



Classical Estimation and Bayes Risk of the Shape Parameter of Pareto Type -I Distribution

M. Geetha¹, R. Selvam²

¹ Assistant Professor (Research Scholar), Department of Statistics, Padmavani Arts and Science College for Women, Salem-11.

² Assistant Professor, Department of Statistics, Govt. Arts College (Autonomous), Salem-7.

ABSTRACT: In this paper, Pareto type –I distribution is proposed to compare the classical estimators such as the Maximum Likelihood Estimator (MLE), Uniformly Minimum Variance Unbiased Estimator (UMVUE) and Minimum Mean Square Error Estimator (MiniMSE). The Bayes risk can be obtained by using non – informative and informative prior under different loss functions such as Square Error Loss Function (SELF), Quadratic Loss Function (QLF), Precautionary Loss Function (PLF) and Entropy Loss Function (ELF) through simulation techniques. As per the result, it is observed that the MiniMSE is the best among the other proposed estimators. It is also found that the Bayes risk under QLF is least one among all the other loss functions namely SELF, PLF and ELF using informative prior.

Keywords: Bayes Estimator, Bayes Risk, Classical Estimation, Informative and non – informative priors, Loss function and Pareto distribution.

I. INTRODUCTION

The Pareto probability distribution is a simple model for non- negative data with positively Skewed distribution. This distribution was introduced by Wilfredo Pareto (1848-1923) especially for wealth distribution of the population of a city within a given area. The use of the Pareto distribution as model to analyses stock prize and instability in business and economic, field of bio medical science, risk factor in insurance company, migration of population, survival time in quadratic system, Geophysical phenomena in society, reliability and life testing. Al Omari Ahmed, hadecel salim, Al-Kutubi and Noor akma ibrhim, (2010), studied the comparison of the Bayesian estimation with maximum likelihood estimation. Al Omari Mohammed Ahmed and Noor akma ibrhim, (2011), have been studied the performance of MLE and Bayes estimation of survival function using non informative prior with right censored data. Sankudey and Sudhansu A.Maiti (2012), have studied the Bayes estimators of Rayleigh parameter and its associated risk based on extended Jeffrey's prior under the assumptions of both symmetric and asymmetric loss function. R.K.Radha, (2015), studied the Bayesian analysis of exponential distribution using informative prior. Kawsar Fatima and S.P Ahmad (2018) have been studied the Bayes estimation of shape parameter of Exponentiated moment exponential distribution using informative and non – informative priors under different loss functions. Gaurav Shukla, Umesh Chandra and Vinod kumar (2020), derived and examined the expression for risk function under three different loss function. It is remarkable that the development of appropriate Bayesian inference procure has been very limited. Bayesian inference in the Pareto type I distribution for the special case in which the scale parameter is known.

The probability density function (pdf) of Pareto type – I distribution is defined as

$$f(t; \alpha, \theta) = \begin{cases} \frac{\theta \alpha^\theta}{t^{\theta+1}}; & t > \alpha; \theta > 0; \alpha > 0 \end{cases} \quad (1.1)$$

Where t is a random variable, θ is the shape parameter and α is the scale parameter, which is known.

The moments of Pareto type –I distribution, were given by

$$\text{Mean, } E(t) = \frac{\alpha \theta}{\theta - 1}; \theta > 1$$

$$\text{Variance, } V(t) = \frac{\theta \alpha^2}{(\theta - 1)^2 (\theta - 2)}; \theta > 2$$

II. Classical Estimation

Classical estimation is an important estimation technique in statistics. In this section a specific method of estimation such as Maximum Likelihood Estimation (MLE), Uniformly Minimum Variance Unbiased Estimation (UMVUE), Mini Mean Square Error Estimation (MiniMSE), are considered to estimate the shape parameter of Pareto type I distribution.

2.1 Maximum Likelihood Estimator

Let $t_1, t_2, \dots, t_n$ be a set of 'n' random variables from Pareto Type I distribution with Parameters θ and α having the probability density function defined in (1.1), then the likelihood function,

$$\begin{aligned} L &= \prod_{i=1}^n f(t_i; \alpha, \theta) \\ &= \prod_{i=1}^n \frac{\theta \alpha^\theta}{t_i^{\theta+1}} \\ L &= \theta^n e^{n \theta \log \alpha} e^{-(\theta+1) \sum_{i=1}^n \log t_i} \\ \log L &= n \log \theta + n \theta \log \alpha - \theta \sum_{i=1}^n \log t_i \\ \frac{\partial \log L}{\partial \theta} &= \frac{n}{\theta} + n \log \alpha - \sum_{i=1}^n \log t_i \end{aligned}$$

Using the Maximization Likelihood Principle, we get the estimated value of θ as

$$\hat{\theta} = \left[\frac{\sum_{i=1}^n \log t_i}{n} - \log \alpha \right]^{-1}$$

and $\hat{\alpha} = \min_{1 \leq i \leq n} (t_i)$

In case of frequency distribution

$$\hat{\theta} = \left[\frac{\sum_{i=1}^n f_i \log t_i}{N} - \log \alpha \right]^{-1} \quad (2.1.1)$$

Where $N = \sum f$

2.2 Uniformly Minimum Variance Unbiased Estimator

The distribution whose density functions have the following general form

$f(t, \theta) = a(\theta) \cdot b(t) e^{[c(\theta)d(t)]}$ is known as one parameter exponential family of distribution.

The Pareto distribution belongs to the exponential family of distribution has the density function (1.1) can be written as

$$\begin{aligned} f(t, \theta) &= \theta e^{\theta \log \alpha - (\theta+1) \log_e t} \\ &= \theta e^{\theta \log_e \alpha - \theta \log_e t - \log_e t} \\ &= \theta e^{-\log_e t} e^{-\theta \log_e (t/\alpha)} \end{aligned}$$

Where,

$$a(\theta) = \theta, \quad b(t) = e^{\log_e t}, \quad c(\theta) = -\theta, \quad d(t) = \log_e \left(\frac{t}{\alpha} \right)$$

Therefore, the statistic $P = \sum_{i=1}^n \log_e \left(\frac{t_i}{\alpha} \right)$ is a complete sufficient statistic for θ .

It is easy to show that the statistic P is distributed as Gamma Distribution with Parameters n and θ .

