

A New Asymmetric Loss Function: Estimation of Parameter of Exponential Distribution

Dinesh Kumar, Pradip Kumar*, Sanjay Kumar Singh and Umesh Singh

Department of Statistics, Banaras Hindu University, Uttar Pradesh- 221005, India

Received: 29 Oct. 2018,, Revised: 23 Nov. 2018, Accepted: 17 Dec. 2018

Published online: 1 Jan. 2019

Abstract: A new asymmetric loss function, as a generalization of squared error loss function (SELF) have been proposed. For application point of view, exponential distribution with scale parameter θ have been considered and Bayes estimators of θ are obtained under SELF, LINEX loss function, General entropy Loss Function (GELF) and proposed loss function. The simulation study is carried out to compare their performances (in sense of having smallest posterior risks).

Keywords: Bayes estimators, symmetric and asymmetric loss functions, posterior risk.

1 Introduction

An important element in the Bayesian paradigm of estimation is the specification of suitable loss function depending on the situations. Initially, the loss function used was SELF, which is appropriate in the case when over estimation and under estimation of equal magnitudes are of equal importance. But in the situation, when over estimation is more serious than under estimation of the same magnitude and vice-versa, the use of SELF may lead to a wrong decision. For example, in dam construction, an under estimation of the height of the dam is usually much more serious than an over estimation. To sort-out this problems, Varian (1975) introduced a very useful asymmetric loss function, in his applied study of real estate assessment and it is called LINEX loss function. It rises approximately exponentially on one side of zero and approximately linearly on the other side of zero. However, use of these loss functions is suitable for estimation of location parameters only. For estimation of scale parameter, the suggested alternatives were weighted SELF and modified LINEX loss function, which are asymmetric loss functions. However, a drawback of weighted SELF is that it always penalizes over estimation heavily than under estimation and behavior of modified LINEX loss function is same as that of LINEX loss function. Several authors have used above loss functions to estimate the parameters of a number of distributions. For example, Asgharzadeh, and Farsipour (2008) discussed estimation of exponential mean time to failure with the help of symmetric and different asymmetric loss functions, Dey and Maiti (2010) have considered Maxwell distribution and derived Bayes estimators of its parameters using informative as well as non-informative priors under different loss functions. Kazmi et al. (2012) have obtained Bayes estimators of the parameters of Maxwell distribution under different loss functions and they have also carried out simulation study to know the performance of the considered Bayes estimators in terms of their posterior risks. Singh et al. (2013) have considered (Inverse Weibull Distribution) IWD and derived MLE and Bayes estimators of its parameters under SELF and GELF using gamma as informative prior.

In the present piece of work, we have introduced a new asymmetric loss function and we will call it as Quadratic Type Squared Error Loss Function (QTSELF) and we defined it as,

$$L_{QT}(\delta, \theta) = (\delta - \theta)^2 e^{c(\delta - \theta)} \quad (1)$$

where $c \in \mathbb{R}$ is the loss parameter. When $c > 0$, over estimation causes more serious than under estimation and when $c < 0$, under estimation is more serious than over estimation. For $c = 0$, it reduces to SELF. The seriousness can be seen in the Figure 1 for the fixed positive and negative values of $c = \pm 0.5$.

* Corresponding author e-mail: pradipkumar2612@gmail.com

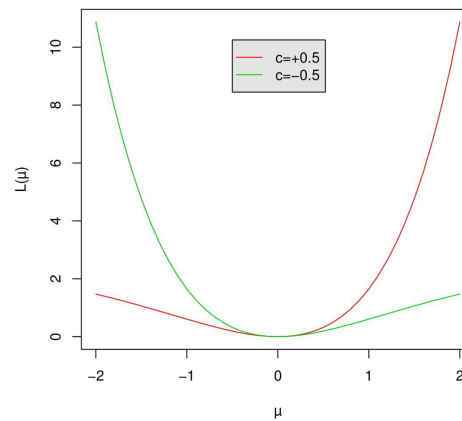


Fig. 1: Plots of $L(\mu) = L_{QT}(\delta, \theta)$, $\mu = \delta - \theta$ for fixed $c = \pm 0.5$.

