\lambda}{n} \sum_{i=1}^d x_{i:n}^{-\alpha} \log x_{i:n} + \frac{1}{n} \sum_{i=1}^{n-d} \frac{\log z_i}{1 + z_i^\alpha} + \frac{1}{n} \sum_{i=1}^{n-d} \log z_i - \frac{\lambda}{n} \sum_{i=1}^{n-d} z_i^{-\alpha} \log z_i \right]^{-1} \quad (2.5)$$

$$\lambda = \left[\frac{1}{2(1 + \lambda)} + \frac{1}{2n} \sum_{i=1}^d x_{i:n}^{-\alpha} + \frac{1}{2n} \sum_{i=1}^{n-d} z_i^{-\alpha} \right]^{-1}. \quad (2.6)$$

In order to apply the E-M iteration method, we proceed as follows. In the expectation step (E-step), the missing data are replaced by their expected values, while in the maximization step (M-step), the log-likelihood function is maximized with respect to the observed data and the expected value of the censored data updating the values of the estimators. Let $(\alpha^{(k)}, \lambda^{(k)})$ be the estimate of (α, λ) at k^{th} stage, then the estimators at the $(k+1)^{th}$ stage are obtained as

$$\begin{aligned} \lambda^{(k+1)} &= \left[\frac{1}{2(1 + \lambda^{(k)})} + \frac{1}{2n} \sum_{i=1}^d x_{i:n}^{-\alpha^{(k)}} + \frac{1}{2} \left(1 - \frac{d}{n}\right) D(c, \alpha^{(k)}, \lambda^{(k)}) \right]^{-1} \\ \alpha^{(k+1)} &= \left[\frac{1}{n} \sum_{i=1}^d \frac{\log x_{i:n}}{1 + x_{i:n}^{\alpha^{(k)}}} + \frac{1}{n} \sum_{i=1}^d \log x_{i:n} - \frac{\lambda^{(k+1)}}{n} \sum_{i=1}^d x_{i:n}^{-\alpha^{(k)}} \log x_{i:n} \right. \\ &\quad \left. + \left(1 - \frac{d}{n}\right) A(c, \alpha^{(k)}, \lambda^{(k)}) + \left(1 - \frac{d}{n}\right) B(c, \alpha^{(k)}, \lambda^{(k)}) - \lambda^{(k+1)} \left(1 - \frac{d}{n}\right) C(c, \alpha^{(k)}, \lambda^{(k)}) \right]^{-1}. \end{aligned}$$

where

$$\begin{aligned} A(c, \alpha, \lambda) &= E\left(\frac{\log Z}{1 + Z^\alpha} \mid Z > c\right) = \frac{1}{1 - F_X(c; \alpha, \lambda)} \int_c^\infty \frac{\log x}{1 + x^\alpha} f_X(x) dx, \\ B(c, \alpha, \lambda) &= E\left(\log Z \mid Z > c\right) = \frac{1}{1 - F_X(c; \alpha, \lambda)} \int_c^\infty \log x f_X(x) dx, \\ C(c, \alpha, \lambda) &= E\left(Z^{-\alpha} \log Z \mid Z > c\right) = \frac{1}{1 - F_X(c; \alpha, \lambda)} \int_c^\infty x^{-\alpha} \log x f_X(x) dx, \\ D(c, \alpha, \lambda) &= E\left(Z^{-\alpha} \mid Z > c\right) = \frac{1}{1 - F_X(c; \alpha, \lambda)} \int_c^\infty x^{-\alpha} f_X(x) dx. \end{aligned}$$

A stopping criteria, such as the distance between two consecutive iterations is as smaller as possible, that is $\sqrt{(\alpha^{(k+1)} - \alpha^{(k)})^2 + (\lambda^{(k+1)} - \lambda^{(k)})^2} < \epsilon$, where ϵ is a preassigned small positive number, has been used to check the convergence of the iteration scheme. Let $\hat{\alpha}_{ML}$ and $\hat{\lambda}_{ML}$ be the solutions obtained by using the E-M method, and these are the MLEs of the parameters α and λ respectively.

Next, we will use the missing value principle to obtain the observed Fisher information matrix (see Louis (1982)) and consequently determine the ACIs for the parameters α and λ numerically. Using the missing value principle, the information matrix can be written as

$$I_X(\theta) = I_W(\theta) - I_{W|X}(\theta), \quad (2.7)$$

where $\theta = (\alpha, \lambda)$, X = observed sample value, W = complete sample value, $W|X$ = missing sample value. The expressions $I_W(\theta)$, $I_X(\theta)$ and $I_{W|X}(\theta)$ denote the information matrices based on complete data,

incomplete data and missing data respectively. The expressions $I_W(\theta)$ and $I_{W|X}(\theta)$ are respectively given by

$$I_W(\theta) = -E\left[\frac{\partial^2 L_c(W; \theta)}{\partial \theta^2}\right] \quad \text{and} \quad I_{W|X}(\theta) = -(n-d)E_{Z|X}\left[\frac{\partial^2 \log f(z|X, \theta)}{\partial \theta^2}\right].$$

In our case the above two matrices (of order 2×2) can be written as

$$I_W(\theta) = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \quad \text{and} \quad I_{W|X}(\theta) = (n-d) \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix},$$

where

$$a_{11} = -E\left[\frac{\partial^2 L_c}{\partial \alpha^2}\right], a_{12} = -E\left[\frac{\partial^2 L_c}{\partial \alpha \partial \lambda}\right] = a_{21}, a_{22} = -E\left[\frac{\partial^2 L_c}{\partial \lambda^2}\right]$$

and

$$b_{11} = -E_{Z|X}\left[\frac{\partial^2 \log f_Z(z|X, \theta)}{\partial \alpha^2}\right], b_{22} = -E_{Z|X}\left[\frac{\partial^2 \log f_Z(z|X, \theta)}{\partial \lambda^2}\right],$$

$$b_{12} = -E_{Z|X}\left[\frac{\partial^2 \log f_Z(z|X, \theta)}{\partial \alpha \partial \lambda}\right] = b_{21}.$$

The expressions a_{ij} s and b_{ij} s could not be obtained in closed forms. Hence we have used numerical methods to get the approximate values numerically. The details of these expressions are given in the appendix.

Utilizing all the above expressions, the observed information matrix and consequently the observed variance-covariance matrix are computed. It seems difficult to obtain the exact confidence intervals for the parameters α and λ as the exact distributions of the MLEs of the parameters α and λ could not be obtainable. However, it is possible to obtain the ACIs utilizing the large sample theory for the MLEs. According to this principle, the asymptotic sampling distribution of $(\hat{\alpha}_{ML} - \alpha, \hat{\lambda}_{ML} - \lambda)^T \sim N_2(0, [I_X(\hat{\alpha}_{ML}, \hat{\lambda}_{ML})]^{-1})$ as $n \rightarrow \infty$. Utilizing this information, one gets the two-sided $100 \times (1 - \psi)\%$ ACIs respectively for α and λ as

$$\left\{ \begin{array}{ll} \hat{\alpha}_L = \hat{\alpha}_{ML} - z_{\psi/2} \sqrt{Var(\hat{\alpha}_{ML})}, & \hat{\alpha}_U = \hat{\alpha}_{ML} + z_{\psi/2} \sqrt{Var(\hat{\alpha}_{ML})} \\ \hat{\lambda}_L = \hat{\lambda}_{ML} - z_{\psi/2} \sqrt{Var(\hat{\lambda}_{ML})}, & \hat{\lambda}_U = \hat{\lambda}_{ML} + z_{\psi/2} \sqrt{Var(\hat{\lambda}_{ML})} \end{array} \right\},$$

where z_{ψ} is the ψ^{th} percentile of the standard normal distribution and

$$[I_X(\hat{\alpha}_{ML}, \hat{\lambda}_{ML})]^{-1} = \begin{bmatrix} Var(\hat{\alpha}_{ML}) & Cov(\hat{\alpha}_{ML}, \hat{\lambda}_{ML}) \\ Cov(\hat{\lambda}_{ML}, \hat{\alpha}_{ML}) & Var(\hat{\lambda}_{ML}) \end{bmatrix}.$$

Looking at the expressions of $\hat{\alpha}_L$ and $\hat{\lambda}_L$, it is quite possible that the lower bounds of the ACIs may be negative, where the true value of the parameters are positive. To avoid this situation, one may consider lower bounds as the $\max(0, \hat{\alpha}_L)$ and $\max(0, \hat{\lambda}_L)$ instead of $\hat{\alpha}_L$ and $\hat{\lambda}_L$ for obtaining the ACIs of α and λ respectively. Another approach is to use the normal approximation on the log-transformation of the MLEs of the parameters instead of the normal approximation on the MLEs (see Meeker and Escobar (1998)). In this case, the asymptotic sampling distribution of $(\log \hat{\alpha}_{ML} - \log \alpha, \log \hat{\lambda}_{ML} - \log \lambda) \sim N_2(0, V_X(\log \hat{\alpha}_{ML}, \log \hat{\lambda}_{ML}))$ as $n \rightarrow \infty$ where $V_X(\log \hat{\alpha}_{ML}, \log \hat{\lambda}_{ML})$ is the variance-covariance matrix

of $\log \hat{\alpha}_{ML}$ and $\log \hat{\lambda}_{ML}$. The variance-covariance matrix $V_X(\log \hat{\alpha}_{ML}, \log \hat{\lambda}_{ML})$ can be derived easily using the well-known delta method (see Greene (2018)) as

$$V_X(\log \hat{\alpha}_{ML}, \log \hat{\lambda}_{ML}) = \begin{bmatrix} \text{Var}(\hat{\alpha}_{ML})/\hat{\alpha}_{ML}^2 & \text{Cov}(\hat{\alpha}_{ML}, \hat{\lambda}_{ML})/(\hat{\alpha}_{ML}\hat{\lambda}_{ML}) \\ \text{Cov}(\hat{\lambda}_{ML}, \hat{\alpha}_{ML})/(\hat{\alpha}_{ML}\hat{\lambda}_{ML}) & \text{Var}(\hat{\lambda}_{ML})/\hat{\lambda}_{ML}^2 \end{bmatrix}.$$

Therefore based on the log-transformed MLEs, the two-sided $100 \times (1 - \psi)\%$ modified ACIs (denote it as MACIs) of α and λ are respectively obtained as

$$\left\{ \begin{array}{l} \hat{\alpha}_L = \hat{\alpha}_{ML} \cdot \exp\left(-\frac{z_{\psi/2} \sqrt{\text{Var}(\hat{\alpha}_{ML})}}{\hat{\alpha}_{ML}}\right), \quad \hat{\alpha}_U = \hat{\alpha}_{ML} \cdot \exp\left(\frac{z_{\psi/2} \sqrt{\text{Var}(\hat{\alpha}_{ML})}}{\hat{\alpha}_{ML}}\right) \\ \hat{\lambda}_L = \hat{\lambda}_{ML} \cdot \exp\left(-\frac{z_{\psi/2} \sqrt{\text{Var}(\hat{\lambda}_{ML})}}{\hat{\lambda}_{ML}}\right), \quad \hat{\lambda}_U = \hat{\lambda}_{ML} \cdot \exp\left(\frac{z_{\psi/2} \sqrt{\text{Var}(\hat{\lambda}_{ML})}}{\hat{\lambda}_{ML}}\right) \end{array} \right\}.$$

Remark 2.1. In Section 4, we have computed the MLEs, the ACIs and the MACIs for the parameters α and λ using the Monte-Carlo simulation procedure for several choices of sample sizes and censoring schemes, numerically.

3 Bayesian Estimation

The Bayes estimators play an important role when some prior information regarding the behavior of the parameters is available to the statistician for inference purposes. These Bayes estimators may dominate certain classical estimators, such as the MLE. In view of the above, we try to investigate the performances of certain Bayes estimators when type-i hybrid censored samples are available from the GIL distribution. In this section, we derive the Bayes estimators for the associated parameters α and λ concerning the GEL function (1.3).

It should be noted that the GEL function is an asymmetric loss function and was introduced by Calabria and Pulcini (1994). It is also interesting to note that this loss function behaves like the squared error loss function when $w = -1$. Further for $w = 1$ it behaves like the weighted squared error loss function of the form $(\delta - \theta)^2/\theta$. Here w is the shape parameter of this loss function. We further noted that, for $w < 0$, underestimation is more serious than overestimation, and the opposite is true when $w > 0$.

In order to derive the Bayes estimators, we assume that the parameters α and λ are statistically independent. We consider $G_\alpha(a, b)$ and $G_\lambda(p, q)$ as priors for the parameters α and λ respectively. Here $G(a, b)$ denotes a gamma population with shape parameter a and scale parameter b . The joint posterior density function of (α, λ) , given the sample $\underline{X}$ can be written as

$$\pi(\alpha, \lambda | \underline{x}) = k G_\alpha(a, b) G_\lambda(p, q) l(\alpha, \lambda | \underline{x}) = k g(\alpha, \lambda, \underline{x}), \quad (3.1)$$

where

$$k^{-1} = \int_0^\infty \int_0^\infty g(\alpha, \lambda, \underline{x}) d\alpha d\lambda,$$

and

$$\begin{aligned} g(\alpha, \lambda, \underline{x}) &= \frac{\alpha^{d+a-1} \lambda^{2d+p-1}}{(\lambda^*)^d} \left[1 - \left\{ 1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right\} e^{-\lambda c^{-\alpha}} \right]^{n-d} e^{-\lambda(q + \sum_{i=1}^d x_{i:n}^{-\alpha}) - b\alpha} \\ &\quad \times \prod_{i=1}^d (1 + x_{i:n}^{-\alpha}) x_{i:n}^{-\alpha-1}. \end{aligned} \quad (3.2)$$

The Bayes estimators of α and λ concerning the GEL function are given by

$$\hat{\alpha}_{EB} = \left[E(\alpha^{-w} | \tilde{X} = \tilde{x}) \right]^{-\frac{1}{w}} = \left[k \int_0^\infty \int_0^\infty \alpha^{-w} g(\alpha, \lambda, \tilde{x}) d\alpha d\lambda \right]^{-\frac{1}{w}}, \quad (3.3)$$

and

$$\hat{\lambda}_{EB} = \left[E(\lambda^{-w} | \tilde{X} = \tilde{x}) \right]^{-\frac{1}{w}} = \left[k \int_0^\infty \int_0^\infty \lambda^{-w} g(\alpha, \lambda, \tilde{x}) d\alpha d\lambda \right]^{-\frac{1}{w}}. \quad (3.4)$$

It doesn't seem easy to evaluate the above ((3.3) and (3.4)) ratio of integrals analytically. Authors frequently used two popular approximation methods, such as the Lindley's method (Lindley (1980)) and Tierney and Kadane's method (Tierney and Kadane (1986)) for evaluating the above ratio of integrals. Tierney and Kadane's method has certain advantages over Lindley's method, in the sense that it requires dealing with only first and second-order derivatives of the log-likelihood function, unlike the case of Lindley's method where one needs to evaluate the third-order derivatives. However, in certain situations, these two approximation methods produce very similar results. In the following subsection, we prefer to use Tierney and Kadane's approximation method for obtaining Bayes estimators.