If $T \sim \text{Pareto}(\alpha, \theta)$, then $\log_e(t/\alpha) \sim \text{exponential}(\theta)$

and $P = \sum_{i=1}^n \log_e \left(\frac{t_i}{\alpha} \right) \sim \text{Gamma}(n, \theta)$ with pdf

$$g(p) = \frac{\theta^n}{\Gamma(n)} p^{n-1} e^{-\theta p}; \quad p \geq 0, \theta > 0$$

Consider,

$$\begin{aligned} E\left(\frac{1}{p}\right) &= \int_0^\infty \frac{1}{p} g(p) dp \\ &= \int_0^\infty \frac{1}{p} \frac{\theta^n}{\Gamma(n)} p^{n-1} e^{-\theta p} dp \\ &= \frac{\theta^n}{\Gamma(n)} \frac{\Gamma(n-1)}{\theta^{n-1}} \end{aligned}$$

$$E\left(\frac{1}{p}\right) = \frac{\theta}{n-1}$$

$$E\left(\frac{n-1}{p}\right) = \theta$$

$\Rightarrow \frac{n-1}{p}$ is unbiased estimator for θ , P represents a complete sufficient statistics for θ . Thus, by theorem of Lehmann- Scheffe, the UMVUE of θ is given by

$$\hat{\theta}_{UMVUE} = \frac{n-1}{p}$$

2.3 Minimum Mean Square Error Estimator

The Minimum Mean Squared Error estimator (MinMSE) can be found in the class of estimators of the form $\frac{c}{p}$

(2.3.1)

Therefore

$$\begin{aligned} MSE_{\theta}\left(\frac{c}{p}\right) &= E\left[\left(\frac{c}{p} - \theta\right)^2\right] \\ &= E\left[\left(\frac{c}{p}\right)^2\right] + \theta^2 - 2\theta E\left[\left(\frac{c}{p}\right)\right] \\ &= c^2 E\left[\left(\frac{1}{p}\right)^2\right] + \theta^2 - 2c\theta E\left[\left(\frac{1}{p}\right)\right] \end{aligned}$$

Using maximum likelihood principle, we get

$$\begin{aligned} \frac{\partial}{\partial c} MSE_{\theta}\left(\frac{c}{p}\right) &= 0 \Rightarrow 2cE\left[\left(\frac{1}{p}\right)^2\right] - 2\theta E\left[\left(\frac{1}{p}\right)\right] \\ \therefore 2cE\left[\left(\frac{1}{p}\right)^2\right] - 2\theta E\left[\left(\frac{1}{p}\right)\right] &= 0 \\ 2cE\left[\left(\frac{1}{p}\right)^2\right] &= 2\theta E\left[\left(\frac{1}{p}\right)\right] \end{aligned}$$

Where,

$$\begin{aligned} c &= \frac{\theta E\left[\left(\frac{1}{p}\right)\right]}{E\left[\left(\frac{1}{p}\right)^2\right]} \\ E\left[\left(\frac{1}{p}\right)^2\right] &= \int_0^\infty \left(\frac{1}{p}\right)^2 g(p) dp \\ &= \int_0^\infty \left(\frac{1}{p}\right)^2 \frac{\theta^n}{\Gamma(n)} p^{n-1} e^{-\theta p} dp \\ &= \frac{\theta^n}{\Gamma(n)} \frac{\Gamma(n-2)}{\theta^{n-2}} \\ E\left[\left(\frac{1}{p}\right)^2\right] &= \frac{\theta^2}{(n-1)(n-2)} \end{aligned}$$

Substituting the equations (2.2.1), (2.3.3) in (2.3.2), we have

$$c = \frac{\theta E\left[\left(\frac{1}{p}\right)\right]}{E\left[\left(\frac{1}{p}\right)^2\right]} = \frac{\frac{\theta}{n-1}}{\frac{\theta^2}{(n-1)(n-2)}}$$

$$\Rightarrow c = (n-2)$$

Therefore,

$$\hat{\theta}_{MinMSE} = \frac{n-2}{p} \quad (\text{from 2.3.1})$$

(2.3.2)

2.4 Mean Squared Error (MSE) for three classical estimators

2.4.1 MSE for Maximum Likelihood estimation

$$\begin{aligned} \text{MSE}(\hat{\theta}_{MLE}) &= \text{MSE}_{\theta}\left(\frac{n}{p}\right) \\ &= \text{var}\left(\frac{n}{p}\right) + \left[E\left(\frac{n}{p} - \theta\right)\right]^2 \\ &= n^2 \text{var}\left(\frac{1}{p}\right) + \left[E\left(\frac{n}{p} - \theta\right)\right]^2 \end{aligned}$$

Where

$$\begin{aligned} \text{var}\left(\frac{1}{p}\right) &= E\left[\left(\frac{1}{p}\right)^2\right] - \left[E\left(\frac{1}{p}\right)\right]^2 \\ &= \frac{\theta^2}{(n-1)(n-2)} - \left[\frac{\theta}{n-1}\right]^2 \\ \text{var}\left(\frac{1}{p}\right) &= \frac{\theta^2}{(n-1)^2(n-2)} \end{aligned}$$

Consider,

$$\begin{aligned} \left[E\left(\frac{n}{p} - \theta\right)\right]^2 &= E\left[\left(\frac{n}{p}\right)^2 + \theta^2 - 2\theta\left(\frac{n}{p}\right)\right] \\ &= E\left(\frac{n}{p}\right)^2 + E(\theta^2) - 2\theta E\left(\frac{n}{p}\right) \\ &= n^2 \frac{\theta^2}{(n-1)^2} + \theta^2 - 2\theta n \frac{\theta}{n-1} \\ \left[E\left(\frac{n}{p} - \theta\right)\right]^2 &= \frac{\theta^2}{(n-1)^2} \\ \therefore \text{MSE}(\hat{\theta}_{MLE}) &= \frac{n^2 \theta^2}{(n-1)^2(n-2)} + \frac{\theta^2}{(n-1)^2} \\ \text{MSE}(\hat{\theta}_{MLE}) &= \frac{\theta^2(n+2)}{(n-1)(n-2)} \end{aligned}$$

(2.4.1)

2.4.2 Mean Squared Error for Uniformly Minimum Variance Unbiased Estimator

$$\begin{aligned} \text{MSE}(\hat{\theta}_{UMVUE}) &= \text{var}\left(\frac{n-1}{p}\right) + \left[E\left(\frac{n-1}{p}\right) - \theta\right]^2 \\ &= (n-1)^2 \text{var}\left(\frac{1}{p}\right) + \left[E\left(\frac{n-1}{p}\right) - \theta\right]^2 \quad \dots (2.4.5) \\ \text{MSE}(\hat{\theta}_{UMVUE}) &= (n-1)^2 \frac{\theta^2}{(n-1)^2(n-2)} \\ \text{MSE}(\hat{\theta}_{UMVUE}) &= \frac{\theta^2}{(n-2)} \end{aligned}$$

(2.4.2)

2.4.3 Mean squared error for Minimum Mean Square Error

$$\begin{aligned} \text{MSE}(\hat{\theta}_{MinMSE}) &= \text{var}\left(\frac{n-2}{p}\right) + \left[E\left(\frac{n-2}{p}\right) - \theta\right]^2 \\ &= (n-2)^2 \text{var}\left(\frac{1}{p}\right) + \left[E\left(\frac{n-2}{p}\right) - \theta\right]^2 \\ &= (n-2)^2 \frac{\theta^2}{(n-1)^2(n-2)} + \frac{\theta^2}{(n-1)^2} \\ &= \frac{\theta^2(n-2)}{(n-1)^2} + \frac{\theta^2}{(n-1)^2} \\ \text{MSE}(\hat{\theta}_{MinMSE}) &= \frac{\theta^2}{(n-1)} \end{aligned}$$

(2.4.3)

From equations (2.4.1), (2.4.2) and (2.4.3), we have

$$\text{MSE}(\hat{\theta}_{MinMSE}) \leq \text{MSE}(\hat{\theta}_{UMVUE}) \leq \text{MSE}(\hat{\theta}_{MLE})$$

(2.4.4)

From (2.4.4) it is observed that the Minimum Mean Squared Error (MinMSE) is the best estimator than the Maximum Likelihood Estimator (MLE) and the Uniformly Minimum Variance Unbiased Estimator (UMVUE).