Now the Bayes estimator of θ under the proposed loss function (1) is the decision rule δ which minimizes the corresponding Bayes risk, viz

$$R_{QT}^B(\delta, \theta) = \int_{\underline{x}} \int_{\theta} L_{QT}(\delta, \theta) h(\theta | \underline{x}) m(\underline{x}) d\theta d\underline{x}$$

and its minimization is same as that of the following integral

$$R_{QT}(\delta, \theta) = \int_{\theta} L_{QT}(\delta, \theta) h(\theta | \underline{x}) d\theta \quad (2)$$

where, $R_{QT}(\delta, \theta)$ is the posterior risk under the proposed loss function (1).

Now, by the principle of maxima and minima, for the optimum value(s) of $R_{QT}(\delta, \theta)$ w.r.t. δ , we must have,

$$\begin{aligned} \frac{\partial}{\partial \delta} R_{QT}(\delta, \theta) &= 0 \\ \Rightarrow \frac{\partial}{\partial \delta} \int_{\theta} (\delta - \theta)^2 e^{c(\delta - \theta)} h(\theta | \underline{x}) d\theta &= 0 \end{aligned}$$

After simplification, it reduces to

$$A\delta^2 + B\delta + C = 0 \quad (3)$$

where,

$$\begin{aligned} A &= cE(e^{-c\theta} | \underline{x}) \\ B &= 2[E(e^{-c\theta} | \underline{x}) - cE(\theta e^{-c\theta} | \underline{x})] \\ C &= cE(\theta^2 e^{-c\theta} | \underline{x}) - 2E(\theta e^{-c\theta} | \underline{x}) \end{aligned}$$

Now, as (3) is a quadratic equation in δ , so that Shri Dharacharya method can be applied to solve it. If $B^2 - 4AC < 0$, then any real solution of (3) can not be obtained and consequently we go for re-generation of the data $\underline{x}$. On the other hand, if $B^2 - 4AC = 0$, then unique solution of it will be there, say it is $\delta = \delta_0 = -\frac{B}{2A}$ and it will be Bayes estimator of θ provided $\delta_0, \theta \in \Theta$. Finally, if $B^2 - 4AC > 0$, then we will have two real solutions, say $\delta = \delta_1, \delta_2$ of (3) and if $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_1} > 0$ and $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_2} < 0$, then δ_1 will be the desired Bayes estimator of θ ; if $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_1} < 0$ and $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_2} > 0$, then δ_2 will be the Bayes estimator of θ ; if

$\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_1} > 0$ and $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_2} > 0$, then the decision rule for which $R(\delta, \theta)$ is the most least will be the Bayes estimator of θ ; lastly if $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_1} < 0$ and $\left[\frac{\partial^2}{\partial \delta^2} R_{QT}(\delta, \theta)\right]_{\delta=\delta_2} < 0$, then we will again go to the re-generation of the sample $\underline{X}$.

If $\hat{\delta}_{QT}$ be the Bayes estimator of θ , thus obtained, then its posterior risk will be given by

$$R_{QT}(\hat{\delta}_{QT}, \theta) = e^{c\hat{\delta}_{QT}} \left[\hat{\delta}_{QT}^2 E(e^{-c\theta} | \underline{X}) + E(\theta^2 e^{-c\theta} | \underline{X}) - 2\hat{\delta}_{QT} E(\theta e^{-c\theta} | \underline{X}) \right] \quad (4)$$

Let us consider the appropriate model for the given data $\underline{X} = (x_1, x_2, x_3, \dots, x_n)$ is exponential distribution with probability density function (pdf)

$$f(x|\theta) = \theta e^{-\theta x}; \quad x > 0, \quad \theta > 0 \quad (5)$$

and hence the joint pdf of $x_1, x_2, x_3, \dots, x_n$ i.e. pdf of $\underline{X}$ is given by

$$\begin{aligned} L(\underline{X}|\theta) &= \prod_{i=1}^n f(x_i|\theta) \\ &= \theta^n e^{-\theta \sum_{i=1}^n x_i} \end{aligned} \quad (6)$$

2 Different Priors of θ and Corresponding Posteriors:

The Bayesian deduction requires appropriate choice of priors for the parameters. Arnold and Press (1983) pointed out that, from Bayesian viewpoint, there is clearly no way in which one can say that one prior is better than any other. The prior information is a totally subjective assessment of an expert before any data has been observed. Therefore, it is very difficult to consider the informative prior due to assessment of their hyper-parameters. Also, Berger (1985) argues that when information is not in compact form the Bayesian analysis using non-informative prior or single most suitable consideration. So, here we consider two non-informative (the uniform and the Jeffreys) priors and one informative prior as the gamma prior.