3.1 Tierney and Kadane's Approximation Method

In this subsection, we derive the Bayes estimators for the parameters α and λ using the approximation method proposed by Tierney and Kadane (1986) and the GEL function.

The Bayes estimators of a function of θ , say $u(\theta)$, with respect to the GEL function using the Tierney and Kadane's approximation method is obtained as

$$\hat{u}_{TK} = \left\{ E(u(\theta)^{-w} | \text{data}) \right\}^{-\frac{1}{w}} = \left(\left[\frac{\det \Sigma_\star}{\det \Sigma_o} \right]^{\frac{1}{2}} \exp \{ n(L_\star(\theta_\star) - L_o(\theta_o)) \} \right)^{-\frac{1}{w}},$$

where $L_o(\theta_o) = \max\{L(\theta)\}_{\theta_o}$, $L_\star(\theta_\star) = \max\{L_\star(\theta)\}_{\theta_\star}$, $L_o(\theta) = \frac{1}{n}[L(\theta) + \log(v(\theta))]$, $L_\star(\theta) = L_o(\theta) + \frac{1}{n} \log(u(\theta))$, $L(\theta)$ is the log-likelihood function and $v(\theta)$ is the logarithm of the prior density function. In our model, $\theta = (\alpha, \lambda)$ and

$$\begin{aligned} L_o(\alpha, \lambda) = & \frac{1}{n} \left[d(\log \alpha + 2 \log \lambda - \log \lambda^*) + \sum_{i=1}^d \log(1 + x_{i:n}^{-\alpha}) - \alpha^* \sum_{i=1}^d \log x_{i:n} \right. \\ & - \lambda \sum_{i=1}^d x_{i:n}^{-\alpha} + (n-d) \log \left(1 - \left(1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right) e^{-\lambda c^{-\alpha}} \right) \\ & \left. + (a-1) \log \alpha + (p-1) \log \lambda - b\alpha - q\lambda \right]. \end{aligned}$$

Let $(\hat{\alpha}_0, \hat{\lambda}_0)$ be the choice of (α, λ) that maximizes $L_o(\alpha, \lambda)$. Utilizing $(\hat{\alpha}_0, \hat{\lambda}_0)$ we obtain the determinant of Σ_o as

$$\det \Sigma_o = \left[\frac{\partial^2 L_o}{\partial \alpha^2} \frac{\partial^2 L_o}{\partial \lambda^2} - \frac{\partial^2 L_o}{\partial \alpha \partial \lambda} \frac{\partial^2 L_o}{\partial \lambda \partial \alpha} \right]_{(\hat{\alpha}_0, \hat{\lambda}_0)}^{-1}, \quad (3.5)$$

where the terms involved in (3.5) have been computed numerically. In order to evaluate the desired Bayes estimator of α with respect to the GEL function, we choose $u(\alpha, \lambda) = \alpha^{-w}$ and consequently obtain

$$L_\star^\alpha(\alpha, \lambda) = L_o(\alpha, \lambda) - \frac{w}{n} \log \alpha. \quad (3.6)$$

Further, let $(\hat{\alpha}_*, \hat{\lambda}_*)$ be the choice for which the function $L_*^\alpha(\alpha, \lambda)$ is maximized. Utilizing $(\hat{\alpha}_*, \hat{\lambda}_*)$ we obtain

$$\det \Sigma_*^\alpha = \left[\left(\frac{\partial^2 L_*^\alpha}{\partial \alpha^2} \right) \left(\frac{\partial^2 L_*^\alpha}{\partial \lambda^2} \right) - \left(\frac{\partial^2 L_*^\alpha}{\partial \alpha \partial \lambda} \right) \left(\frac{\partial^2 L_*^\alpha}{\partial \lambda \partial \alpha} \right) \right]_{(\hat{\alpha}_*, \hat{\lambda}_*)}^{-1}, \quad (3.7)$$

where the terms involved in (3.7) have been computed numerically.

Finally, the desired approximate Bayes estimator of α using Tierney and Kadane (1986) method turns out to be

$$\hat{\alpha}_{TK} = \left(\left[\frac{\det \Sigma_*^\alpha}{\det \Sigma_o} \right]^{\frac{1}{2}} \exp \{n(L_*^\alpha(\theta_*) - L_o(\theta_o))\} \right)^{-\frac{1}{w}}. \quad (3.8)$$

In a very similar manner, we obtain the approximate Bayes estimator of λ , taking $u(\alpha, \lambda) = \lambda^{-w}$. The approximate Bayes estimator of λ is obtained as

$$\hat{\lambda}_{TK} = \left(\left[\frac{\det \Sigma_*^\lambda}{\det \Sigma_o} \right]^{\frac{1}{2}} \exp \{n(L_*^\lambda(\theta_*) - L_o(\theta_o))\} \right)^{-\frac{1}{w}}. \quad (3.9)$$

Remark 3.1. *Tierney and Kadane's method for finding approximate Bayes estimators does not help in obtaining confidence intervals of the parameters. However, the MCMC sampling technique can be used to obtain the Bayes estimators as well as the HPD credible intervals for the parameters α and λ . This we will discuss in the next subsection.*

3.2 Estimation Using MCMC Method

In this subsection, we consider the MCMC method for computing Bayes estimators of the parameters α and λ of the GIL distribution. We use an MCMC method, such as the Gibbs sampling along with the Metropolis-Hastings method, to generate samples from the conditional posterior densities of the concerned parameters.

The posterior joint density of α and λ is given by

$$\begin{aligned} \pi(\alpha, \lambda | \underline{x}) &\propto \frac{\alpha^{d+a-1} \lambda^{2d+q-1}}{(\lambda^*)^d} \left\{ 1 - \left(1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right) e^{-\lambda c^{-\alpha}} \right\}^{n-d} \\ &\times e^{-b\alpha - \lambda(q + \sum_{i=1}^d x_{i:n}^{-\alpha})} \times \prod_{i=1}^d x_{i:n}^{-\alpha^*} (1 + x_{i:n}^{-\alpha}). \end{aligned} \quad (3.10)$$

Hence the conditional posterior densities of α and λ are respectively given by

$$\begin{aligned} \pi_1(\alpha | \lambda, \underline{x}) &\propto \alpha^{d+a-1} \left\{ 1 - \left(1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right) e^{-\lambda c^{-\alpha}} \right\}^{n-d} \\ &\times e^{-b\alpha - \lambda(q + \sum_{i=1}^d x_{i:n}^{-\alpha})} \times \prod_{i=1}^d x_{i:n}^{-\alpha^*} (1 + x_{i:n}^{-\alpha}), \end{aligned} \quad (3.11)$$

and

$$\pi_2(\lambda | \alpha, \underline{x}) \propto \frac{\lambda^{2d+p-1}}{(\lambda^*)^d} \left\{ 1 - \left(1 + \frac{\lambda}{\lambda^*} c^{-\alpha} \right) e^{-\lambda c^{-\alpha}} \right\}^{n-d} e^{-\lambda(q + \sum_{i=1}^d x_{i:n}^{-\alpha})}. \quad (3.12)$$

The algorithm for Gibbs sampling procedure consists of the following steps.

Step 1. Start with $j = 1$ and initial value of the parameter as (α_0, λ_0) .

Step 2. Using Metropolis-Hastings procedure, generate $\alpha_j \sim \pi_1(\bullet|\lambda_{j-1})$.

Step 3. Using Metropolis-Hastings procedure, generate $\lambda_j \sim \pi_2(\bullet|\alpha_j)$.

Step 4. Repeat the steps 2 and 3 for all $j = 1, 2, \dots, M$ and obtain $(\alpha_1, \lambda_1), (\alpha_2, \lambda_2), \dots, (\alpha_M, \lambda_M)$.

Step 5. The Bayes estimators of α and λ using the GEL function are respectively given by

$$\hat{\alpha}_{MC} = \left(\frac{1}{M - M_0} \sum_{i=M_0+1}^M \alpha_i^{-w} \right)^{-\frac{1}{w}} \text{ and } \hat{\lambda}_{MC} = \left(\frac{1}{M - M_0} \sum_{i=M_0+1}^M \lambda_i^{-w} \right)^{-\frac{1}{w}}, \quad (3.13)$$

where M_0 is the burn-in period.

The details of the Gibbs sampling procedure along with Metropolis-Hastings algorithm can be seen in Chib and Greenberg (1995) and Saraiva and Suzuki (2017). In order to generate $\alpha_j \sim \pi_1(\bullet|\lambda_{j-1}, X)$, in Step 2, we first generate $\alpha^* \sim N(\alpha_{j-1}, \sigma_\alpha^2)$. Then compute $A_\alpha = \frac{\pi_1(\alpha^*|\lambda_{j-1}, X)f_N(\alpha_{j-1}; \alpha^*, \sigma_\alpha^2)}{\pi_1(\alpha_{j-1}|\lambda_{j-1}, X)f_N(\alpha_{j-1}; \alpha_{j-1}, \sigma_\alpha^2)}$ and consequently compute $\psi_\alpha = \min(1, A_\alpha)$. Next, generate $u_\alpha \sim U(0, 1)$. If $u_\alpha < \psi_\alpha$, then accept α^* , and update $\alpha_j = \alpha^*$, otherwise reject α^* and $\alpha_j = \alpha_{j-1}$. Here σ_α^2 is the variance of α obtained from the inverse of the observed Fisher information matrix. In a very similar manner, we have generated $\lambda_j \sim \pi_2(\bullet|\alpha_{j-1}, X)$ in Step 3. Using this sample, we obtain the approximate $100(1 - \psi)\%$ HPD credible intervals for the parameters α and λ using the method of Chen and Shao (1999).

4 Computational Results

In the previous sections, we have proposed several point and interval estimators, such as the MLEs (obtained via EM-algorithm), Bayes estimators (obtained using both MCMC and Tierney and Kadane's approximation methods), ACIs, **MACIs**, HPD credible intervals for α and λ using type-i hybrid samples from the GIL distribution. We have also considered bootstrap confidence intervals, such as bootstrap-t (boot-t) and bootstrap-p (boot-p) for the associated parameters. The details of the algorithm for computing these two bootstrap confidence intervals have been omitted here for brevity, however interested readers may look at the book by Efron and Tibshirani (1994) for complete understanding of the methods.

It is to be noted that none of the estimators has closed-form expressions. Hence, an analytical comparison of performances among these is not possible. However, for application purposes, it is essential to evaluate their performances, which may be handy in real-life applications. In this section, taking advantage of the superior computational facilities available now-a-days, we would like to compare their performances numerically. In order to evaluate and compare the proposed point estimators, we use two criteria, namely the bias

$$E(\delta - \theta)$$

and the MSE

$$E(\delta - \theta)^2,$$

where δ is an estimator for the parameter θ . The confidence intervals are compared through the CPs, ALs and the PCD values.

In order to compare the estimators numerically, we have generated 20,000 random samples from the GIL distribution with parameters (α, λ) and different choices of n , r ($r \leq n$) and T . A high level of accuracy has been achieved in our simulation in the sense that the standard error of the simulation is very small which is of the order 10^{-3} . The Bayes estimators are computed with respect to the GEL

function using both informative prior (IP) and non-informative prior (NIP). In the case of informative prior (IP), the hyper parameters in the gamma priors have been chosen as $a = 3$, $b = 1$, $p = 5$, $q = 4$. In the case of non-informative prior, we assume that $a = 0$, $b = 0$, $p = 0$, $q = 0$ as there is no prior information available to the practitioner. For Metropolis-Hastings algorithm we consider Normal distribution as the proposal distribution to generate samples for α and λ . MLEs of the parameters are considered as the initial value (α_0, λ_0) accordingly for MCMC method. Further we considered $M = 55000$, $M_0 = 5000$ and 10 as jumping period for MCMC method.

The parameter w in the GEL function describes the shape of the loss function. Moreover, the GEL function behaves like the squared error (SE) loss function for $w = -1$ and like a special case of the weighted squared error loss function for $w = 1$. In order to see the effect of loss function on the estimators, we have conveniently taken $w = -1, -0.5, 0.5, 1.0$ in our simulation study. In these cases the true values of the parameters have been taken as $\alpha = 3$ and $\lambda = 2$. We use the notations $\hat{\alpha}_{TK1}(\hat{\lambda}_{TK1})$, $\hat{\alpha}_{TK2}(\hat{\lambda}_{TK2})$, $\hat{\alpha}_{TK3}(\hat{\lambda}_{TK3})$, and $\hat{\alpha}_{TK4}(\hat{\lambda}_{TK4})$ for the Bayes estimators of $\alpha(\lambda)$ using the Tierney and Kadane's method for the choices of $w = -1, -0.5, 0.5$ and 1.0 respectively in the tables. Similarly, the notations $\hat{\alpha}_{MC1}(\hat{\lambda}_{MC1})$, $\hat{\alpha}_{MC2}(\hat{\lambda}_{MC2})$, $\hat{\alpha}_{MC3}(\hat{\lambda}_{MC3})$, and $\hat{\alpha}_{MC4}(\hat{\lambda}_{MC4})$ for the Bayes estimators of $\alpha(\lambda)$ using the MCMC method for the choices of $w = -1, -0.5, 0.5$ and 1.0 respectively have been used. Further, we have used the notations HPD1 and HPD2 for HPD intervals of parameters using informative and non-informative priors respectively.

The simulation study has been conducted for several combinations of n, r, T , however, for illustration purposes, we present the estimates, MSEs and biases (in the case of point estimation) and intervals, coverage probabilities and average lengths (in the case of interval estimation) for some choices of n, r and T in the tables. In Tables 4.1-4.2, we present the estimates, MSEs and biases of several point estimators of α . Similarly, in Tables 4.3-4.4, the estimates, MSEs and biases of all the proposed point estimators of λ have been presented. In Tables 4.1-4.4, from the fifth column onwards, each cell contains three values, namely the estimates, MSEs and the biases of the estimators, respectively.

The interval estimators of α , such as the ACI, MACI, boot-p confidence interval, boot-t confidence interval and the HPD credible intervals (HPD1 and HPD2) with their CP and AL for the same combinations of n, r and T are presented in Table 4.5. Similarly, the interval estimators with CP and AL of the parameter λ have been presented in Table 4.6. The coverage probabilities and average lengths of all the proposed confidence intervals are computed at $\psi = 0.05$ level of significance for all combinations of n, r , and T . In Tables 4.5-4.6, from the fourth column onwards, each cell contains three values, namely the interval estimates, CPs and the ALs of the confidence intervals, respectively.