2.5 Bayesian Estimation and Bayes Risk using non – informative prior

Bayesian estimation is an estimation of an unknown parameter θ that minimizes the expected loss for all observations x of X . The Bayes approach is an average case analysis by taking the average risk of an estimator for all the parameters involved in the distribution under study. Suppose we take the prior probability distribution π , on the parameter space ω then the average risk is defined as

$$R_{\pi}(\hat{\theta}) = E_{\theta, x}[L(\theta, \hat{\theta})]$$

and the Bayes risk for a prior π is the minimum that the average risk can achieve

$$\hat{R}_{\pi} = \inf_{\theta} [R_{\pi}(\hat{\theta})]$$

In Bayesian analysis, when prior knowledge about the parameter is not available, it is possible to make use of the non-informative prior. Since we have no knowledge on the parameters, we may use Jeffrey's prior which is the square root of the Fisher information matrix of parameter per observation.

The Jeffrey's prior is defined as

$$g_1(\theta) \propto \sqrt{I(\theta)} = b\sqrt{I(\theta)}$$

Where $I(\theta) = -nE\left(\frac{\partial^2 \log L}{\partial \theta^2}\right)$

Therefore

$$\begin{aligned} g(\theta) &= b\sqrt{-nE\left(\frac{\partial^2 \log L}{\partial \theta^2}\right)} \\ \log L &= \theta \log \alpha - (\theta + 1) \log x \\ \frac{d \log L}{d\theta} &= \frac{1}{\theta} + \log \alpha - \log x \\ \frac{d^2 \log L}{d\theta^2} &= -\frac{1}{\theta^2} \\ \therefore E\left(\frac{d^2 \log L}{d\theta^2}\right) &= -\frac{1}{\theta^2} \\ g(\theta) &= \frac{b}{\theta} \sqrt{n}, \theta > 0 \end{aligned}$$

Assuming that $t_1, t_2, \dots, t_n$ be the n independent observation which follows the Pareto Type I Distribution with probability density function, given in (1.1) the value of α is known and θ is the only unknown parameter, we shall obtain the posterior probability density function for θ using Jeffrey's prior distribution.

Posterior probability density (pdf) function for θ is

$$\begin{aligned} h(\theta/t_1, t_2, \dots, t_n) &= \frac{g(\theta)L(\theta; t_1, t_2, \dots, t_n)}{\int_0^{\infty} g(\theta)L(\theta; t_1, t_2, \dots, t_n)d\theta} \\ &= \frac{\left(\frac{1}{\theta}\right) b\sqrt{n} \cdot \theta^n e^{n\theta \log \alpha} e^{-(\theta+1)\sum \log x}}{\int_0^{\infty} \frac{1}{\theta} b\sqrt{n} \cdot \theta^n e^{n\theta \log \alpha} e^{-(\theta+1)\sum \log x} d\theta} \\ &= \frac{(\theta^{-1} \theta^n e^{n\theta \log \alpha - (\theta+1)\sum \log x})}{\int_0^{\infty} \theta^{-1} \theta^n e^{n\theta \log \alpha - (\theta+1)\sum \log x} d\theta} \\ &= \frac{(\theta^{n-1} e^{-\theta(-n \log \alpha + \sum \log x)})}{\int_0^{\infty} \theta^{n-1} e^{-\theta(-n \log \alpha + \sum \log x)} d\theta} \end{aligned}$$

Where,

$$\begin{aligned} p &= \sum_{i=1}^n \log(t_i/\alpha) \\ &= \frac{(\theta^{n-1} e^{-\theta p})}{\int_0^{\infty} \theta^{n-1} e^{-\theta p} d\theta} \end{aligned}$$

$$h(\theta/t_1, t_2, \dots, t_n) = \frac{p^n}{\Gamma(n)} \theta^{n-1} e^{-\theta p}$$

The posterior density function of Jeffrey's prior is

$$h(\theta/t_1, t_2, \dots, t_n) = \frac{p^n}{\Gamma(n)} \theta^{n-1} e^{-\theta p}$$

III. Bayes estimation and Bayes Risk under different loss function

3.1.1 Bayes estimator under the Squared Error Loss function (SELF)

The SELF is defined as

$$L(\hat{\theta}, \theta)_{SELF} = (\hat{\theta} - \theta)^2$$

The Bayes estimator under SELF is

$$\begin{aligned}\hat{\theta}_{SELF} &= E(\theta) = \int_0^\infty \theta h(\theta/t_1, t_2, \dots, t_n) d\theta \\ &= \frac{\left(\frac{1}{\theta}\right) b \sqrt{n} \theta^n e^{n\theta \log \alpha} e^{-(\theta+1)\Sigma \log x}}{\int_0^\infty \left(\frac{1}{\theta}\right) b \sqrt{n} \theta^n e^{n\theta \log \alpha} e^{-(\theta+1)\Sigma \log x} d\theta} \\ &= \frac{p^n}{\sqrt{n}} \cdot \frac{\sqrt{n+1}}{p^{n+1}} \\ \hat{\theta}_{SELF} &= E(\theta) = \frac{n}{p}\end{aligned}$$

(3.1.1)

3.1.2 Bayes Risk under Squared Error Loss Function

The Bayes risk $R(\theta, \hat{\theta})$ under SELF is defined as the expected loss under SELF

$$\begin{aligned}R(\theta, \hat{\theta}) &= [E(\theta) - \theta]^2 \\ \Rightarrow R(\theta, \hat{\theta}) &= \frac{n^2}{p^2} + \theta^2 - 2 \frac{n\theta}{p} \\ R(\theta, \hat{\theta}) &= E[L(\theta, \hat{\theta})] \\ &= E\left[\left(\frac{n}{p} - \theta\right)^2\right] \\ \Rightarrow R(\theta, \hat{\theta}) &= \frac{n^2}{p^2} + \theta^2 - 2 \frac{n\theta}{p}\end{aligned}$$

(3.1.2)

3.1.3 Bayes Estimation under Quadratic Loss Function

The quadratic loss function is defined as

$$L(\hat{\theta}, \theta) = \left(\frac{\theta - \hat{\theta}}{\theta}\right)^2 = \left(1 - \frac{\hat{\theta}}{\theta}\right)^2$$