2.1 Uniform Prior:

Let us consider prior for θ as standard uniform distribution having pdf is

$$\pi_U(\theta) \propto 1 \quad ; 0 < \theta < \infty \quad (7)$$

and hence the corresponding posterior distribution has the following pdf

$$\begin{aligned} h_U(\theta|\underline{X}) &\propto L(\underline{X}|\theta) \pi_U(\theta) \\ &\propto \theta^n e^{-\theta \sum_{i=1}^n x_i} \\ &\propto \theta^n e^{-\theta \sum_{i=1}^n x_i} \\ &= \frac{\theta^n e^{-\theta \sum_{i=1}^n x_i}}{\int_0^\infty \theta^n e^{-\theta \sum_{i=1}^n x_i} d\theta} \\ h_U(\theta|\underline{X}) &= \frac{(\sum_{i=1}^n x_i)^{n+1}}{\Gamma(n+1)} \theta^n e^{-\theta \sum_{i=1}^n x_i} \end{aligned} \quad (8)$$

which is obviously gamma distribution with parameters $n+1$ and $\sum_{i=1}^n x_i$.

2.2 Jeffreys Prior:

The most widely used method to choose non-informative prior is suggested by Jeffreys (1961) and is termed as Jeffreys prior, which is given by

$$\pi_J(\theta) = [I(\theta)]^{1/2} \quad (9)$$

where $I(\theta)$ is Fisher information and is given by

$$I(\theta) = -E \left[\frac{\partial^2}{\partial \theta^2} \ln L(\underline{x}; \theta) \right] \quad (10)$$

Using (6) in (10), we get

$$\begin{aligned} I(\theta) &= -E \left[\frac{\partial^2}{\partial \theta^2} \left(n \ln \theta - \theta \sum_{i=1}^n x_i \right) \right] \\ &= \frac{n}{\theta^2} \end{aligned} \quad (11)$$

using (11) in (9), we get

$$\pi_J(\theta) = \frac{\sqrt{n}}{\theta} \quad (12)$$

and hence the corresponding posterior pdf of θ given $\underline{x}$ is obtained as follows

$$\begin{aligned} h_J(\theta|\underline{x}) &= \frac{L(\underline{x}; \theta) \pi_J(\theta)}{\int_0^\infty L(\underline{x}; \theta) \pi_J(\theta) d\theta} \\ &= \frac{(\sum_{i=1}^n x_i)^n}{\Gamma(n)} \theta^{n-1} e^{-\sum_{i=1}^n x_i \theta} \end{aligned} \quad (13)$$

which is again gamma distribution with parameter n and $\sum_{i=1}^n x_i$.

2.3 Conjugate Prior:

Let the prior for θ be gamma distribution with the parameters a and b having pdf as

$$\pi_C(\theta) = \frac{b^a}{\Gamma(a)} \theta^{a-1} e^{-b\theta} \quad ; \theta > 0, a, b > 0 \quad (14)$$

and hence the posterior pdf of θ given $\underline{x}$ is obtained as follows

$$\begin{aligned} h_C(\theta|\underline{x}) &= \frac{L(\underline{x}; \theta) \pi_C(\theta)}{\int_0^\infty L(\underline{x}; \theta) \pi_C(\theta) d\theta} \\ &= \frac{(b + \sum_{i=1}^n x_i)^{n+a}}{\Gamma(n+a)} \theta^{n+a-1} e^{-(b + \sum_{i=1}^n x_i) \theta} \end{aligned} \quad (15)$$

which is again gamma distribution with parameter $n+a$ and $b + \sum_{i=1}^n x_i$ i.e. the prior gamma(a, b) considered for θ is a conjugate prior.