In our simulation study regarding the comparison of confidence intervals, we observed that except the ACI of λ , none of the intervals attains the desirable nominal confidence level $1 - \psi$. However all the interval estimators attain a level very close to the nominal level $1 - \psi$, hence it is not advisable to ignore them all. We propose two criteria for choosing the confidence intervals. (i) Fix a lower acceptable threshold for CP, collect those intervals which satisfy this threshold, and then evaluate the CIs in terms of their ALs. (ii) Choose the set of acceptable confidence intervals (which have attained the acceptable lower threshold CP), and then apply the unified criterion of 'Probability Coverage Density' defined by

$$\text{Probability Coverage Density (PCD)} = \frac{\text{Coverage Probability(CP)}}{\text{Average Length(AL)}}$$

to choose the confidence intervals in terms of highest PCD. We refer to Khatun et al. (2020) for the justification on using PCD criteria for comparing the confidence intervals. In Table 4.7, we report the PCD values of all the proposed interval estimators for α and λ with specific combinations of n, r and T . The simulation study has been carried with computer programs through R-software. The program codes will be available to the readers on request upon the authors.

The following conclusions can be drawn from our simulation study as well as the Tables 4.1-4.7, which we describe in the form of some remarks.

Remark 4.1. *The performances of all the point estimators in terms of bias and the MSEs can be described as follows.*

(a) *It is noticed that for fixed values of r and T when n increases, the bias and the MSEs of all the estimators of both α and λ decrease.*

(b) *While estimating the parameter α , the Bayes estimators using Tierney and Kadane's approximation method (denoted as $\hat{\alpha}_{TK4}$ with IP) and MCMC method (denoted as $\hat{\alpha}_{MC4}$ with IP) compete with the MLE in terms of their biases. However, for estimating the parameter λ , the Bayes estimators using Tierney and Kadane's method (denoted as $\hat{\lambda}_{TK4}$ with IP) and MCMC method (denoted as $\hat{\lambda}_{MC4}$ with IP) perform better than the MLE and compete each other in terms of their biases.*

(c) *The Bayes estimators for α using Tierney and Kadane's method (denoted as $\hat{\alpha}_{TK4}$ with IP) and MCMC method (denoted as $\hat{\alpha}_{MC4}$ with IP) compete with each other and have better performance than others in terms of MSE. However, for estimating the parameter λ , the Bayes estimator using Tierney and Kadane's method (denoted as $\hat{\lambda}_{TK1}$ with IP) and MCMC method (denoted as $\hat{\lambda}_{MC1}$ with IP) too compete with each other and have better performance than others in terms of the MSE.*

Remark 4.2. *The performances of all the confidence intervals in terms of CP, AL and PCD can be described as follows.*

(a) *If we fix the nominal level as 0.95 strictly, the performance of ACI in terms of CP is the best and is the only estimator attaining the nominal level for estimating λ . However, none of the confidence intervals attains the nominal level while estimating α .*

(b) *Consider the interval estimation of α . If we choose a slightly lower acceptable threshold for CP, say 0.90 (whereas our target is 0.95), then the CIs in terms of CPs can be ordered (best being the CI with the highest CP) as HPD1, HPD2, ACI, **MACI** then followed by boot-t and boot-p (in fact boot-t and boot-p compete with each other). In terms of AL, the four intervals HPD1, HPD2, ACI and the **MACI** compete with each other in terms of having the shortest length, followed by either boot-t or boot-p.*

(c) *Consider the interval estimation of λ . If we fix the nominal level slightly lower, say 0.90, then the interval estimators can be ordered in terms of CPs as ACI is the best, next **MACI**, HPD1, HPD2, then followed by boot-t and boot-p. In terms of shortest AL, the confidence intervals are ordered as HPD1 being the shortest length, then HPD2, ACI, **MACI**, followed by boot-t and boot-p.*

(d) *If we choose the intervals in terms of both CP and AL, we observe that ACI, **MACI**, HPD1 and HPD2 perform better than the other two intervals. In order to get the best interval, we use the PCD criteria. In terms of PCD values, we observe that the interval HPD1 has the best performance, followed by HPD2, ACI, and **MACI** for estimating λ . In estimating α , the interval HPD1 and HPD2 have the best performance, except for a few combinations of n , r and T where ACI and **MACI** have slightly better performance than HPD1 and HPD2.*

Table 4.1: The Average Estimates, MSEs and Biases of Various Estimators of $\alpha = 3$ for Different Choices of n , T and r

n	T	r	Prior	$\hat{\alpha}_{ML}$	$\hat{\alpha}_{TK1}$	$\hat{\alpha}_{TK2}$	$\hat{\alpha}_{TK3}$	$\hat{\alpha}_{TK4}$	$\hat{\alpha}_{MC1}$	$\hat{\alpha}_{MC2}$	$\hat{\alpha}_{MC3}$	$\hat{\alpha}_{MC4}$
40	2	30	IP	2.988	3.105	3.090	3.060	3.045	3.105	3.090	3.060	3.045
				0.205	0.199	0.194	0.186	0.183	0.199	0.195	0.187	0.183
				-0.012	0.105	0.090	0.060	0.045	0.105	0.090	0.060	0.045
		40	NIP		3.118	3.103	3.071	3.056	3.118	3.103	3.071	3.056
					0.227	0.221	0.212	0.207	0.228	0.223	0.213	0.209
					0.118	0.103	0.071	0.056	0.118	0.103	0.071	0.056
	4	30	IP	3.041	3.073	3.060	3.032	3.019	3.072	3.059	3.031	3.018
				0.201	0.168	0.165	0.160	0.159	0.169	0.167	0.162	0.160
				0.041	0.073	0.060	0.032	0.019	0.072	0.059	0.031	0.018
		40	NIP		3.079	3.065	3.037	3.023	3.079	3.065	3.037	3.023
					0.187	0.184	0.179	0.177	0.187	0.184	0.179	0.177
					0.079	0.065	0.037	0.023	0.079	0.065	0.037	0.023
50	2	30	IP	2.993	3.109	3.094	3.064	3.049	3.109	3.094	3.064	3.049
				0.197	0.197	0.192	0.184	0.180	0.197	0.192	0.184	0.180
				-0.007	0.109	0.094	0.064	0.049	0.109	0.094	0.064	0.049
		40	NIP		3.135	3.119	3.088	3.072	3.135	3.120	3.088	3.072
					0.230	0.224	0.213	0.209	0.231	0.225	0.214	0.210
					0.135	0.119	0.088	0.072	0.135	0.120	0.088	0.072
		50	IP	3.097	3.050	3.038	3.014	3.001	3.068	3.056	3.032	3.020
				0.173	0.147	0.144	0.140	0.138	0.150	0.147	0.143	0.141
				0.097	0.069	0.057	0.033	0.021	0.068	0.056	0.032	0.020
		60	NIP		3.076	3.064	3.039	3.026	3.101	3.088	3.063	3.050
					0.177	0.174	0.167	0.165	0.179	0.175	0.169	0.166
					0.100	0.088	0.063	0.050	0.101	0.088	0.063	0.050
	4	40	IP	3.006	3.084	3.072	3.050	3.038	3.083	3.072	3.049	3.038
				0.150	0.143	0.140	0.136	0.134	0.143	0.141	0.136	0.134
				0.006	0.084	0.072	0.050	0.038	0.083	0.072	0.049	0.038
		50	NIP		3.095	3.083	3.060	3.048	3.095	3.084	3.060	3.048
					0.164	0.161	0.155	0.153	0.164	0.161	0.156	0.153
					0.095	0.083	0.060	0.048	0.095	0.084	0.060	0.048
50	4	40	IP	3.034	3.072	3.061	3.040	3.029	3.072	3.061	3.040	3.029
				0.147	0.133	0.131	0.127	0.126	0.133	0.131	0.128	0.126
				0.034	0.072	0.061	0.040	0.029	0.072	0.061	0.040	0.029
		50	NIP		3.078	3.067	3.045	3.034	3.078	3.067	3.044	3.033
					0.152	0.149	0.145	0.144	0.152	0.150	0.146	0.144
					0.078	0.067	0.045	0.034	0.078	0.067	0.044	0.033
		60	IP	2.999	3.071	3.059	3.037	3.025	3.071	3.060	3.037	3.026
				0.139	0.136	0.133	0.129	0.128	0.136	0.134	0.130	0.128
				-0.001	0.071	0.059	0.037	0.025	0.071	0.060	0.037	0.026
		70	NIP		3.094	3.082	3.059	3.047	3.094	3.083	3.059	3.048
					0.164	0.161	0.155	0.153	0.165	0.161	0.156	0.153
					0.094	0.082	0.059	0.047	0.094	0.083	0.059	0.048
50	4	50	IP	3.085	3.065	3.056	3.036	3.026	3.066	3.056	3.037	3.027
				0.138	0.124	0.122	0.119	0.117	0.124	0.122	0.119	0.118
				0.085	0.065	0.056	0.036	0.026	0.066	0.056	0.037	0.027
		60	NIP		3.079	3.069	3.049	3.039	3.079	3.069	3.049	3.039
					0.134	0.132	0.128	0.126	0.134	0.132	0.128	0.126
					0.079	0.069	0.049	0.039	0.079	0.069	0.049	0.039

Table 4.2: The Average Estimates, MSEs and Biases of Various Estimators of $\alpha = 3$ for Different Choices of n , T and r

n	T	r	Prior	$\hat{\alpha}_{ML}$	$\hat{\alpha}_{TK1}$	$\hat{\alpha}_{TK2}$	$\hat{\alpha}_{TK3}$	$\hat{\alpha}_{TK4}$	$\hat{\alpha}_{MC1}$	$\hat{\alpha}_{MC2}$	$\hat{\alpha}_{MC3}$	$\hat{\alpha}_{MC4}$
60	2	50	IP	2.999	3.060	3.050	3.032	3.023	3.060	3.051	3.033	3.023
				0.126	0.115	0.114	0.111	0.110	0.116	0.114	0.111	0.110
				-0.001	0.060	0.050	0.032	0.023	0.060	0.051	0.033	0.023
		60	NIP		3.068	3.058	3.040	3.030	3.068	3.059	3.040	3.031
					0.127	0.125	0.122	0.120	0.127	0.125	0.122	0.121
					0.068	0.058	0.040	0.030	0.068	0.059	0.040	0.031
	4	50	IP	3.007	3.051	3.042	3.024	3.015	3.051	3.042	3.024	3.015
				0.128	0.108	0.107	0.105	0.104	0.108	0.107	0.105	0.104
				0.007	0.051	0.042	0.024	0.015	0.051	0.042	0.024	0.015
		60	NIP		3.059	3.049	3.031	3.022	3.059	3.050	3.032	3.022
					0.119	0.117	0.115	0.114	0.119	0.118	0.115	0.114
					0.059	0.049	0.031	0.022	0.059	0.050	0.032	0.022
80	2	50	IP	3.014	3.063	3.054	3.036	3.027	3.064	3.054	3.036	3.027
				0.115	0.113	0.111	0.108	0.107	0.114	0.112	0.109	0.108
				0.014	0.063	0.054	0.036	0.027	0.064	0.054	0.036	0.027
		60	NIP		3.071	3.062	3.044	3.034	3.072	3.063	3.044	3.035
					0.124	0.122	0.119	0.117	0.125	0.123	0.119	0.118
					0.071	0.062	0.044	0.034	0.072	0.063	0.044	0.035
		70	IP	3.077	3.061	3.053	3.037	3.029	3.061	3.053	3.037	3.029
				0.111	0.100	0.098	0.096	0.095	0.100	0.099	0.096	0.095
				0.077	0.061	0.053	0.037	0.029	0.061	0.053	0.037	0.029
		80	NIP		3.061	3.052	3.036	3.028	3.061	3.053	3.037	3.028
					0.110	0.108	0.106	0.105	0.110	0.109	0.106	0.105
					0.061	0.052	0.036	0.028	0.061	0.053	0.037	0.028
	4	50	IP	2.986	3.038	3.031	3.018	3.011	3.038	3.031	3.018	3.011
				0.097	0.085	0.084	0.083	0.082	0.085	0.084	0.083	0.082
				-0.014	0.038	0.031	0.018	0.011	0.038	0.031	0.018	0.011
		60	NIP		3.046	3.039	3.026	3.019	3.046	3.039	3.026	3.019
					0.086	0.085	0.084	0.083	0.086	0.086	0.084	0.084
					0.046	0.039	0.026	0.019	0.046	0.039	0.026	0.019
		70	IP	2.988	3.037	3.030	3.017	3.010	3.037	3.030	3.017	3.010
				0.097	0.083	0.082	0.081	0.081	0.083	0.083	0.081	0.081
				-0.012	0.037	0.030	0.017	0.010	0.037	0.030	0.017	0.010
		80	NIP		3.044	3.037	3.024	3.017	3.044	3.037	3.024	3.017
					0.085	0.084	0.082	0.082	0.085	0.084	0.083	0.082
					0.044	0.037	0.024	0.017	0.044	0.037	0.024	0.017
	8	50	IP	3.016	3.043	3.036	3.023	3.017	3.043	3.036	3.023	3.017
				0.084	0.082	0.081	0.080	0.079	0.082	0.081	0.080	0.079
				0.016	0.043	0.036	0.023	0.017	0.043	0.036	0.023	0.017
		60	NIP		3.048	3.042	3.029	3.022	3.048	3.042	3.029	3.022
					0.085	0.084	0.082	0.082	0.085	0.084	0.082	0.082
					0.048	0.042	0.029	0.022	0.048	0.042	0.029	0.022
	16	50	IP	3.049	3.039	3.033	3.021	3.015	3.039	3.033	3.021	3.015
				0.079	0.073	0.073	0.071	0.071	0.074	0.073	0.072	0.071
				0.049	0.039	0.033	0.021	0.015	0.039	0.033	0.021	0.015
		60	NIP		3.042	3.036	3.024	3.018	3.042	3.036	3.024	3.018
					0.077	0.077	0.075	0.075	0.078	0.077	0.076	0.075
					0.042	0.036	0.024	0.018	0.042	0.036	0.024	0.018