Consider, the risk function $R_Q(\hat{\theta}, \theta)$ to estimate the parameter θ under the Quadratic Loss function is defined as

$$\begin{aligned}R_Q(\hat{\theta}, \theta) &= E\left(1 - \frac{\hat{\theta}}{\theta}\right)^2 \\ R_Q(\hat{\theta}, \theta) &= \int_0^\infty \left(1 - \frac{\hat{\theta}}{\theta}\right)^2 h(\theta/t_1, t_2, \dots, t_n) d\theta \\ \text{Let } \frac{dR_Q(\hat{\theta}, \theta)}{d\hat{\theta}} &= 0 \\ \Rightarrow 2 \int_0^\infty \left(1 - \frac{\hat{\theta}}{\theta}\right) \left(-\frac{1}{\theta}\right) h(\theta/t_1, t_2, \dots, t_n) d\theta \\ \Rightarrow \int_0^\infty \left(1 - \frac{\hat{\theta}}{\theta}\right)^2 \left(-\frac{1}{\theta}\right) h(\theta/t_1, t_2, \dots, t_n) d\theta \\ \Rightarrow \hat{\theta} \int_0^\infty \left(\frac{1}{\theta^2}\right) h(\theta/t) d\theta - \int_0^\infty \left(\frac{1}{\theta}\right) h(\theta/t) d\theta &= 0 \\ \Rightarrow \hat{\theta}_{QLF} &= \frac{E\left(\frac{1}{\theta}\right)}{\left(\frac{1}{\theta^2}\right)}\end{aligned}$$

(3.1.3)

Where

$$\begin{aligned}E\left(\frac{1}{\theta}\right) &= \int_0^\infty \left(\frac{1}{\theta}\right) h(\theta/t_1, t_2, \dots, t_n) d\theta \\ &= \int_0^\infty \left(\frac{1}{\theta}\right) \frac{p^n}{\sqrt{n}} \theta^{n-1} e^{-\theta p} d\theta\end{aligned}$$

$$\begin{aligned}
&= \frac{p^n \sqrt{n-1}}{\sqrt{n} p^{n-1}} \\
E\left(\frac{1}{\theta}\right) &= \frac{p}{n-1} \\
E\left(\frac{1}{\theta^2}\right) &= \int_0^\infty \left(\frac{1}{\theta^2}\right) h(\theta/t_1, t_2, \dots, t_n) d\theta \\
&= \int_0^\infty \left(\frac{1}{\theta^2}\right) \frac{p^n}{\sqrt{n}} \theta^{n-1} e^{-\theta p} d\theta \\
&= \frac{p^n}{\sqrt{n}} \int_0^\infty \left(\frac{1}{\theta^2}\right) \theta^{n-1} e^{-\theta p} d\theta \\
&= \frac{p^n \sqrt{n-2}}{\sqrt{n} p^{n-2}} \\
E\left(\frac{1}{\theta^2}\right) &= \frac{p^2}{(n-1)(n-2)} \\
\text{Substituting } E\left(\frac{1}{\theta}\right) \text{ and } E\left(\frac{1}{\theta^2}\right) \text{ in (3.1.2), we get} \\
\therefore \hat{\theta}_{QLF} &= \frac{(n-2)}{p} \\
(3.1.4)
\end{aligned}$$

3.1.4 Bayes Risk under Quadratic Loss Function

The Bayes risk $R(\theta, \hat{\theta})$ under the quadratic loss function is defined as the expected loss under QLF

$$\begin{aligned}
(\text{ie}) \quad R(\theta, \hat{\theta}) &= E[L(\theta, \hat{\theta})] \\
\Rightarrow \quad R(\theta, \hat{\theta}) &= \left(\frac{\theta - \hat{\theta}}{\theta}\right)^2 \\
\text{Where, } \hat{\theta} &= \frac{(n-2)}{p} \\
\Rightarrow \quad R(\theta, \hat{\theta}) &= 1 - \left(\frac{\frac{(n-2)}{p}}{\theta}\right)^2 \\
&= 1 - \left(\frac{(n-2)}{p\theta}\right)^2 \\
&= 1 + \left(\frac{(n-2)}{\theta p}\right)^2 - \left(\frac{(n-2)}{p\theta}\right) \\
R(\theta, \hat{\theta}) &= 1 + \left\{ \frac{(n-2)}{\theta p} \left[\frac{n-2-2\theta p}{\theta p} \right] \right\} \\
(3.1.5)
\end{aligned}$$

3.1.5 Bayes Estimation under Precautionary Loss Function

The Precautionary Loss Function is defined as

$$L(\hat{\theta}, \theta) = \frac{(\hat{\theta} - \theta)^2}{\hat{\theta}}$$

The Bayes estimator under a precautionary loss function is defined as

$$\begin{aligned}
\hat{\theta}_{PLE} &= [E(\theta^2)]^{\frac{1}{2}} \\
\text{Where, } E(\theta^2) &= \int_0^\infty \theta^2 h(\theta/t_1, t_2, \dots, t_n) d\theta \\
&= \int_0^\infty \theta^2 \frac{p^n}{\sqrt{n}} \theta^{n-1} e^{-\theta p} d\theta \\
&= \frac{p^n}{\sqrt{n}} \int_0^\infty \theta^{n+2-1} e^{-\theta p} d\theta
\end{aligned}$$

$$\begin{aligned}
&= \frac{p^n \sqrt{n+2}}{\sqrt{n} p^{n+2}} \\
E(\theta^2) &= \frac{1}{\sqrt{n}} \frac{(n+1)n\sqrt{n}}{p^2} \\
\therefore \hat{\theta}_{PLF} &= \left[\frac{n(n+1)}{p^2} \right]^{1/2} \\
(3.1.6)
\end{aligned}$$

3.1.6 Bayes Risk under Precautionary Loss Function

The Bayes Risk $R(\theta, \hat{\theta})$ under Precautionary Loss Function is defined as the expected loss under PLF

$$(ie) R(\theta, \hat{\theta}) = \frac{(\hat{\theta} - \theta)^2}{\hat{\theta}}, \text{ where } \hat{\theta} = \left[\frac{n(n+1)}{p^2} \right]^{1/2}$$

$$R(\theta, \hat{\theta}) = \frac{\left[\frac{\sqrt{n(n+1)}}{p} - \theta \right]^2}{\frac{\sqrt{n(n+1)}}{p}}$$

$$\begin{aligned} R(\theta, \hat{\theta}) &= \left[\frac{\sqrt{n(n+1)}}{p} - \theta \right]^2 X \frac{p}{\sqrt{n(n+1)}} \\ &= \left(\frac{\sqrt{n(n+1)}}{p} \right)^2 + \theta^2 - 2\theta \frac{\sqrt{n(n+1)}}{p} X \frac{p}{\sqrt{n(n+1)}} \\ \Rightarrow R(\theta, \hat{\theta}) &= \frac{[n(n+1)]^{1/2}}{p} + \frac{\theta^2 p}{[n(n+1)]^{1/2}} - 2\theta \end{aligned}$$

3.1.7 Bayes Estimation under Entropy Loss Function

The Entropy Loss Function is defined as

$$L(\delta) = b[\delta^p - p \log \delta - 1]$$

$$\text{Where } \delta = \frac{\hat{\theta}}{\theta}$$

Where minimum occurs at $\hat{\theta} = \theta$. Also the loss function $L(\delta)$ has been used in the original form having $p=1$.