Table 4.3: The Average Estimates, MSEs and Biases of Various Estimators of $\lambda = 2$ for Different Choices of n , T and r

n	T	r	Prior	$\hat{\lambda}_{ML}$	$\hat{\lambda}_{TK1}$	$\hat{\lambda}_{TK2}$	$\hat{\lambda}_{TK3}$	$\hat{\lambda}_{TK4}$	$\hat{\lambda}_{MC1}$	$\hat{\lambda}_{MC2}$	$\hat{\lambda}_{MC3}$	$\hat{\lambda}_{MC4}$
40	2	30	IP	2.041	1.944	1.936	1.922	1.914	1.943	1.936	1.921	1.914
				0.073	0.056	0.056	0.057	0.058	0.056	0.056	0.057	0.058
				0.041	-0.056	-0.064	-0.078	-0.086	-0.057	-0.064	-0.079	-0.086
		40	NIP		2.047	2.039	2.022	2.013	2.047	2.039	2.022	2.013
					0.083	0.081	0.078	0.077	0.083	0.081	0.079	0.077
					0.047	0.039	0.022	0.013	0.047	0.039	0.022	0.013
	4	30	IP	2.055	1.954	1.946	1.932	1.924	1.953	1.946	1.931	1.924
				0.080	0.055	0.055	0.056	0.056	0.055	0.055	0.056	0.057
				0.055	-0.046	-0.054	-0.068	-0.076	-0.047	-0.054	-0.069	-0.076
		40	NIP		2.044	2.036	2.019	2.011	2.044	2.036	2.019	2.011
					0.078	0.076	0.074	0.073	0.078	0.076	0.074	0.073
					0.044	0.036	0.019	0.011	0.044	0.036	0.019	0.011
50	2	30	IP	2.039	1.942	1.934	1.920	1.912	1.941	1.934	1.920	1.912
				0.070	0.054	0.054	0.055	0.056	0.054	0.054	0.055	0.056
				0.039	-0.058	-0.066	-0.080	-0.088	-0.059	-0.066	-0.080	-0.088
		40	NIP		2.042	2.033	2.016	2.008	2.042	2.033	2.016	2.008
					0.084	0.082	0.079	0.078	0.084	0.082	0.080	0.079
					0.042	0.033	0.016	0.008	0.042	0.033	0.016	0.008
		50	IP	2.052	1.935	1.928	1.914	1.906	1.948	1.940	1.926	1.919
				0.078	0.051	0.052	0.053	0.053	0.052	0.053	0.054	0.054
				0.052	-0.051	-0.059	-0.073	-0.080	-0.052	-0.060	-0.074	-0.081
		40	NIP		2.022	2.013	1.997	1.988	2.038	2.030	2.013	2.005
					0.075	0.074	0.072	0.071	0.076	0.075	0.073	0.072
					0.038	0.029	0.013	0.004	0.038	0.030	0.013	0.005
	4	40	IP	2.043	1.964	1.958	1.946	1.940	1.964	1.958	1.946	1.940
				0.060	0.046	0.046	0.046	0.047	0.046	0.046	0.046	0.047
				0.043	-0.036	-0.042	-0.054	-0.060	-0.036	-0.042	-0.054	-0.060
		50	NIP		2.033	2.026	2.013	2.006	2.033	2.026	2.013	2.006
					0.062	0.061	0.060	0.059	0.062	0.061	0.060	0.059
					0.033	0.026	0.013	0.006	0.033	0.026	0.013	0.006
50	2	40	IP	2.036	1.957	1.951	1.939	1.933	1.957	1.951	1.939	1.933
				0.058	0.045	0.045	0.046	0.046	0.045	0.045	0.046	0.046
				0.036	-0.043	-0.049	-0.061	-0.067	-0.043	-0.049	-0.061	-0.067
		50	NIP		2.030	2.023	2.010	2.003	2.030	2.023	2.010	2.003
					0.061	0.060	0.058	0.058	0.061	0.060	0.059	0.058
					0.030	0.023	0.010	0.003	0.030	0.023	0.010	0.003
	4	40	IP	2.034	1.955	1.949	1.937	1.931	1.955	1.949	1.937	1.931
				0.057	0.045	0.046	0.046	0.047	0.045	0.046	0.046	0.047
				0.034	-0.045	-0.051	-0.063	-0.069	-0.045	-0.051	-0.063	-0.069
		50	NIP		2.037	2.030	2.017	2.010	2.037	2.031	2.017	2.010
					0.061	0.060	0.058	0.057	0.061	0.060	0.058	0.058
					0.037	0.030	0.017	0.010	0.037	0.031	0.017	0.010
50	2	40	IP	2.037	1.955	1.949	1.937	1.931	1.955	1.949	1.937	1.931
				0.060	0.044	0.044	0.045	0.046	0.044	0.045	0.045	0.046
				0.037	-0.045	-0.051	-0.063	-0.069	-0.045	-0.051	-0.063	-0.069
		50	NIP		2.034	2.027	2.014	2.007	2.034	2.028	2.014	2.008
					0.059	0.058	0.057	0.056	0.060	0.059	0.057	0.056
					0.034	0.027	0.014	0.007	0.034	0.028	0.014	0.008

Table 4.4: The Average Estimates, MSEs and Biases of Various Estimators of $\lambda = 2$ for Different Choices of n , T and r

n	T	r	Prior	$\hat{\lambda}_{ML}$	$\hat{\lambda}_{TK1}$	$\hat{\lambda}_{TK2}$	$\hat{\lambda}_{TK3}$	$\hat{\lambda}_{TK4}$	$\hat{\lambda}_{MC1}$	$\hat{\lambda}_{MC2}$	$\hat{\lambda}_{MC3}$	$\hat{\lambda}_{MC4}$
60	2	50	IP	2.028	1.963	1.958	1.948	1.943	1.963	1.958	1.948	1.943
				0.047	0.038	0.039	0.039	0.039	0.039	0.039	0.039	0.039
				0.028	-0.037	-0.042	-0.052	-0.057	-0.037	-0.042	-0.052	-0.057
		60	NIP		2.028	2.023	2.012	2.006	2.028	2.023	2.012	2.006
					0.049	0.049	0.048	0.047	0.050	0.049	0.048	0.047
					0.028	0.023	0.012	0.006	0.028	0.023	0.012	0.006
	4	50	IP	2.029	1.963	1.958	1.948	1.943	1.964	1.959	1.949	1.943
				0.047	0.038	0.038	0.039	0.039	0.038	0.038	0.039	0.039
				0.029	-0.037	-0.042	-0.052	-0.057	-0.036	-0.041	-0.051	-0.057
		60	NIP		2.029	2.023	2.012	2.007	2.029	2.023	2.012	2.007
					0.049	0.049	0.047	0.047	0.049	0.049	0.048	0.047
					0.029	0.023	0.012	0.007	0.029	0.023	0.012	0.007
80	2	40	IP	2.029	1.963	1.958	1.948	1.943	1.963	1.958	1.948	1.943
				0.048	0.038	0.039	0.039	0.039	0.039	0.039	0.039	0.040
				0.029	-0.037	-0.042	-0.052	-0.057	-0.037	-0.042	-0.052	-0.057
		50	NIP		2.028	2.023	2.012	2.006	2.028	2.023	2.012	2.006
					0.049	0.049	0.048	0.047	0.050	0.049	0.048	0.048
					0.028	0.023	0.012	0.006	0.028	0.023	0.012	0.006
	4	40	IP	2.034	1.965	1.960	1.950	1.945	1.965	1.960	1.950	1.945
				0.050	0.038	0.039	0.039	0.039	0.039	0.039	0.039	0.040
				0.034	-0.035	-0.040	-0.050	-0.055	-0.035	-0.040	-0.050	-0.055
		50	NIP		2.029	2.024	2.012	2.007	2.029	2.024	2.013	2.007
					0.049	0.048	0.047	0.047	0.049	0.048	0.047	0.047
					0.029	0.024	0.012	0.007	0.029	0.024	0.013	0.007
80	2	40	IP	2.043	1.964	1.958	1.946	1.940	1.964	1.958	1.946	1.940
				0.060	0.046	0.046	0.046	0.047	0.046	0.046	0.046	0.047
				0.043	-0.036	-0.042	-0.054	-0.060	-0.036	-0.042	-0.054	-0.060
		50	NIP		2.033	2.026	2.013	2.006	2.033	2.026	2.013	2.006
					0.062	0.061	0.060	0.059	0.062	0.061	0.060	0.059
					0.033	0.026	0.013	0.006	0.033	0.026	0.013	0.006
	4	40	IP	2.036	1.957	1.951	1.939	1.933	1.957	1.951	1.939	1.933
				0.058	0.045	0.045	0.046	0.046	0.045	0.045	0.046	0.046
				0.036	-0.043	-0.049	-0.061	-0.067	-0.043	-0.049	-0.061	-0.067
		50	NIP		2.030	2.023	2.010	2.003	2.030	2.023	2.010	2.003
					0.061	0.060	0.058	0.058	0.061	0.060	0.059	0.058
					0.030	0.023	0.010	0.003	0.030	0.023	0.010	0.003
80	4	40	IP	2.034	1.955	1.949	1.937	1.931	1.955	1.949	1.937	1.931
				0.057	0.045	0.046	0.046	0.047	0.045	0.046	0.046	0.047
				0.034	-0.045	-0.051	-0.063	-0.069	-0.045	-0.051	-0.063	-0.069
		50	NIP		2.037	2.030	2.017	2.010	2.037	2.031	2.017	2.010
					0.061	0.060	0.058	0.057	0.061	0.060	0.058	0.058
					0.037	0.030	0.017	0.010	0.037	0.031	0.017	0.010
80	4	40	IP	2.037	1.955	1.949	1.937	1.931	1.955	1.949	1.937	1.931
				0.060	0.044	0.044	0.045	0.046	0.044	0.045	0.045	0.046
				0.037	-0.045	-0.051	-0.063	-0.069	-0.045	-0.051	-0.063	-0.069
		50	NIP		2.034	2.027	2.014	2.007	2.034	2.028	2.014	2.008
					0.059	0.058	0.057	0.056	0.060	0.059	0.057	0.056
					0.034	0.027	0.014	0.007	0.034	0.028	0.014	0.008

Table 4.5: The Confidence Intervals with CPs and ALs of $\alpha = 3$ for Different Choices of n , T and r

n	T	r	ACI	MACI	Boot-p	Boot-t	HPD1	HPD2
40	2	30	(2.206,3.769)	(2.302,3.884)	(2.375,4.161)	(2.376,4.164)	(2.297,3.935)	(2.292,3.965)
			0.914	0.927	0.922	0.923	0.946	0.943
			1.563	1.582	1.786	1.788	1.638	1.673
		40	(2.269,3.813)	(2.360,3.922)	(2.419,4.094)	(2.419,4.093)	(2.309,3.860)	(2.301,3.879)
			0.915	0.915	0.925	0.926	0.947	0.943
			1.544	1.562	1.675	1.674	1.551	1.578
	4	30	(2.211,3.776)	(2.301,3.881)	(2.393,4.176)	(2.391,4.172)	(2.301,3.939)	(2.305,3.987)
			0.919	0.927	0.930	0.931	0.949	0.944
			1.565	1.580	1.783	1.781	1.638	1.682
		40	(2.334,3.859)	(2.429,3.975)	(2.510,4.125)	(2.510,4.125)	(2.349,3.813)	(2.362,3.863)
			0.951	0.936	0.931	0.932	0.959	0.946
			1.525	1.546	1.615	1.615	1.464	1.501
50	2	40	(2.314,3.699)	(2.384,3.778)	(2.443,3.976)	(2.445,3.977)	(2.380,3.805)	(2.380,3.830)
			0.923	0.919	0.928	0.928	0.951	0.943
			1.385	1.394	1.533	1.532	1.425	1.450
		50	(2.345,3.723)	(2.408,3.793)	(2.469,3.939)	(2.468,3.937)	(2.389,3.774)	(2.383,3.788)
			0.924	0.916	0.928	0.929	0.953	0.944
			1.378	1.385	1.470	1.469	1.385	1.405
	4	40	(2.309,3.688)	(2.391,3.787)	(2.456,3.966)	(2.454,3.967)	(2.372,3.788)	(2.382,3.825)
			0.935	0.933	0.936	0.937	0.952	0.945
			1.379	1.396	1.510	1.513	1.460	1.443
		50	(2.406,3.765)	(2.476,3.847)	(2.550,3.967)	(2.551,3.969)	(2.418,3.729)	(2.420,3.752)
			0.945	0.936	0.946	0.948	0.954	0.951
			1.359	1.371	1.417	1.418	1.311	1.332
60	2	50	(2.373,3.625)	(2.432,3.693)	(2.480,3.834)	(2.480,3.832)	(2.429,3.704)	(2.428,3.721)
			0.921	0.921	0.923	0.924	0.950	0.933
			1.252	1.261	1.354	1.352	1.275	1.293
		60	(2.383,3.630)	(2.445,3.703)	(2.486,3.805)	(2.487,3.804)	(2.427,3.685)	(2.428,3.703)
			0.922	0.918	0.924	0.925	0.951	0.934
			1.247	1.257	1.319	1.317	1.258	1.275
	4	50	(2.386,3.641)	(2.450,3.715)	(2.512,3.853)	(2.511,3.851)	(2.436,3.704)	(2.436,3.721)
			0.941	0.943	0.943	0.943	0.955	0.950
			1.255	1.265	1.341	1.340	1.268	1.285
		60	(2.458,3.696)	(2.511,3.754)	(2.581,3.859)	(2.582,3.858)	(2.468,3.667)	(2.464,3.671)
			0.949	0.941	0.949	0.950	0.958	0.953
			1.238	1.243	1.278	1.276	1.199	1.207
80	2	70	(2.449,3.524)	(2.500,3.583)	(2.530,3.661)	(2.528,3.661)	(2.499,3.588)	(2.502,3.602)
			0.921	0.924	0.929	0.929	0.945	0.944
			1.075	1.083	1.131	1.133	1.089	1.100
		80	(2.451,3.525)	(2.502,3.584)	(2.531,3.650)	(2.529,3.651)	(2.499,3.586)	(2.501,3.598)
			0.922	0.919	0.926	0.928	0.946	0.944
			1.074	1.082	1.119	1.122	1.087	1.097
	4	70	(2.477,3.555)	(2.528,3.613)	(2.575,3.694)	(2.574,3.695)	(2.511,3.585)	(2.512,3.596)
			0.941	0.942	0.943	0.943	0.957	0.951
			1.078	1.085	1.119	1.121	1.074	1.084
		80	(2.518,3.580)	(2.569,3.639)	(2.615,3.694)	(2.612,3.695)	(2.530,3.559)	(2.527,3.565)
			0.947	0.941	0.948	0.949	0.960	0.955
			1.062	1.070	1.079	1.083	1.029	1.038