$L(\delta)$ can be written as

$$L(\delta) = b[\delta - \log \delta - 1]; b > 0$$

The Bayes estimator under the entropy loss functions is given by

$$\hat{\theta}_{ELF} = \left[E \left(\frac{1}{\theta} \right) \right]^{-1} = \frac{1}{E \left(\frac{1}{\theta} \right)}$$

(3.1.7)

$$\begin{aligned} \text{Where } E \left(\frac{1}{\theta} \right) &= \int_0^\infty \left(\frac{1}{\theta} \right) h(\theta/t_1, t_2, \dots, t_n) d\theta \\ &= \int_0^\infty \left(\frac{1}{\theta} \right) \frac{p^n}{\sqrt{n}} \theta^{n-1} e^{-\theta p} d\theta \\ &= \frac{p^n}{\sqrt{n}} \int_0^\infty \theta^{-1} \theta^{n-1} e^{-\theta p} d\theta \\ &= \frac{p^n \sqrt{n-1}}{\sqrt{n} p^{n-1}} \\ E(1/\theta) &= \frac{p}{(n-1)} \\ \hat{\theta}_{ELF} &= \frac{n-1}{p} \end{aligned}$$

3.1.8 Bayes Risk under Entropy Loss Function

The Bayes risk $R(\theta, \hat{\theta})$ under entropy loss function is defined as the expected loss under ELF, which is given by

$$\begin{aligned} R(\theta, \hat{\theta}) &= E[L(\delta)] \\ \Rightarrow R(\theta, \hat{\theta}) &= b \left[\frac{\left(\frac{n-1}{p} \right)}{\theta} - \log_e \left(\frac{\left(\frac{n-1}{p} \right)}{\theta} \right) - 1 \right] \\ \Rightarrow R(\theta, \hat{\theta}) &= b \left[\frac{n-1}{\theta p} - \log_e \left(\frac{n-1}{\theta p} \right) - 1 \right] \end{aligned}$$

(3.1.7)

IV. Bayes estimation and Bayes risk using Informative prior

Assuming that θ has informative prior as exponential prior which takes the following form

$$g(\theta) = \frac{1}{\lambda} e^{-\theta/\lambda}; \theta, \lambda > 0$$

The posterior pdf of exponential prior is defined as

$$h(\theta/t) = \frac{L(t_1, t_2, \dots, t_n)g(\theta)}{\int_0^\infty L(t_1, t_2, \dots, t_n)g(\theta)d\theta}$$

$$= \frac{\theta^n e^{-(\theta+1)p} \frac{1}{\lambda} e^{-\theta/\lambda}}{\int_0^\infty \theta^n e^{-(\theta+1)p} \frac{1}{\lambda} e^{-\theta/\lambda} d\theta}$$

Consider,

$$\int_0^\infty \theta^n e^{-(\theta+1)p} \frac{1}{\lambda} e^{-\theta/\lambda} d\theta = \int_0^\infty \theta^n e^{-\theta p} e^{-p} \frac{1}{\lambda} e^{-\theta/\lambda} d\theta$$

$$= \int_0^\infty \theta^n e^{-\theta p - \theta/\lambda} \frac{1}{\lambda} e^{-p} d\theta$$

$$= \frac{e^{-p}}{\lambda} \frac{\sqrt{n+1}}{(p+1/\lambda)^{n+1}}$$

∴ The posterior pdf of exponential prior is

$$h(\theta/t_1, t_2, \dots, t_n) = \frac{(p+1/\lambda)^{n+1}}{\sqrt{n+1}} \theta^n e^{-\theta} (p+1/\lambda)$$

4.1 Bayes Estimation and Bayes Risk using informative prior under different loss function

4.1.1 Bayes estimation under Squared Error Loss Function (SELF)

The SELF is defined as $L(\theta, \hat{\theta}) = (\hat{\theta} - \theta)^2$

The Bayes estimator under squared error loss function is

$$\hat{\theta}_{SELF} = E(\theta) = \int_0^\infty \theta h(\theta/t_1, t_2, \dots, t_n) d\theta$$

$$= \int_0^\infty \theta \frac{(p+1/\lambda)^{n+1}}{\sqrt{n+1}} \theta^n e^{-\theta} d\theta$$

$$\hat{\theta}_{SELF} = E(\theta) = \frac{n+1}{p+1/\lambda}$$

(4.1.1)

4.1.2 Bayes risk under squared error loss function

The Bayes risk $R(\theta, \hat{\theta})$ under SELF is defined as the expected loss under SELF

$$(ie) R(\theta, \hat{\theta}) = E[L(\hat{\theta}, \theta)] = (\hat{\theta} - \theta)^2 = [E(\theta) - \theta]^2$$

$$= \left[\frac{n+1}{(p+1/\lambda)} - \theta \right]^2 \text{ using (4.1.1)}$$

$$R(\theta, \hat{\theta}) = \left\{ \frac{1}{(p+1/\lambda)^2} (n^2 + 1 + 2n) + \theta^2 - \frac{2\theta(n+1)}{(p+1/\lambda)} \right\}$$

(4.1.2)

4.1.3 Bayes Estimation under Quadratic loss function

The Quadratic loss function is defined as

$$L(\theta, \hat{\theta}) = \left(\frac{\theta - \hat{\theta}}{\theta} \right)^2 = \left(1 - \frac{\hat{\theta}}{\theta} \right)^2$$

Consider the risk function $R(\theta, \hat{\theta})$ to estimate the parameter θ under quadratic loss function

$$\text{where, } R(\theta, \hat{\theta}) = E \left(1 - \frac{\hat{\theta}}{\theta} \right)^2$$

$$= \int_0^\infty \left(1 - \frac{\hat{\theta}}{\theta} \right)^2 h(\theta/t_1, t_2, \dots, t_n) d\theta$$

$$\Rightarrow \hat{\theta}_{QLF} = \frac{E(1/\theta)}{E(1/\theta^2)}$$

Where,

$$E \left(\frac{1}{\theta} \right) = \int_0^\infty \left(\frac{1}{\theta} \right) h(\theta/t_1, t_2, \dots, t_n) d\theta$$

$$= \int_0^\infty \left(\frac{1}{\theta} \right) \frac{(p+1/\lambda)^{n+1}}{\sqrt{n+1}} \theta^n e^{-\theta} d\theta$$

$$= \frac{(p+1/\lambda)^{n+1}}{\sqrt{n+1}} \int_0^\infty \theta^{n-1} e^{-\theta} d\theta$$

$$E\left(\frac{1}{\theta}\right) = \frac{(p+1/\lambda)}{n}$$

Also,

$$\begin{aligned} E\left(\frac{1}{\theta^2}\right) &= \int_0^\infty \left(\frac{1}{\theta^2}\right) h(\theta/x) d\theta \\ &= \int_0^\infty \left(\frac{1}{\theta^2}\right) \frac{(p+1/\lambda)^{n+1}}{\sqrt{n+1}} \theta^n e^{-\theta(p+1/\lambda)} d\theta \\ &= \frac{(p+1/\lambda)^n (p+1/\lambda)}{\sqrt{n+1}} \frac{\sqrt{n+1}}{(p+1/\lambda)^{n+1}} \\ &= \frac{(p+1/\lambda)^2}{n(n-1)} \end{aligned}$$

∴ The Bayes estimator under quadratic loss function is

$$\hat{\theta}_{QLF} = \frac{(n-1)}{(p+1/\lambda)} \quad (4.1.3)$$

4.1.4 Bayes risk under the quadratic loss function

The Bayes risk $R(\theta, \hat{\theta})$ under Quadratic loss function defined as the expected loss under QLF,

$$R(\theta, \hat{\theta}) = E[L(\hat{\theta}, \theta)] = \left(\frac{\theta - \hat{\theta}}{\theta}\right)^2 = \left(1 - \frac{\hat{\theta}}{\theta}\right)^2$$

$$\text{Where } \hat{\theta} = \frac{(n-1)}{(p+1/\lambda)}$$

$$\begin{aligned} \Rightarrow R(\theta, \hat{\theta}) &= \left[1 - \frac{\left(\frac{n-1}{p+1/\lambda}\right)}{\theta}\right]^2 \\ \Rightarrow R(\theta, \hat{\theta}) &= 1 + \left[\frac{n-1}{\theta(p+1/\lambda)}\right]^2 - 2 \frac{n-1}{\theta(p+1/\lambda)} \\ &= 1 + \left\{ \frac{1}{\theta(p+1/\lambda)} \left[\frac{(n-1)^2}{\theta(p+1/\lambda)} - 2(n-1) \right] \right\} \end{aligned}$$

Therefore,

$$R_{QLF}(\theta, \hat{\theta}) = 1 + \left\{ \frac{1}{\theta(p+1/\lambda)} \left[\frac{n^2+1-2n-2n\theta(p+1/\lambda)+\theta(p+1/\lambda)}{\theta(p+1/\lambda)} \right] \right\} \quad (4.1.4)$$

4.1.5 Bayes Estimation under Precautionary loss function

The Precautionary loss function is defined as

$$L_{PLF}(\theta, \hat{\theta}) = \left(\frac{\hat{\theta} - \theta}{\theta}\right)^2 = \left(1 - \frac{\theta}{\hat{\theta}}\right)^2$$

Consider the risk function $R(\theta, \hat{\theta})$ to estimate the parameter θ under quadratic loss function

$$\text{Where, } R(\theta, \hat{\theta}) = E[L(\hat{\theta}, \theta)]$$

$$\Rightarrow R(\theta, \hat{\theta}) = \int_0^\infty L(\hat{\theta}, \theta) h(\theta/x) d\theta$$

$$\begin{aligned} &= \int_0^\infty \left(1 - \frac{\theta}{\hat{\theta}}\right) \left(-\frac{1}{\hat{\theta}^2}\right) h(\theta/x) d\theta = 0 \\ \Rightarrow \frac{1}{\hat{\theta}^3} \int_0^\infty \theta h(\theta/x) d\theta &= \frac{1}{\hat{\theta}^2} \int_0^\infty h(\theta/x) d\theta \end{aligned}$$

∴ The Bayes estimator under precautionary loss function is

$$\hat{\theta}_{PLF} = \frac{(n-1)}{(p+1/\lambda)} \quad (4.1.5)$$

4.1.6 Bayes Risk under Precautionary loss function

The Bayes risk $R(\theta, \hat{\theta})$ under precautionary loss function is defined as the expected loss under PLF

$$(\text{ie}) R(\theta, \hat{\theta}) = E(L(\theta, \hat{\theta})) = \frac{(\hat{\theta} - \theta)^2}{\hat{\theta}}$$

Where $\hat{\theta} = \frac{n+1}{p+\frac{1}{\lambda}}$

$$\therefore R(\theta, \hat{\theta}) = \frac{1}{\left(\frac{n+1}{p+\frac{1}{\lambda}}\right)} \left[\left(\frac{n+1}{p+\frac{1}{\lambda}}\right)^2 + \theta^2 - 2\theta \left(\frac{n+1}{p+\frac{1}{\lambda}}\right) \right] \quad (4.1.6)$$

4.1.7 Bayes Estimation under Entropy loss function

The entropy loss function is defined as

$$L_{BE}(\delta) = b[\delta^p - p \log \delta - 1]; b > 0$$

$$\text{Where } \delta = \frac{\hat{\theta}}{\theta}$$

Where minimum occurs at $\hat{\theta} = \theta$. Also the loss function $L(\delta)$ has been used in the original form having $p=1$.

$L(\delta)$ can be written as

$$L(\delta) = b[\delta - \log \delta - 1]; b > 0$$

The Bayes estimator under the entropy loss function is defined as

$$\begin{aligned} \hat{\theta}_{ELF} &= [E(1/\theta)]^{-1} = \frac{1}{E(1/\theta)} \\ &= \frac{1}{\frac{(p+\frac{1}{\lambda})}{n}} \end{aligned}$$

$\therefore$ The Bayes estimator under entropy loss function is

$$\hat{\theta}_{ELF} = \frac{n}{(p+\frac{1}{\lambda})} \quad (4.1.7)$$

4.1.8 Bayes risk under Entropy Loss Function

The Bayes risk $R(\theta, \hat{\theta})$ under entropy loss function is defined as the expected loss under ELF

$$(ie) R(\theta, \hat{\theta}) = E(L(\delta)) = E[b[\delta - \log \delta - 1]]; b > 0$$

$$\text{Where } \delta = \frac{\hat{\theta}}{\theta} \text{ and } \hat{\theta} = \frac{n}{(p+\frac{1}{\lambda})}$$

$$\begin{aligned} \therefore R(\theta, \hat{\theta}) &= b \left[\frac{\left(\frac{n}{(p+\frac{1}{\lambda})}\right)}{\theta} - \log_e \left(\frac{\left(\frac{n}{(p+\frac{1}{\lambda})}\right)}{\theta} \right) - 1 \right] \\ R(\theta, \hat{\theta}) &= b \left[\frac{n}{\theta(p+\frac{1}{\lambda})} - \log_e \left(\frac{n}{\theta(p+\frac{1}{\lambda})} \right) - 1 \right] \end{aligned} \quad (4.1.8)$$

V. Result and discussion

5.1. Comparison of classical estimation

In this study, we choose a sample size of $n=25, 50$ and 100 to represent the small median and large data set. The classical estimation of the shape parameter of the Pareto type I distribution using MLE, UMVUE and Minimum mean square error were obtained by using simulation technique and presented in table -5.1

Table 5.1 Classical estimation of the shape parameter $\theta = 0.1$ and Shape parameter $\alpha = 0.2$ & 0.4

Form the table 5.1, it is observed that MiniMSE is the best among the other proposed estimators such as MLE and UMVUE.