Table 4.6: The Confidence Intervals with CPs and ALs of $\lambda = 2$ for Different Choices of n , T and r

n	T	r	ACI	MACI	Boot-p	Boot-t	HPD1	HPD2
40	2	30	(1.527,2.555) 0.959 1.028	(1.587,2.627) 0.950 1.039	(1.613,2.743) 0.933 1.130	(1.614,2.745) 0.930 1.131	(1.499,2.407) 0.932 0.908	(1.556,2.558) 0.932 1.002
		40	(1.537,2.572) 0.956 1.035	(1.591,2.631) 0.952 1.040	(1.633,2.765) 0.920 1.132	(1.632,2.761) 0.919 1.129	(1.505,2.420) 0.939 0.915	(1.555,2.555) 0.941 1.000
		30	(1.526,2.552) 0.960 1.026	(1.588,2.627) 0.951 1.039	(1.612,2.743) 0.939 1.131	(1.611,2.742) 0.938 1.131	(1.497,2.406) 0.933 0.909	(1.553,2.553) 0.933 1.000
	4	40	(1.536,2.568) 0.957 1.032	(1.591,2.632) 0.946 1.041	(1.631,2.751) 0.925 1.120	(1.631,2.750) 0.924 1.119	(1.503,2.412) 0.945 0.909	(1.552,2.546) 0.942 0.994
		40	(1.584,2.502) 0.953 0.918	(1.625,2.545) 0.953 0.921	(1.659,2.643) 0.932 0.984	(1.660,2.641) 0.935 0.981	(1.558,2.385) 0.935 0.827	(1.598,2.486) 0.936 0.888
	2	50	(1.579,2.493) 0.955 0.914	(1.626,2.547) 0.951 0.921	(1.656,2.629) 0.929 0.973	(1.657,2.629) 0.932 0.972	(1.553,2.377) 0.936 0.824	(1.596,2.482) 0.936 0.886
		40	(1.578,2.491) 0.957 0.913	(1.626,2.548) 0.951 0.922	(1.650,2.629) 0.930 0.979	(1.652,2.628) 0.930 0.976	(1.550,2.374) 0.930 0.824	(1.601,2.491) 0.938 0.890
		50	(1.580,2.494) 0.953 0.914	(1.629,2.553) 0.941 0.924	(1.658,2.626) 0.928 0.968	(1.657,2.625) 0.925 0.968	(1.551,2.374) 0.939 0.823	(1.599,2.486) 0.943 0.887
60	2	50	(1.614,2.443) 0.956 0.829	(1.652,2.487) 0.949 0.835	(1.678,2.550) 0.935 0.872	(1.678,2.551) 0.934 0.873	(1.591,2.349) 0.941 0.758	(1.632,2.440) 0.942 0.808
		60	(1.614,2.444) 0.957 0.830	(1.655,2.491) 0.952 0.836	(1.680,2.550) 0.936 0.870	(1.680,2.550) 0.934 0.870	(1.591,2.350) 0.940 0.759	(1.633,2.440) 0.943 0.807
		50	(1.614,2.444) 0.955 0.830	(1.655,2.491) 0.952 0.837	(1.679,2.551) 0.934 0.872	(1.679,2.552) 0.932 0.873	(1.591,2.349) 0.942 0.758	(1.632,2.439) 0.942 0.807
	4	60	(1.618,2.450) 0.951 0.832	(1.656,2.494) 0.945 0.837	(1.684,2.554) 0.929 0.870	(1.685,2.554) 0.928 0.869	(1.593,2.351) 0.936 0.758	(1.634,2.440) 0.945 0.806
		70	(1.660,2.372) 0.956 0.712	(1.691,2.408) 0.951 0.717	(1.711,2.445) 0.938 0.734	(1.711,2.444) 0.936 0.733	(1.642,2.304) 0.939 0.662	(1.677,2.373) 0.936 0.696
	2	80	(1.660,2.373) 0.955 0.713	(1.696,2.415) 0.951 0.719	(1.712,2.445) 0.937 0.733	(1.711,2.444) 0.936 0.733	(1.642,2.304) 0.938 0.662	(1.678,2.373) 0.935 0.695
		70	(1.661,2.374) 0.954 0.713	(1.695,2.414) 0.949 0.719	(1.712,2.446) 0.935 0.734	(1.711,2.445) 0.934 0.734	(1.642,2.304) 0.936 0.662	(1.680,2.376) 0.943 0.696
		80	(1.662,2.375) 0.950 0.713	(1.695,2.414) 0.948 0.719	(1.714,2.445) 0.934 0.731	(1.713,2.445) 0.932 0.732	(1.642,2.304) 0.939 0.662	(1.681,2.377) 0.944 0.696

Table 4.7: PCD Values for Various CIs of $\alpha = 3$, and $\lambda = 2$ with Different Choices of n , T and r

n	T	r	α						λ					
			ACI	MACI	Boot-p	Boot-t	HPD1	HPD2	ACI	MACI	Boot-p	Boot-t	HPD1	HPD2
40	2	30	0.585	0.586	0.516	0.516	0.578	0.564	0.933	0.914	0.826	0.822	1.026	0.930
		40	0.593	0.586	0.552	0.553	0.611	0.598	0.924	0.913	0.813	0.814	1.026	0.941
	4	30	0.587	0.586	0.522	0.523	0.579	0.561	0.936	0.915	0.830	0.829	1.026	0.933
		40	0.624	0.605	0.576	0.577	0.655	0.630	0.927	0.909	0.826	0.826	1.040	0.948
50	2	40	0.666	0.660	0.605	0.606	0.667	0.650	1.038	1.035	0.947	0.953	1.131	1.054
		50	0.671	0.661	0.631	0.632	0.688	0.672	1.045	1.032	0.955	0.959	1.136	1.056
	4	40	0.678	0.669	0.620	0.619	0.652	0.655	1.048	1.032	0.950	0.953	1.129	1.054
		50	0.695	0.683	0.668	0.669	0.728	0.714	1.043	1.019	0.959	0.956	1.141	1.063
60	2	50	0.736	0.730	0.682	0.683	0.745	0.722	1.153	1.137	1.072	1.070	1.241	1.166
		60	0.739	0.730	0.701	0.702	0.756	0.733	1.153	1.139	1.076	1.074	1.238	1.169
	4	50	0.750	0.745	0.703	0.704	0.753	0.739	1.151	1.139	1.071	1.068	1.243	1.167
		60	0.767	0.757	0.743	0.745	0.799	0.790	1.143	1.129	1.068	1.068	1.235	1.172
80	2	70	0.857	0.853	0.821	0.820	0.868	0.858	1.343	1.327	1.278	1.277	1.418	1.345
		80	0.858	0.854	0.828	0.827	0.870	0.861	1.339	1.322	1.278	1.277	1.417	1.345
	4	70	0.873	0.868	0.843	0.841	0.891	0.877	1.338	1.320	1.274	1.272	1.414	1.355
		80	0.892	0.880	0.879	0.876	0.933	0.920	1.332	1.319	1.278	1.273	1.418	1.356

5 Bayesian Prediction

In this section, we discuss the prediction for censored observation or observation from the future sample. Basically, two types of predictions are discussed in the literature, one-sample prediction and two-sample prediction for predicting future order statistics. Below we discuss these two types of prediction procedures for obtaining future samples under type-i hybrid censoring scheme.

5.1 One-Sample Prediction

Firstly, we consider that $\underline{X} = (X_{1:n}, X_{2:n}, \dots, X_{d:n})$ be an informative (observed) type-i hybrid censored sample of size d from the GIL distribution. Further, we assume that $\underline{Y} = (X_{d+1}, X_{d+2:n}, \dots, X_{n:n})$ denotes ordered lifetimes of the remaining $n - d$ censored units. Under one-sample prediction, we usually predict the lifetime of the s -th censored order statistic $Y_s = X_{s:n}$ (for $d < s \leq n$). We noticed that, the conditional density function of Y_s given type-i hybrid censored sample $\underline{X}$, is of the form

$$f_1(y_s | \underline{x}, \alpha, \lambda) = K \sum_{j=0}^{s-d-1} (-1)^{s-d-1-j} \binom{s-d-1}{j} (1 - F(c))^{j-n+d} (1 - F(y_s))^{n-d-1-j} f(y_s), \quad y_s \geq c,$$

where $K = (s - d) \binom{n-d}{s-d}$. The corresponding one-sample survival function obtained as

$$S_1(t | \underline{x}, \alpha, \lambda) = \frac{P(Y_s > t | \underline{x}, \alpha, \lambda)}{P(Y_s > c | \underline{x}, \alpha, \lambda)}$$

where

$$P(Y_s > t | \underline{x}, \alpha, \lambda) = K \sum_{j=0}^{s-d-1} (-1)^{s-d-1-j} \binom{s-d-1}{j} (1 - F(c))^{j-n+d} \frac{(1 - F(t))^{n-d-j}}{n-d-j}$$

$$P(Y_s > c | \underline{x}, \alpha, \lambda) = K \sum_{j=0}^{s-d-1} (-1)^{s-d-1-j} \binom{s-d-1}{j} \frac{1}{n-d-j}.$$

Therefore, the associated one-sample posterior predictive density function and the corresponding posterior predictive survival function are respectively obtained as

$$f_1^*(y_s | \underline{x}) = \int_0^\infty \int_0^\infty f_1(y_s | \underline{x}, \alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda$$

and

$$S_1^*(t | \underline{x}) = \int_0^\infty \int_0^\infty S_1(t | \underline{x}, \alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda.$$

The s -th censored order statistic Y_s can be estimated under GEL function as

$$\hat{y}_s = \left\{ \int_c^\infty y^{-w} f_1^*(y | \underline{x}) dy \right\}^{-\frac{1}{w}} = \left\{ \int_0^\infty \int_0^\infty Q_1(c, \alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda \right\}^{-\frac{1}{w}}, \quad (5.1)$$

where

$$\begin{aligned} Q_1(c, \alpha, \lambda) &= \int_c^\infty y^{-w} f_1(y | \underline{x}, \alpha, \lambda) dy \\ &= K \sum_{j=0}^{s-d-1} (-1)^{s-d-1-j} \binom{s-d-1}{j} (1 - F(c))^{j-n+d} \int_c^\infty y^{-w} (1 - F(y))^{n-d-1-j} f(y) dy \\ &= K \sum_{j=0}^{s-d-1} (-1)^{s-d-1-j} \binom{s-d-1}{j} (1 - F(c))^{j-n+d} \int_{F(c)}^1 \{F^{-1}(u)\}^{-w} (1 - u)^{n-d-1-j} du. \end{aligned}$$

It seems difficult to evaluate the above integral analytically as the function $F^{-1}(\cdot)$ does not have a closed-form expression. We use the quantile function for GIL distribution which is defined in R-package ‘Lindley’ to evaluate the integration numerically. Further, it is not easy to evaluate the integrations given in (5.1). Hence, we use the posterior samples of α and λ obtained through the Gibbs sampler to approximate the integrations numerically. So, we have the approximate value as

$$\hat{y}_s = \left\{ \frac{1}{M - M_0} \sum_{i=M_0+1}^M Q_1(c, \alpha_i, \lambda_i) \right\}^{-\frac{1}{w}}. \quad (5.2)$$

The $100(1 - \psi)\%$ equal-tailed credible predictive interval (L, U) for the censored observation y_s can be derived by solving the two equations

$$S_1^*(L | \underline{x}) = 1 - \frac{\psi}{2}, \quad S_1^*(U | \underline{x}) = \frac{\psi}{2}. \quad (5.3)$$

The $100(1 - \psi)\%$ HPD predictive interval for the censored observation y_s can be derived by solving the equations

$$S_1^*(L_j | \underline{x}) = 1 - \frac{j}{M - M_0}, \quad S_1^*(U_j | \underline{x}) = 1 - \frac{j + [(1 - \psi)(M - M_0)]}{M - M_0}, \quad (5.4)$$

for $j = 1, 2, \dots, ((M - M_0) - [(1 - \psi)(M - M_0)])$ with $[p]$ denoting the greatest integer less than or equal to p and then choosing the interval (L_{j^*}, U_{j^*}) with minimum length among $\{(L_j, U_j), j = 1, 2, \dots, ((M - M_0) - [(1 - \psi)(M - M_0)])\}$.

5.2 Two-Sample Prediction

In two-sample prediction, we wish to predict the observations of a future sample $\underline{Z} = (Z_1, Z_2, \dots, Z_\nu)$ draw from the GIL distribution, which is independent from the observed type-i hybrid sample $\underline{X}$. If $Z_{(1)}, Z_{(2)}, \dots, Z_{(\nu)}$ denote the corresponding order statistics, then for any s -th order statistic $Y_s = Z_{(s)}$ the associated density function can be obtained as

$$f_2(y_s | \underline{x}, \alpha, \lambda) = s \binom{\nu}{s} \sum_{j=0}^{s-1} (-1)^{s-1-j} \binom{s-1}{j} (1 - F(y_s))^{\nu-1-j} f(y_s), \quad y_s > 0.$$

The associated predictive survival function is then obtained as

$$S_2(t | \underline{x}, \alpha, \lambda) = P(Y_s > t),$$

where

$$P(Y_s > t) = s \binom{\nu}{s} \sum_{j=0}^{s-1} (-1)^{s-1-j} \binom{s-1}{j} \frac{(1 - F(y_s))^{\nu-j}}{\nu - j}.$$

Thus, the two-sample posterior predictive density and the corresponding posterior predictive survival function, respectively, can be obtained as

$$f_2^*(y_s | \underline{x}) = \int_0^\infty \int_0^\infty f_2(y_s | \underline{x}, \alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda,$$

and

$$S_2^*(t | \underline{x}) = \int_0^\infty \int_0^\infty S_2(t | \underline{x}, \alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda.$$

Therefore, the estimator for Y_s can be obtained under the GEL function as

$$\hat{y}_s = \left\{ \int_0^\infty y^{-w} f_2^*(y | \underline{x}) dy \right\}^{-\frac{1}{w}} = \left\{ \int_0^\infty \int_0^\infty Q_2(\alpha, \lambda) \pi(\alpha, \lambda | \underline{x}) d\alpha d\lambda \right\}^{-\frac{1}{w}}, \quad (5.5)$$

where

$$\begin{aligned} Q_2(\alpha, \lambda) &= \int_0^\infty y^{-w} f_2(y | \underline{x}, \alpha, \lambda) dy \\ &= \nu \binom{\nu}{s} \sum_{j=0}^{s-1} (-1)^{s-1-j} \binom{s-1}{j} \int_0^1 \{F^{-1}(u)\}^{-w} (1-u)^{\nu-1-j} du. \end{aligned}$$

In a similar manner as mentioned in the case of one-sample prediction, we can estimate the value of y_s as

$$\hat{y}_s = \left\{ \frac{1}{M - M_0} \sum_{i=M_0+1}^M Q_2(\alpha_i, \lambda_i) \right\}^{-\frac{1}{w}}. \quad (5.6)$$

Finally the two-sided $100(1 - \psi)\%$ credible interval for the future observation y_s can be constructed by solving the following equations

$$S_2^*(L | \underline{x}) = 1 - \frac{\psi}{2}, \quad S_2^*(U | \underline{x}) = \frac{\psi}{2}. \quad (5.7)$$

The $100(1 - \psi)\%$ HPD predictive intervals can be obtained by proceeding in a similar manner to the one-sample prediction case.

6 Application with Real-Life Dataset

In this section, we consider two real-life data sets and show that GIL distribution is a good fit for these data sets compared to some other distributions. Various estimation and prediction methodologies have been demonstrated using type-i hybrid censored samples based on these datasets.

Example 6.1 In this example, we consider the death rates due to COVID-19 of Italy, from the period 17th July 2020 to 29th August 2020 for 44 days. The data set has been taken from the website [<https://covid19.who.int/WHO-COVID-19-global-data.csv>]. The death rates have been considered with population size 1000. The data set is given as: 0.09581, 0.05264, 0.06693, 0.01433, 0.06203, 0.07153, 0.04287, 0.04756, 0.02375, 0.02372, 0.02369, 0.02368, 0.05204, 0.02835, 0.01415, 0.04237, 0.02351, 0.03757, 0.05632, 0.02345, 0.04681, 0.02803, 0.01398, 0.06049, 0.00929, 0.01855, 0.02777, 0.04619, 0.02765, 0.01379, 0.72409, 0.01830, 0.01828, 0.02281, 0.03184, 0.02719, 0.04061, 0.01347, 0.03126, 0.01779, 0.01772, 0.05724, 0.02188, 0.03914.

To check the randomness of the data, we employ the run test at 5% level of significance. The p-value obtained is 0.2223, indicating that we can not reject the null hypothesis. We have fitted the generalized inverse Lindley distribution (GILD), inverse Lindley distribution (ILD), inverse gamma distribution (IGD), inverse Weibull distribution (IWD) and generalized inverse exponential distribution (GIED) to this data set using the Kolmogorov-Smirnov test at 5% significance level. We have reported the test statistics with p-values, Akaike information criteria (AIC), corrected AIC (AICc), Bayesian information criterion (BIC) of all the distributions in Table 6.1.

From Table 6.1, it is clear that the GILD gives a better fitting compared to the other four distributions. Here we consider two censoring conditions, such as Scheme-1 with $(r, T) = (30, 0.04)$ and Scheme-2 with $(r, T) = (40, 0.06)$. Since there is no prior information available for the parameters, we use non-informative prior for computing the Bayes estimators. Using the type-i hybrid censored samples, we have computed all the point and interval estimators (see Tables 6.2-6.5.) Further, in Figure 1 the trace and auto-correlation plots for the generated posterior samples of the parameters α and λ are given. The trace plot guarantees the well-mixed behavior or randomness of MCMC samples, whereas the auto-correlation plot guarantees the convergence of the method.

Based on our computational results, we recommend to use $\hat{\alpha}_{TK4} = 1.16276, 1.25508$, $\hat{\alpha}_{MC4} = 1.16746, 1.24874$, for α . For estimating λ we recommend to use $\hat{\lambda}_{TK1} = 0.03204, 0.02233$, and $\hat{\lambda}_{MC1} = 0.03125, 0.02205$. The confidence intervals (0.86838, 1.51148) or (0.98210, 1.55490) are preferred for α whereas the intervals (0.00337, 0.07214) or (0.00392, 0.04807) are recommended for λ .

Table 6.1: MLE, KS-test, AIC, AICc and BIC for different models for real data

Model	MLE		KS		AIC	AICc	BIC
	α	λ	Statistics	p-value			
IGD	2.77934	0.07351	0.08070	0.91466	-213.77260	-213.47990	-210.20420
ILD	0.05160	-	0.13502	0.36586	-212.84910	-212.75380	-211.06490
IWD	1.84659	0.00098	0.07488	0.95046	-214.81870	-214.52610	-211.25040
GIED	3.00223	0.04808	0.08965	0.84018	-212.80840	-212.51570	-209.24000
GILD	1.25921	0.01947	0.07244	0.96214	-215.01220	-214.71950	-211.44380

Table 6.2: Various point estimates for α with two different schemes

(r, T)	$\hat{\alpha}_{ML}$	$\hat{\alpha}_{TK1}$	$\hat{\alpha}_{TK2}$	$\hat{\alpha}_{TK3}$	$\hat{\alpha}_{TK4}$	$\hat{\alpha}_{MC1}$	$\hat{\alpha}_{MC2}$	$\hat{\alpha}_{MC3}$	$\hat{\alpha}_{MC4}$
(30, 0.04)	1.19495	1.18673	1.17869	1.17570	1.16276	1.19109	1.18523	1.17341	1.16746
(40, 0.06)	1.27137	1.26641	1.25991	1.25803	1.25508	1.26577	1.26152	1.25301	1.24874

Table 6.3: Various point estimates for λ with two different schemes

(r, T)	$\hat{\lambda}_{ML}$	$\hat{\lambda}_{TK1}$	$\hat{\lambda}_{TK2}$	$\hat{\lambda}_{TK3}$	$\hat{\lambda}_{TK4}$	$\hat{\lambda}_{MC1}$	$\hat{\lambda}_{MC2}$	$\hat{\lambda}_{MC3}$	$\hat{\lambda}_{MC4}$
(30, 0.04)	0.02530	0.03204	0.02907Z	0.02434	0.02165	0.03125	0.02831	0.02291	0.02048
(40, 0.06)	0.01866	0.02233	0.02052	0.01769	0.01624	0.02205	0.02042	0.01734	0.01591

Table 6.4: Various interval estimates for α with two different schemes

(r, T)	Bounds	ACI	MACI	Boot-p	Boot-t	HPD
(30, 0.04)	Lower Bound	0.87392	0.91443	0.92077	0.92246	0.86838
	Upper Bound	1.51598	1.56265	1.64114	1.83658	1.51148
(40, 0.06)	Lower Bound	0.98489	1.01488	1.05896	1.08284	0.98210
	Upper Bound	1.55786	1.59270	1.66288	1.74637	1.55490

Table 6.5: Various interval estimates for λ with two different schemes

(r, T)	Bounds	ACI	MACI	Boot-p	Boot-t	HPD
(30, 0.04)	Lower Bound	0.00000	0.00722	0.00373	0.00278	0.00337
	Upper Bound	0.05708	0.08839	0.08757	0.07938	0.07214
(40, 0.06)	Lower Bound	0.00000	0.00610	0.00453	0.00357	0.00392
	Upper Bound	0.03951	0.05704	0.04220	0.04228	0.04807

Table 6.6: One- and Two-Sample Bayesian prediction for the real data-set

	(r, T)	s	$\hat{y}_s$				HPD			Credible Interval		
			w1	w2	w3	w4	L	U	Length	L	U	Length
1-S	(30, 0.04)	d+1	0.04179	0.04177	0.04173	0.04171	0.03999	0.04552	0.00553	0.040	0.047	0.007
		d+2	0.04374	0.04369	0.04360	0.04356	0.04005	0.04943	0.00938	0.040	0.051	0.011
		d+3	0.04589	0.04581	0.04566	0.04559	0.04051	0.05364	0.01313	0.041	0.056	0.015
		d+4	0.04829	0.04816	0.04793	0.04783	0.04113	0.05817	0.01704	0.042	0.061	0.019
		d+5	0.05097	0.05079	0.05047	0.05031	0.04188	0.06313	0.02125	0.043	0.067	0.024
	(40, 0.06)	d+1	0.06568	0.06554	0.06530	0.06518	0.06001	0.07775	0.01774	0.060	0.083	0.023
		d+2	0.07311	0.07275	0.07209	0.07179	0.06013	0.09396	0.03383	0.061	0.102	0.041
		d+3	0.08353	0.08272	0.08131	0.08069	0.06134	0.11707	0.05573	0.064	0.129	0.065
		d+4	0.09990	0.09806	0.09502	0.09373	0.06320	0.15526	0.09206	0.068	0.178	0.110
		d+5	0.13257	0.12703	0.11908	0.11602	0.06617	0.24076	0.17459	0.075	0.294	0.219
2-S	(30, 0.04)	1	0.01012	0.00998	0.00969	0.00955	0.00600	0.01500	0.00900	0.006	0.015	0.009
		2	0.01185	0.01174	0.01151	0.01139	0.00700	0.01600	0.00900	0.008	0.017	0.009
		3	0.01311	0.01300	0.01279	0.01268	0.00900	0.01800	0.00900	0.009	0.018	0.009
		4	0.01399	0.01406	0.01385	0.01375	0.01000	0.01900	0.00900	0.010	0.019	0.009
		5	0.01786	0.01495	0.01481	0.01470	0.01100	0.02000	0.00900	0.011	0.020	0.009
	(40, 0.06)	1	0.01046	0.01035	0.01011	0.00999	0.00600	0.01500	0.00900	0.006	0.015	0.009
		2	0.01214	0.01205	0.01185	0.01176	0.00800	0.01600	0.00800	0.008	0.017	0.009
		3	0.01333	0.01325	0.01307	0.01298	0.00900	0.01800	0.00900	0.009	0.018	0.009
		4	0.01435	0.01425	0.01408	0.01399	0.01000	0.01900	0.00900	0.010	0.019	0.009
		5	0.01530	0.01517	0.01498	0.01489	0.01100	0.02000	0.00900	0.011	0.020	0.009

Using one-sample prediction, we predict the lifetimes of the first five units from the remaining censored observations. To predict the censored observations, we consider the GEL function with $w = -1, -0.5, 0.5, 1$. Further, the equal-tail and HPD predictive intervals are computed with 95% confidence level. Similarly, we predict a future sample of size ν (size of the real data set) using the two-sample prediction procedure. We compute the HPD and equal-tail predictive intervals with the same confidence level, that is, 95%. We report the predicted values in Table 6.6, where 1-S indicates one-sample prediction and 2-S indicates two-sample prediction. In the table, we denote $w = -1, 0.5, 0.5$, and 1 as $w1, w2, w3$, and $w4$, respectively. From Table 6.6, we notice that the lengths of the predictive intervals increase as s increases for both one and two-sample predictions. The lengths of the HPD predictive intervals are smaller than the equal-tail predictive intervals. In the one-sample prediction, the equal-tail and HPD contain the true observation of the given full data set.

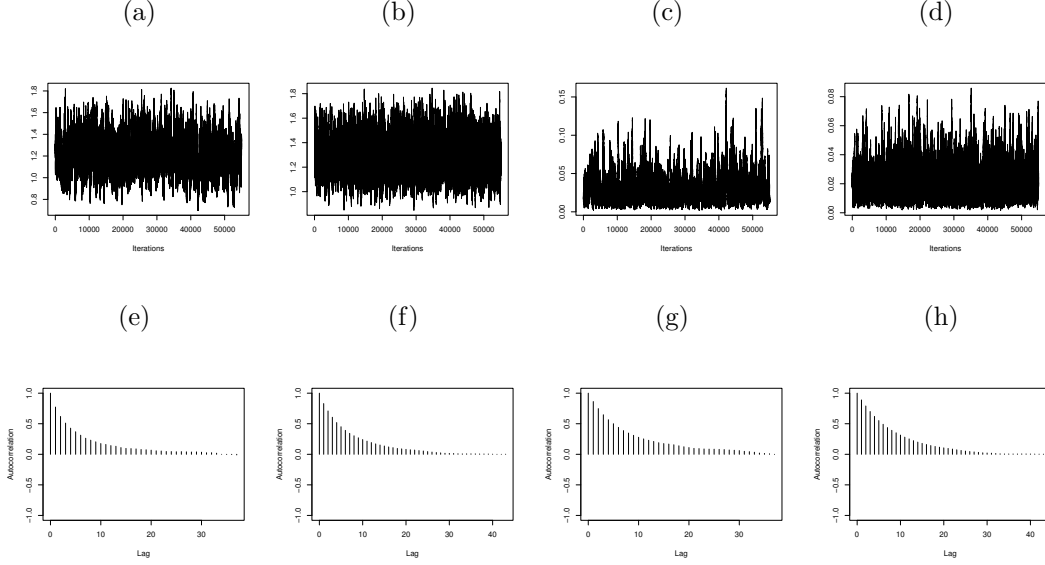


Figure 1: (a),(b) represents trace-plots for α and (c),(d) represents trace-plots for λ under the censoring Scheme-1, Scheme-2 respectively and (e), (f), (g), (h) represents corresponding autocorrelation plots

Example 6.2 This data set is collected from Wu and Wu (2005). We consider the data set presented in Wu and Wu (2005). The data set is regarding the duration of remission achieved by four drugs in the treatment of leukaemia for four groups, each having 20 patients. Here we consider the third group data set for our purposes. The data are: 4.158, 4.025, 5.17, 11.909, 4.912, 4.629, 3.955, 6.735, 3.14, 12.446, 8.777, 6.321, 3.256, 8.25, 3.759, 5.205, 3.071, 3.147, 9.773, 10.218.

Like Example 6.1, we used the run test to check for randomness and obtained the p-value as 0.6459. We have fitted the GILD and the other four popular distributions mentioned in Example 6.1 using the Kolmogorov-Smirnov test at 5% significance level, which is reported in Table 6.7.

From the results, it is clear that the GILD gives a better fitting compared to the other four distributions. Using Scheme-1 $((r, T) = (15, 9))$ and the Scheme-2 $((r, T) = (18, 11))$ we have computed all the estimators which are reported in Tables 6.8-6.11. Further, in Figure 2 the trace and auto-correlation plots for the generated posterior samples of the parameters α and λ are given. The trace plot guarantees the randomness of MCMC samples, whereas the auto-correlation plot guarantees the convergence of the method.

Based on our computational results, we recommend to use $\hat{\alpha}_{TK4} = 2.35378, 2.49175$, $\hat{\alpha}_{MC4} = 2.30438, 2.53022$, for α . For estimating λ we recommend to use $\hat{\lambda}_{TK1} = 52.36935, 60.09215$ and $\hat{\lambda}_{MC1} = 45.01995, 62.76743$. The confidence intervals (1.54984, 3.31618) or (1.68854, 3.51278) are preferred for α where as the intervals (7.10718, 109.41670) or (8.07375, 143.76080) are recommended for λ .