Distribution	Sample size (n)	shape=0.1					
		scale=0.2			scale=0.4		
		MLE	UMVUE	MINMSE	MLE	UMVUE	MINMSE
PARETO	25	0.005472	0.004938	0.000768	0.019878	0.0188798	0.0030587
	50	0.001739	0.001667	0.0000815	0.006957	0.0066667	0.0003261
	75	0.000833	0.000816	0.0000180	0.003333	0.0032653	0.00007225

5.2 Bayes estimation and Bayes Risk of the shape parameter of Pareto type – I distribution

We choose a sample of size $n = 25, 50, 75$ and 100 to represent small, medium and large data set. The Bayes risk of the shape parameter for Pareto type – I distribution is estimated using non – informative prior (Jefferey's prior) and informative prior (Exponential prior) under different loss functions. The value of shape parameter $\theta = 1, 2, 3$ and $\alpha = 4, 5, 6$. The results are obtained through simulation technique and presented in the table form (5.2.1) to (5.2.8).

5.2 Bayes estimation and Bayes Risk of the shape parameter of Pareto type – I distribution

Table. 5.2.1 Bayes Estimation under Square Error Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	2.481 4	2.461 2	2.210 2	1.961 2	1.956 2	1.946 2	1.934 8	1.924 8	1.9116
	50	2.322 4	2.321 2	2.110 2	1.954 2	1.944 2	1.931 2	1.922 8	1.911 2	1.9002
	75	1.981 2	1.974 2	1.961 2	1.943 2	1.933 2	1.922 0	1.911 8	1.900 2	1.8998
	100	1.971 0	1.962 0	1.954 2	1.931 0	1.922 0	1.911 0	1.900 8	1.890 2	1.8848
Exponential	25	2.311 0	2.310 1	2.101 2	1.953 2	1.944 2	1.936 8	1.924 6	1.911 2	1.9002
	50	2.228 4	2.210 2	2.000 1	1.941 0	1.933 6	1.922 8	1.911 2	1.900 6	1.8842
	75	1.972 0	1.961 2	1.934 2	1.932 0	1.922 0	1.911 6	1.900 2	1.900 0	1.8762
	100	1.962 0	1.954 2	1.921 0	1.922 0	1.911 0	1.900 6	1.900 0	1.890 0	1.8662

Table.5.2.2. Bayes Risk under Square Error Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	1.543 2	1.521 0	0.9810	0.971 0	0.981 0	0.9712	1.101 2	0.901 2	0.894 2
	50	1.493 0	1.481 0	0.9772	0.961 2	0.977 2	0.9662	1.091 2	0.894 9	0.884 0
	75	0.981 0	0.978 0	0.9612	0.991 0	0.961 0	0.9542	0.984 2	0.874 2	0.844 2

	100	0.971 0	0.961 0	0.9510	0.954 2	0.951 2	0.9442	0.974 0	0.864 0	0.831 2
Exponential	25	1.423 3	1.514 2	0.9712	0.961 2	0.951 2	0.9412	1.100 2	0.801 2	0.821 2
	50	1.412 0	1.477 2	0.9612	0.951 2	0.941 2	0.9312	1.081 2	0.721 2	0.811 1
	75	0.977 2	0.961 2	0.9512	0.941 2	0.931 0	0.9212	0.971 2	0.731 0	0.731 2
	100	0.961 2	0.951 2	0.9412	0.931 2	0.921 2	0.9010	0.921 2	0.721 0	0.721 2

Table 5.2.3. Bayes Estimation under Quadratic Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	1.421 2	1.314 8	1.3046	1.284 6	1.261 6	1.2546	1.274 2	1.254 2	1.2462
	50	1.401 0	1.284 6	1.2747	1.261 2	1.241 2	1.2462	1.264 2	1.244 2	1.2362
	75	1.214 1	1.114 2	1.1040	1.094 2	1.009 2	1.0082	1.094 1	1.009 2	1.0008
	100	0.984 6	0.976 8	0.9662	0.954 6	0.941 2	0.9316	0.944 2	0.931 2	0.9210
Exponential	25	1.314 2	1.284 6	1.2762	1.261 2	1.246 2	1.2446	1.260 0	1.249 2	1.2348
	50	1.304 0	1.116 8	1.2662	1.254 2	1.224 6	1.2146	1.254 9	1.239 2	1.2228
	75	1.184 2	1.104 0	1.0042	1.002 2	1.008 5	1.0068	1.009 8	1.009 0	1.0007
	100	0.976 8	0.966 8	0.9556	0.942 6	0.934 6	0.9262	0.931 2	0.921 2	0.9100

Table. 5.2.4. Bayes Risk under Quadratic Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	0.821 2	0.792 1	0.771 0	0.692 1	0.654 2	0.621 0	0.610 9	0.591 1	0.4212
	50	0.801 0	0.781 2	0.761 0	0.682 1	0.621 2	0.611 0	0.590 2	0.571 2	0.4102
	75	0.701 2	0.701 2	0.691 2	0.672 1	0.610 0	0.592 2	0.580 1	0.561 0	0.4012
	100	0.691 2	0.681 0	0.671 0	0.662 9	0.592 0	0.582 0	0.570 2	0.551 5	0.4012
Exponential	25	0.731 2	0.726 2	0.710 2	0.691 2	0.599 2	0.588 9	0.591 2	0.581 2	0.4010
	50	0.721 9	0.716 9	0.704 2	0.680 2	0.597 2	0.581 2	0.566 2	0.561 2	0.3912

	75	0.611 2	0.700 2	0.671 2	0.651 2	0.581 2	0.571 9	0.551 2	0.541 2	0.3812
	100	0.604 0	0.671 2	0.670 9	0.650 1	0.570 1	0.566 2	0.541 0	0.531 0	0.3712