Table 6.7: MLE, KS-test, AIC, AICc and BIC for different models for real data

Model	MLE		KS		AIC	AICc	BIC
	α	λ	Statistics	p-value			
IGD:	5.38176	26.98712	0.12343	0.88477	96.65332	97.35920	98.64478
ILD:	5.75683	-	0.32722	0.02083	113.77950	114.00170	114.77520
IWD:	2.71996	57.75340	0.13050	0.84237	95.89812	96.60400	97.88959
GIED:	7.25388	13.07941	0.14171	0.76569	97.36148	98.06736	99.35294
GILD:	2.71922	58.66230	0.13046	0.84266	95.89707	96.60296	97.88854

Table 6.8: Various point estimates for α with two different schemes

(r, T)	$\hat{\alpha}_{ML}$	$\hat{\alpha}_{TK1}$	$\hat{\alpha}_{TK2}$	$\hat{\alpha}_{TK3}$	$\hat{\alpha}_{TK4}$	$\hat{\alpha}_{MC1}$	$\hat{\alpha}_{MC2}$	$\hat{\alpha}_{MC3}$	$\hat{\alpha}_{MC4}$
(15, 9)	2.49465	2.47119	2.45521	2.38773	2.35378	2.39858	2.37616	2.32908	2.30438
(18, 11)	2.60574	2.59106	2.59246	2.51378	2.49175	2.62136	2.59900	2.55341	2.53022

Table 6.9: Various point estimates for λ with two different schemes

(r, T)	$\hat{\lambda}_{ML}$	$\hat{\lambda}_{TK1}$	$\hat{\lambda}_{TK2}$	$\hat{\lambda}_{TK3}$	$\hat{\lambda}_{TK4}$	$\hat{\lambda}_{MC1}$	$\hat{\lambda}_{MC2}$	$\hat{\lambda}_{MC3}$	$\hat{\lambda}_{MC4}$
(15, 9)	43.77638	52.36935	45.59418	35.08498	31.46231	45.01995	41.14720	33.88170	30.61884
(18, 11)	50.57308	60.09215	52.68082	41.61775	37.46894	62.76743	56.34139	45.23280	40.68391

Table 6.10: Various interval estimates for α with two different schemes

(r, T)	Bounds	ACI	MACI	Boot-p	Boot-t	HPD
(15, 9)	Lower Bound	1.47055	1.65510	1.77303	1.76083	1.54984
	Upper Bound	3.51875	3.76123	4.29881	4.31063	3.31618
(18, 11)	Lower Bound	1.62903	1.79082	1.94118	1.84797	1.68854
	Upper Bound	3.58245	3.79307	3.93762	3.91840	3.51278

Table 6.11: Various interval estimates for λ with two different schemes

(r, T)	Bounds	ACI	MACI	Boot-p	Boot-t	HPD
(15, 9)	Lower Bound	0.00000	10.63117	18.23597	19.78022	7.10718
	Upper Bound	105.75830	180.4332	326.19860	385.49690	109.41670
(18, 11)	Lower Bound	0.00000	12.96837	22.22372	20.15436	8.07375
	Upper Bound	119.34920	197.60940	270.92320	238.12220	143.76080

Table 6.12: One- and Two-Sample Bayesian prediction for the real data-set

	(r, T)	s	$\hat{y}_s$				HPD			Credible Interval		
			w1	w2	w3	w4	L	U	Length	L	U	Length
1-S	(15, 9)	d+1	9.14288	9.11631	9.06950	9.04868	8.25016	11.09340	2.84324	8.26963	11.94490	3.67527
		d+2	10.38729	10.30703	10.17075	10.11192	8.26700	14.00582	5.73882	8.57652	15.46412	6.88760
		d+3	12.32785	12.11987	11.78776	11.65092	8.41788	18.72619	10.30831	8.85871	21.57207	12.71336
	(18, 11)	d+4	16.10856	15.46847	14.57171	14.23355	8.70445	28.79972	20.09527	9.55484	35.33792	25.78308
		d+5	16.10856	25.97625	21.59507	20.34107	8.70445	28.79972	20.09527	10.92618	104.9996	94.07345
		d+1	12.96545	12.76877	12.47201	12.35542	10.21843	19.30582	9.08739	10.27145	22.54601	12.27456
2-S	(15, 9)	d+2	22.72324	20.34520	18.06679	17.35939	10.27694	48.08197	37.80503	11.62967	83.73200	72.10233
		1	2.69676	2.67240	2.62126	2.59419	1.69096	3.69745	2.00649	1.70058	3.70931	2.00873
		2	3.07299	3.05125	3.00658	2.98349	2.06373	4.10642	2.04269	2.08532	4.13234	2.04702
		3	3.36050	3.33892	3.29520	3.27297	2.31732	4.45119	2.13387	2.35766	4.50088	2.14322
		4	3.61840	3.59595	3.55121	3.52848	2.52659	4.78271	2.25612	2.58686	4.86101	2.27415
	(18, 11)	5	3.86617	3.84211	3.79383	3.77087	2.71229	5.11870	2.40641	2.79405	5.23155	2.43750
		1	2.76404	2.74304	2.69950	2.67677	1.82406	3.70451	1.88045	1.83789	3.72091	1.88302
		2	3.12047	3.10174	3.06350	3.04393	2.17995	4.08248	1.90253	2.20439	4.11185	1.90746
		3	3.39016	3.37158	3.33402	3.31511	2.41906	4.39927	1.98021	2.45858	4.44821	1.98963
		4	3.63029	3.61103	3.57327	3.55326	2.61508	4.70108	2.08600	2.67087	4.77317	2.10230
		5	3.85952	3.83901	3.79447	3.77820	2.78881	5.00426	2.21545	2.86222	5.10396	2.24174

We predict a future observation using one and two sample predictions. The significance level and the choice of w in the loss function remain the same as considered in Example 6.1. We have reported the predicted values in Table 6.12. For this data-set also the lengths of the HPD predictive intervals are smaller than the equal-tail predictive intervals. Moreover the equal-tail and HPD predictive intervals contain the true observation for one-sample prediction.

7 Estimating Parameters of the GIL Distribution Under Type-II Hybrid Censoring Scheme

In this section, we discuss the estimation procedures for estimating the parameters of the GIL distribution under the type-ii hybrid censoring scheme in order to give a glimpse of an idea about the inferential

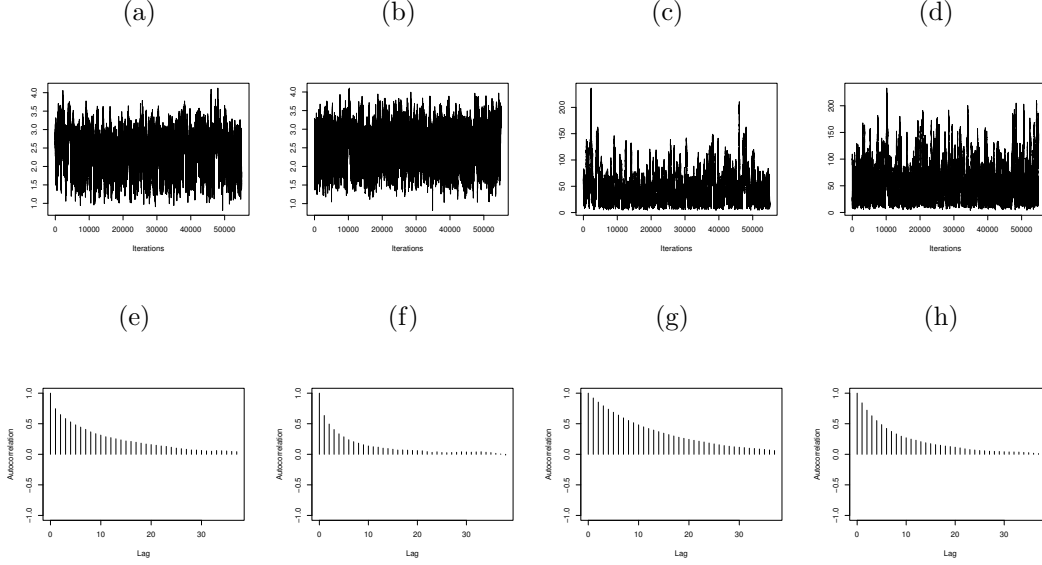


Figure 2: (a),(b) represents trace-plots for α and (c),(d) represents trace-plots for λ under the censoring Scheme-1, Scheme-2 respectively, and (e), (f), (g), (h) represents corresponding auto-correlation plots

procedures. The results obtained under the type-i hybrid censoring can easily be obtained under the type-ii hybrid censoring scheme without making many calculation changes. The only difference will be to generate the samples from GIL distribution under the type-ii hybrid conditions. The details of the calculations and related numerical results will be omitted for brevity.

Note that, Childs et al. (2003) introduced type-ii hybrid censoring scheme, which is a modification of type-i hybrid censoring scheme. In the type-ii hybrid censoring scheme, the experiment terminates at a time which is $\max(X_{r:n}, T)$. The type-ii hybrid censored sample is obtained as follows.

$$\begin{cases} \text{Case-I : } X_{1:n} \leq X_{2:n} \leq \dots \leq X_{r:n}, & \text{if } X_{r:n} \geq T, \\ \text{Case-II : } X_{1:n} \leq X_{2:n} \leq \dots \leq X_{m:n}, & \text{if } r < m < n \text{ with } X_{m:n} \leq T < X_{m+1:n}, \\ \text{Case-III : } X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}, & \text{if } X_{n:n} \leq T. \end{cases} \quad (7.1)$$

This censoring scheme ensures that the observed sample size is not less than r , whereas the type-i hybrid censoring scheme does not guarantee the minimum size of the observed units. Under this censoring scheme, the likelihood function is written as

$$l(\alpha, \lambda) \propto \begin{cases} \prod_{i=1}^r f(x_{i:n}) [1 - F(x_{r:n})]^{(n-r)}, & \text{if Case-I holds true} \\ \prod_{i=1}^m f(x_{i:n}) [1 - F(T)]^{(n-m)} & \text{if Case-II holds true} \\ \prod_{i=1}^n f(x_{i:n}) & \text{if Case-III holds true.} \end{cases}$$

The above Cases can be combined and written in a general form as

$$l(\alpha, \lambda) \propto \prod_{i=1}^{d^*} f(x_{i:n}) [1 - F(c^*)]^{(n-d^*)} \quad (7.2)$$

where

$$c^* = \begin{cases} x_{r:n}, & \text{if Case-I holds} \\ T, & \text{if Case-II holds} \\ x_{n:n}, & \text{if Case-III holds} \end{cases} \quad \text{and} \quad d^* = \begin{cases} r, & \text{if Case-I is true} \\ m, & \text{if Case-II is true} \\ n, & \text{if Case-III is true.} \end{cases}$$

Using the above expressions, the general form of the log-likelihood function using type-ii hybrid censored samples from the GIL distribution is obtained as

$$\begin{aligned} L(\alpha, \lambda) \propto & d^* \log \alpha + 2d^* \log \lambda - d^* \log(1 + \lambda) + \sum_{i=1}^{d^*} \log(1 + x_{i:n}^{-\alpha}) - (1 + \alpha) \sum_{i=1}^{d^*} \log x_{i:n} \\ & - \lambda \sum_{i=1}^{d^*} x_{i:n}^{-\alpha} + (n - d^*) \log \left\{ 1 - \left(1 + \frac{\lambda}{1 + \lambda} c^{*-\alpha} \right) e^{-\lambda c^{*-\alpha}} \right\}. \end{aligned} \quad (7.3)$$

It is easy to observe that the log-likelihood function is very much similar to the log-likelihood function obtained under the type-i hybrid censoring scheme (see (2.1)). Thus taking $d = d^*$ and $c = c^*$ in the equation (2.1) and applying the E-M algorithm, we can obtain the MLEs of the parameters α and λ as obtained under the type-i hybrid censoring scheme. Furthermore, the asymptotic confidence intervals can be obtained by computing the observed Fisher information matrix as discussed there.

The general form of the likelihood function under the type-ii hybrid censoring scheme is the same as the available form of the likelihood function under the type-i hybrid censoring scheme. The other estimation methods, such as the Bayesian estimation using MCMC and Tierney and Kadane's procedures under the gamma prior, can be used to derive the Bayes estimators, which will have the same forms. The confidence intervals, such as the boot-t, boot-p and HPD credible intervals, can also be derived similarly. The details of the calculations and the respective computational results have been omitted for the sake of brevity.

8 Concluding Remarks

Several inverse probability models are available to model real-life datasets showing upside-down bathtub shapes, including the GIL distribution. This distribution has been proposed as an alternative to some popular distributions such as inverse gamma, inverse Weibull, generalized inverse exponential and inverse Gaussian etc., for modeling lifetime datasets by the authors (see Sharma et al. (2016)). In this article, we have investigated the problem of point and interval estimation of the parameters associated with the GIL distribution using type-i hybrid censored samples.

Several point estimators are derived for both the parameters using the GEL function, which is a generalization of the squared error loss function. The MLEs of the parameters could not be obtained in closed-form expressions. Using the E-M algorithm, the approximate MLEs are derived numerically. Bayesian estimation of the parameters is considered with respect to proper and improper priors and the GEL function, whose closed-form expressions do not exist. Utilizing the superior computational facilities available nowadays, we obtained approximate Bayes estimators through some approximation methods, namely Tierney and Kadane's approximation method and the MCMC method. The point estimators are compared through MSE and biases. It has been observed that the Bayes estimators of the parameters using the MCMC method have better performances than other estimators.

Several interval estimators, such as the ACIs, boot-p, boot-t, and HPD credible intervals for the associated parameters, are derived numerically. The interval estimators are compared using CPs, ALs

and a unique measure criterion known as PCD. It has been noticed that the HPD credible intervals for the parameters perform better than other proposed confidence intervals in terms of CPs, ALs, and PCD values. Further, we have predicted a future observation using one and two sample Bayesian prediction procedures. We have considered two real-life datasets that have been satisfactorily modeled using the generalized inverse Lindley distribution. In fact, it is seen that for these two datasets, GIL distribution fits better compared to some other popular distributions. Then the estimation methodologies have been illustrated to show the potential application of our model problem. We hope that the current research work will shed some light on future direction, such as inference of associated parameters of the GIL distribution under some advanced hybrid censoring schemes.

Acknowledgment: The authors would like to sincerely thank the two anonymous reviewers and the associate editor for their valuable suggestions and comments which have helped in improving the presentation of the manuscript. The Second author (Manas Ranjan Tripathy) would like to thank Department of Science and Technology, SERB, [EMR/2017/003078], New Delhi, India for providing some financial support.

References

- Abushal, T. A. (2019). Parameter estimation of Weibull-Exponential distribution under type-i hybrid censored sample. *Journal of King Saud University-Science*, 31(4):1431–1436.
- Basu, S., Singh, S. K., and Singh, U. (2019). Estimation of inverse Lindley distribution using product of spacings function for hybrid censored data. *Methodology and Computing in Applied Probability*, 21(4):1377–1394.
- Calabria, R. and Pulcini, G. (1994). An engineering approach to Bayes estimation for the Weibull distribution. *Microelectronics Reliability*, 34(5):789–802.
- Chen, M.-H. and Shao, Q.-M. (1999). Monte carlo estimation of Bayesian credible and hpd intervals. *Journal of Computational and Graphical Statistics*, 8(1):69–92.
- Chib, S. and Greenberg, E. (1995). Understanding the Metropolis-Hastings algorithm. *The American Statistician*, 49(4):327–335.
- Childs, A., Chandrasekar, B., Balakrishnan, N., and Kundu, D. (2003). Exact likelihood inference based on type-i and type-ii hybrid censored samples from the exponential distribution. *Annals of the Institute of Statistical Mathematics*, 55(2):319–330.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). Maximum likelihood from incomplete data via the e-m algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1):1–22.
- Dey, S. and Pradhan, B. (2014). Generalized inverted exponential distribution under hybrid censoring. *Statistical Methodology*, 18:101–114.
- Draper, N. and Guttman, I. (1987). Bayesian analysis of hybrid life tests with exponential failure times. *Annals of the Institute of Statistical Mathematics*, 39(1):219–225.
- Efron, B. and Tibshirani, R. J. (1994). *An Introduction to the Bootstrap*. Chapman & Hall, New York.
- Epstein, B. (1954). Truncated life tests in the exponential case. *The Annals of Mathematical Statistics*, 25(3):555–564.

- Górny, J. and Cramer, E. (2018). Modularization of hybrid censoring schemes and its application to unified progressive hybrid censoring. *Metrika*, 81(2):173–210.
- Greene, W. H. (2018). *Econometric Analysis*. Prentice Hall, New York.
- Gupta, P. K. and Singh, B. (2013). Parameter estimation of Lindley distribution with hybrid censored data. *International Journal of System Assurance Engineering and Management*, 4(4):378–385.
- Kayal, T., Tripathi, Y. M., and Rastogi, M. K. (2018). Estimation and prediction for an inverted exponentiated Rayleigh distribution under hybrid censoring. *Communications in Statistics-Theory and Methods*, 47(7):1615–1640.
- Khatun, H., Tripathy, M. R., and Pal, N. (2020). Hypothesis testing and interval estimation for quantiles of two normal populations with a common mean. *Communications in Statistics-Theory and Methods*, (<https://doi.org/10.1080/03610926.2020.1845735>).
- Kundu, D. and Pradhan, B. (2009). Estimating the parameters of the generalized exponential distribution in presence of hybrid censoring. *Communications in Statistics- Theory and Methods*, 38(12):2030–2041.
- Lindley, D. V. (1980). Approximate Bayesian methods. *Trabajos de estadística y de investigación operativa*, 31(1):223–245.
- Little, R. J. and Rubin, D. B. (1983). Missing data in large data sets. In *Statistical Methods & the Improvement of Data Quality*, pages 215–243. Elsevier.
- Louis, T. A. (1982). Finding the observed information matrix when using the e-m algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 44(2):226–233.
- Meeker, W. Q. and Escobar, L. A. (1998). *Statistical Methods for Reliability Data*. John Wiley & Sons, New York.
- Nassar, M. and Abo-Kasem, O. (2017). Estimation of the inverse Weibull parameters under adaptive type-II progressive hybrid censoring scheme. *Journal of Computational and Applied Mathematics*, 315:228–239.
- Pandey, A., Kaushik, A., Singh, S. K., and Singh, U. (2021). Statistical analysis for generalized progressive hybrid censored data from Lindley distribution under step-stress partially accelerated life test model. *Austrian Journal of Statistics*, 50(1):105–120.
- Pradhan, B. and Kundu, D. (2009). On progressively censored generalized exponential distribution. *Test*, 18(3):497–515.
- Saraiva, E. F. and Suzuki, A. K. (2017). Bayesian computational methods for estimation of two-parameters Weibull distribution in presence of right-censored data. *Chilean Journal of Statistics*, 8(2):25–43.
- Sen, T., Bhattacharya, R., Tripathi, Y. M., and Pradhan, B. (2019). Inference and optimum life testing plans based on type-ii progressive hybrid censored generalized exponential data. *Communications in Statistics- Simulation and Computation*, 49(12):1–29.

- Sharma, V. K., Singh, S. K., Singh, U., and Merovci, F. (2016). The generalized inverse Lindley distribution: A new inverse statistical model for the study of upside-down bathtub data. *Communications in Statistics-Theory and Methods*, 45(19):5709–5729.
- Singh, S. K., Singh, U., and Sharma, V. K. (2016). Estimation and prediction for type-i hybrid censored data from generalized Lindley distribution. *Journal of Statistics and Management Systems*, 19(3):367–396.
- Sultana, F., Tripathi, Y. M., Rastogi, M. K., and Wu, S.-J. (2018). Parameter estimation for the Kumaraswamy distribution based on hybrid censoring. *American Journal of Mathematical and Management Sciences*, 37(3):243–261.
- Tierney, L. and Kadane, J. B. (1986). Accurate approximations for posterior moments and marginal densities. *Journal of the American Statistical Association*, 81(393):82–86.
- Tripathi, Y. M. and Rastogi, M. K. (2016). Estimation using hybrid censored data from a generalized inverted exponential distribution. *Communications in Statistics- Theory and Methods*, 45(16):4858–4873.
- Wu, S.-F. and Wu, C.-C. (2005). Two stage multiple comparisons with the average for exponential location parameters under heteroscedasticity. *Journal of Statistical Planning and Inference*, 134(2):392–408.

Appendix: From (2.4), we obtain the second order derivatives of $L_c(\alpha, \lambda)$ with respect to α and λ as

$$\begin{aligned}\frac{\partial^2 L_c}{\partial \alpha^2} &= -\frac{n}{\alpha^2} + \sum_{i=1}^d \frac{x_{i:n}^{-\alpha} (\log x_{i:n})^2}{(1 + x_{i:n}^{-\alpha})^2} - \lambda \sum_{i=1}^d \frac{x_{i:n}^{-\alpha} (\log x_{i:n})^2}{(1 + x_{i:n}^{-\alpha})^2} + \sum_{i=1}^{n-d} \frac{z_i^{-\alpha} (\log z_i)^2}{(1 + z_i^{-\alpha})^2} - \lambda \sum_{i=1}^{n-d} \frac{z_i^{-\alpha} (\log z_i)^2}{(1 + z_i^{-\alpha})^2}, \\ \frac{\partial^2 L_c}{\partial \lambda^2} &= -\frac{2n}{\lambda^2} + \frac{n}{(1 + \lambda)^2}, \\ \frac{\partial^2 L_c}{\partial \alpha \partial \lambda} &= \frac{\partial^2 L_c}{\partial \lambda \partial \alpha} = \sum_{i=1}^d x_{i:n}^{-\alpha} \log x_{i:n} + \sum_{i=1}^{n-d} z_i^{-\alpha} \log z_i.\end{aligned}$$

Therefore,

$$\begin{aligned}E\left(\frac{\partial^2 L_c}{\partial \alpha^2}\right) &= -\frac{n}{\alpha^2} + nE\left(\frac{x^{-\alpha} (\log x)^2}{(1 + x^{-\alpha})^2}\right) - n\lambda E\left(x^{-\alpha} (\log x)^2\right), \\ E\left(\frac{\partial^2 L_c}{\partial \lambda^2}\right) &= -\frac{2n}{\lambda^2} + \frac{n}{(1 + \lambda)^2}, \\ E\left(\frac{\partial^2 L_c}{\partial \alpha \partial \lambda}\right) &= E\left(\frac{\partial^2 L_c}{\partial \lambda \partial \alpha}\right) = nE\left(x^{-\alpha} \log x\right),\end{aligned}$$

where

$$\begin{aligned}E\left(\frac{x^{-\alpha} (\log x)^2}{(1 + x^{-\alpha})^2}\right) &= -\frac{1}{\alpha^2(1 + \lambda)} \int_0^1 \frac{\left(\log\left(-\frac{\log u}{\lambda}\right)\right)^2 \log u}{1 - \frac{\log u}{\lambda}} du, \\ E\left(x^{-\alpha} (\log x)^2\right) &= -\frac{1}{\alpha^2(1 + \lambda)} \int_0^1 \left(\log\left(-\frac{\log u}{\lambda}\right)\right)^2 \left(1 - \frac{\log u}{\lambda}\right) \log u \, du, \\ E\left(x^{-\alpha} \log x\right) &= \frac{1}{\alpha(1 + \lambda)} \int_0^1 \log\left(-\frac{\log u}{\lambda}\right) \left(1 - \frac{\log u}{\lambda}\right) \log u \, du.\end{aligned}$$

Again respect to observed data X , for the missing data Z (as mentioned in EM-algorithm), we have

$$\begin{aligned}f_Z(z | X; \alpha, \lambda) &= \frac{f_X(z; \alpha, \lambda)}{1 - F_X(c; \alpha, \lambda)} \\ &= \frac{1}{1 - F_X(c; \alpha, \lambda)} \frac{\alpha \lambda^2}{1 + \lambda} \left(\frac{1 + z^{-\alpha}}{z^{\alpha+1}}\right) e^{-\lambda z^{-\alpha}}, \quad z > c, \quad \alpha > 0, \quad \lambda > 0.\end{aligned}$$

Thus the second order derivatives of $\log f_Z(z | X; \alpha, \lambda)$ with respect to α and λ are obtained as

$$\begin{aligned}\frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \alpha^2} &= \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \alpha^2}}{1 - F_X(c, \alpha, \lambda)} + \left(\frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \alpha}}{1 - F_X(c, \alpha, \lambda)}\right)^2 - \frac{1}{\alpha^2} + \frac{z^{-\alpha} (\log z)^2}{(1 + z^{-\alpha})^2} - \lambda z^{-\alpha} (\log z)^2, \\ \frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \lambda^2} &= \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda^2}}{1 - F_X(c, \alpha, \lambda)} + \left(\frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \lambda}}{1 - F_X(c, \alpha, \lambda)}\right)^2 - \frac{2}{\lambda^2} + \frac{1}{(1 + \lambda)^2}, \\ \frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \alpha \partial \lambda} &= \frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \lambda \partial \alpha} = \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda \partial \alpha}}{1 - F_X(c, \alpha, \lambda)} + \frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \alpha} \frac{\partial F_X(c, \alpha, \lambda)}{\partial \lambda}}{1 - F_X(c, \alpha, \lambda)} + z^{-\alpha} \log z,\end{aligned}$$

where

$$\begin{aligned}
\frac{\partial F_X(c, \alpha, \lambda)}{\partial \alpha} &= \frac{\lambda^2 c^{-\alpha}}{1 + \lambda} (1 + c^{-\alpha}) e^{-\lambda c^{-\alpha}} \log c, \\
\frac{\partial F_X(c, \alpha, \lambda)}{\partial \lambda} &= -\frac{\lambda c^{-\alpha}}{(1 + \lambda)^2} \left((1 + \lambda)(1 + c^{-\alpha}) + 1 \right) e^{-\lambda c^{-\alpha}}, \\
\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \alpha^2} &= \frac{\lambda^2 c^{-\alpha}}{1 + \lambda} \left(\lambda c^{-2\alpha} + (\lambda - 2)c^{-\alpha} - 1 \right) e^{-\lambda c^{-\alpha}} (\log c)^2, \\
\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda^2} &= \frac{c^{-\alpha}}{(1 + \lambda)^3} \left(\lambda(1 + \lambda)^2 (1 + c^{-\alpha}) c^{-\alpha} + (\lambda^2 - 1)c^{-\alpha} - 2 \right) e^{-\lambda c^{-\alpha}}, \\
\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda \partial \alpha} &= \frac{\lambda c^{-\alpha}}{(1 + \lambda)^2} \left((1 + \lambda)(1 + \lambda c^{-\alpha}) + 1 \right) (1 + c^{-\alpha}) e^{-\lambda c^{-\alpha}} \log c.
\end{aligned}$$

Therefore,

$$\begin{aligned}
E_{Z|X} \left(\frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \alpha^2} \right) &= \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \alpha^2}}{1 - F_X(c, \alpha, \lambda)} + \left(\frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \alpha}}{1 - F_X(c, \alpha, \lambda)} \right)^2 - \frac{1}{\alpha^2} + E_{Z|X} \left(\frac{z^{-\alpha} (\log z)^2}{(1 + z^{-\alpha})^2} \right) \\
&\quad - \lambda E_{Z|X} \left(z^{-\alpha} (\log z)^2 \right), \\
E_{Z|X} \left(\frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \lambda^2} \right) &= \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda^2}}{1 - F_X(c, \alpha, \lambda)} + \left(\frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \lambda}}{1 - F_X(c, \alpha, \lambda)} \right)^2 - \frac{2}{\lambda^2} + \frac{1}{(1 + \lambda)^2}, \\
E_{Z|X} \left(\frac{\partial^2 \log f_Z(z | X; \alpha, \lambda)}{\partial \lambda \partial \alpha} \right) &= \frac{\frac{\partial^2 F_X(c, \alpha, \lambda)}{\partial \lambda \partial \alpha}}{1 - F_X(c, \alpha, \lambda)} + \frac{\frac{\partial F_X(c, \alpha, \lambda)}{\partial \alpha} \frac{\partial F_X(c, \alpha, \lambda)}{\partial \lambda}}{1 - F_X(c, \alpha, \lambda)} + E_{Z|X} \left(z^{-\alpha} \log z \right),
\end{aligned}$$

where,

$$\begin{aligned}
E_{Z|X} \left(\frac{z^{-\alpha} (\log z)^2}{(1 + z^{-\alpha})^2} \right) &= -\frac{1}{\alpha^2 (1 + \lambda) (1 - F_X(c; \alpha, \lambda))} \int_{e^{-\lambda c^{-\alpha}}}^1 \frac{\left(\log \left(-\frac{\log u}{\lambda} \right) \right)^2 \log u}{1 - \frac{\log u}{\lambda}} du, \\
E_{Z|X} \left(z^{-\alpha} (\log z)^2 \right) &= -\frac{1}{\alpha^2 (1 + \lambda) (1 - F_X(c; \alpha, \lambda))} \int_{e^{-\lambda c^{-\alpha}}}^1 \left(\log \left(-\frac{\log u}{\lambda} \right) \right)^2 \left(1 - \frac{\log u}{\lambda} \right) \log u \, du, \\
E_{Z|X} \left(z^{-\alpha} \log z \right) &= \frac{1}{\alpha (1 + \lambda) (1 - F_X(c; \alpha, \lambda))} \int_{e^{-\lambda c^{-\alpha}}}^1 \log \left(-\frac{\log u}{\lambda} \right) \left(1 - \frac{\log u}{\lambda} \right) \log u \, du.
\end{aligned}$$