Table. 5.2.5. Bayes Estimation under Precautionary Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	2.984 2	1.984 2	1.974 2	1.961 8	1.956 2	1.946 8	1.938 1	1.921 8	1.901 0
	50	1.994 2	1.974 8	1.961 2	1.954 2	1.944 2	1.934 3	1.921 6	1.911 8	1.891 8
	75	1.974 2	1.964 2	1.954 2	1.934 6	1.931 8	1.921 2	1.901 8	1.891 8	1.871 8
	100	1.824 2	1.804 2	1.804 6	1.708 0	1.801 0	1.796 8	1.791 8	1.881 2	1.870 8
Exponential	25	2.164 2	1.976 2	1.964 6	1.954 2	1.946 8	1.931 7	1.921 8	1.901 0	1.891 2
	50	1.994 6	1.966 0	1.954 2	1.931 2	1.931 3	1.921 8	1.911 0	1.891 0	1.881 2
	75	1.921 6	1.954 2	1.944 2	1.921 2	1.924 8	1.911 8	1.891 2	1.881 2	1.876 2
	100	1.784 2	1.800 2	1.800 0	1.798 0	1.998 0	1.901 8	1.882 1	1.876 2	1.866 2

Table. 5.2.6. Bayes Risk under Precautionary Loss Function for different values of θ and α

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	1.412 1	1.212 2	0.981 2	0.961 2	0.921 2	0.901 2	0.891 2	0.791 2	0.761 0
	50	1.210 2	1.091 2	0.974 2	0.954 2	0.901 3	0.891 2	0.881 2	0.781 0	0.751 2
	75	0.981 2	0.971 2	0.954 1	0.921 2	0.891 2	0.881 0	0.871 0	0.778 0	0.741 2
	100	0.971 2	0.961 2	0.921 2	0.901 0	0.881 0	0.871 2	0.816 1	0.761 0	0.731 0
Exponential	25	1.214 1	1.201 0	0.971 2	0.960 0	0.901 2	0.891 0	0.871 0	0.731 2	0.751 0
	50	1.122 0	1.009 0	0.964 0	0.951 2	0.891 2	0.881 0	0.860 9	0.721 2	0.741 2
	75	0.974 2	0.961 2	0.954 0	0.941 0	0.881 2	0.871 2	0.851 9	0.701 0	0.731 0
	100	0.961 2	0.951 2	0.941 2	0.931 2	0.871 2	0.861 2	0.844 0	0.691 2	0.721 2

Table.5.2.7. Bayes Estimation under Entropy Loss Function for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	2.921 2	2.901 2	2.891 2	2.681 2	2.521 2	2.461 2	2.481 2	2.471 2	2.451 2
	50	2.901 2	2.896 2	2.791 2	2.581 2	2.481 2	2.321 2	2.471 2	2.461 2	2.122 6
	75	2.891 2	2.791 2	2.681 2	2.421 2	2.331 2	2.311 0	2.301 0	2.300 0	2.281 2
	100	1.981 2	1.971 2	1.961 2	1.954 2	1.944 6	1.931 2	1.921 2	1.911 6	1.901 2
Exponential	25	2.681 2	2.891 2	2.521 2	2.461 2	2.481 2	2.431 2	2.301 2	2.311 2	2.330 0
	50	2.521 2	2.791 2	2.461 2	2.521 2	2.461 2	2.301 2	2.201 2	2.121 2	2.101 0
	75	2.561 2	2.681 0	2.321 2	2.401 2	2.301 2	2.281 2	2.101 0	2.000 9	1.988 2
	100	1.961 2	1.961 0	1.951 2	1.944 2	1.931 2	1.921 2	1.901 2	1.921 2	1.911 0

Table. 5.2.8. Bayes Risk under Entropy Loss for different values of θ and α .

Prior	Sample Size	$\theta = 1$			$\theta = 2$			$\theta = 3$		
		$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$	$\alpha = 4$	$\alpha = 5$	$\alpha = 6$
Jeffrey's	25	1.9212	1.9010	1.9000	1.9101	1.9010	0.9812	0.9712	0.9512	0.9410
	50	1.9810	1.9701	1.9610	1.9000	1.8910	0.9712	0.9612	0.9412	0.9312
	75	0.9912	0.9812	0.9712	0.8910	0.9819	0.9612	0.9512	0.9312	0.9212
	100	0.9810	0.9710	0.9610	0.8801	0.8712	0.9510	0.9401	0.9212	0.9000
Exponential	25	0.9812	0.9610	0.8912	0.8819	0.8712	0.8610	0.8512	0.8410	0.8310
	50	0.9712	0.9509	0.8812	0.8712	0.8610	0.8510	0.8412	0.8310	0.8210
	75	0.9610	0.9412	0.8712	0.8612	0.8512	0.8412	0.8310	0.8212	0.8110
	100	0.9510	0.9312	0.8612	0.8510	0.8412	0.8313	0.8220	0.8110	0.8000

VI. Conclusions

In this study, we obtained the classical estimators such as MLE, UMVUE and MiniMSE and the Bayes risk of shape parameter of Pareto type –I distribution using non-informative and informative priors under various loss functions through simulation technique. By comparing the classical estimation, the MiniMSE is the best one among the all others namely MLE & UMVUE. The Bayes risk of the shape parameter under QLF is the least one among all other loss functions, which are proposed for this study. It is also observed that when the sample size is increased, the Bayes risk is decreased. Finally, it is found that the MiniMSE is the best one of the proposed classical estimators and also found that the performance of bayes risk of the shape parameter of Pareto type –I model using informative prior under Quadratic loss function is minimum than the other proposed loss functions.

VII. References

1. Al Omari Mohammed Ahmed; Noor Akma Ibrahim; Comparison of the Bayesian and Maximum Likelihood Estimation for Weibull Distribution. Journal of Mathematics and Statistics 6 (2):100-104 (2010).

2. Al Omari Mohammed Ahmed; Noor Akma Ibrahim; Bayesian Survival Estimator for Weibull Distribution with Censored Data. Journal of Applied Sciences(2011).DOI: 10.3923 / jas.2011.393.396
3. Gaurav Shukla,Umaesh Chandra; Vinod Kumar; Bayes Estimation of the Reliability function of Pareto Distribution under three different loss functions. Journal of Reliability and Statistical Studies, Vol,149-176 (2020).
4. Kawsar Fatima; S.P. Ahmad; Bayesian Approach in Estimation of shape parameter of the Exponentiated moment exponential distribution. Journal of statistical Theory and Applications, Vol,No.2 359 -374 (2018).
5. R.K.Radha; Bayesian Analysis of Exponential Distribution using informative prior (IJSR), ISSN (online): 2319-7064 (2015).
6. Sanku day; Sudhansu S.Maiti; Bayesian estimation of the parameter of Rayleigh distribution under the extended Jeffery's prior. J.App. Stat. Anul, Vol. 5.Issue1,44-59 (2012) DOI 10.1285/i20705948v5nlp44.
7. Arnold,B.C.,2015. PARETO DISTRIBTUON, Second Edition, Monograph on Statistics and Applied Probability, CRC Press, New York .
8. Gupta.S.C and Kappor, V.K, 2000. Fundamentals of Mathematical Statistics, Sultan Chand and Sons. Educational Publishers, New Delhi.
9. Rao, C.R. 1973.LINEAR STATISTICAL INFERENCE AND ITS APPLICATIONS, Willey Eastern, New Delhi.