Now, to have an idea about the shapes of the conjugate prior and corresponding posterior pdf's for different confidence levels in the guessed value of θ as its true value, we randomly generate a sample from exponential distribution for fixed values $n = 20$, $\theta = 2$, $M=2$, $V=0.1$ (showing a higher confidence in the guess value) and $V= 150$ (showing a weak confidence in the guess value). The sample thus generated is,

$\underline{X} = (0.44454268, 0.42825063, 0.90121861, 0.02375386, 0.24978369, 0.0216787, 0.21266969, 0.46129670, 1.47181423, 0.45747565, 1.58447415, 0.53206546, 0.08320718, 0.06530584, 0.10585582, 0.60946017, 0.30534982, 0.25279858, 1.32600740, 0.28972998)$

The graphs are shown below-

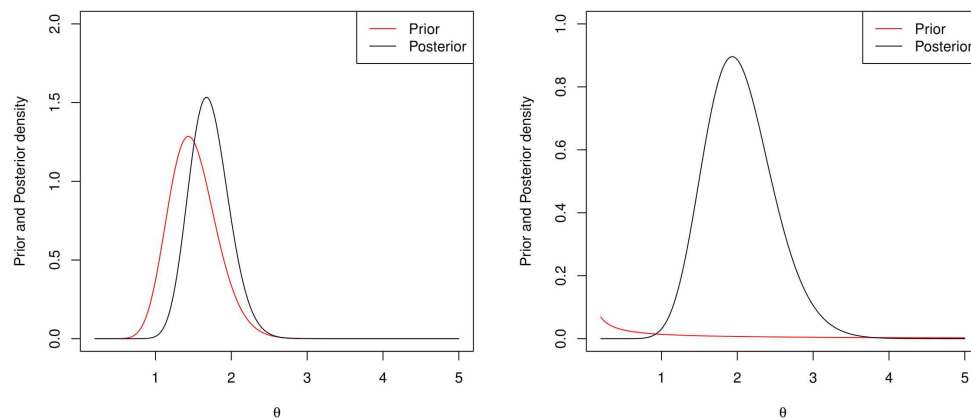


Fig. 2: Conjugate Prior and the corresponding Posterior pdf of θ for $n=20$, $\theta=2$, $M=2$, and $V=0.1, 150$ respectively.

3 Bayes Estimators and Posterior Risks Under Different Loss Functions

In decision theory, the loss criterion is specified in order to obtain best estimator for which Bayes risk is minimum. The simplest form of loss function is squared error loss function (SELF) which is suitable when over estimation and under estimation of same magnitudes have equal importance. However, in most of the real situations, this assumption is not possible. Sometime, over estimation is more serious than under estimation and vice-versa. Motivated to this, several asymmetric loss function has been introduced in literature, and we have also proposed a new loss function as defined in Section 1. Here, we will derive Bayes estimators of θ under SELF, LINEX, GELF and our proposed asymmetric loss function and compare their performance in term of Posterior risks.

3.1 Squared Error Loss Function

The squared error loss function was proposed by Legendre (1805) and Gauss (1810) to develop least-square theory. The squared error loss function (SELF) is defined as

$$L_S(\delta, \theta) = (\delta - \theta)^2 \quad (16)$$

Now, we consider Bayes estimator under SELF to search a decision rule $\hat{\delta}_S$ for which Bayes risk is minimum i.e.

$$R_S(\delta, \theta) = \int_{\theta} (\delta - \theta)^2 h(\theta|\underline{X}) d\theta \quad (17)$$

And the Bayes estimator of θ under above loss function is the decision rule δ which minimizes the posterior risk (17) and same is given by

$$\hat{\delta}_S = E_{\theta}(\theta|\underline{X}) \quad (18)$$

which is nothing but the posterior mean.

Putting the value of (18) in (17), we get the corresponding posterior risk as follows,

$$R_S(\hat{\delta}_S, \theta) = \hat{\delta}_S^2 - 2\hat{\delta}_S E(\theta|\underline{X}) + E(\theta^2|\underline{X}) \quad (19)$$

Equation (19) is equal to the posterior variance. Thus, under SELF posterior risk of Bayes estimator is equal to posterior variance.

The Bayes estimators & corresponding posterior risks of the parameter θ of $\exp(\theta)$ -distribution under SELF $L_S(\delta, \theta)$ for different priors, viz uniform, Jefferys & conjugate priors are given by

$$\hat{\delta}_{SU} = \frac{n+1}{\sum_{i=1}^n x_i}, \quad \hat{\delta}_{SJ} = \frac{n}{\sum_{i=1}^n x_i}, \quad \hat{\delta}_{SC} = \frac{n+a}{(b+\sum_{i=1}^n x_i)} \quad \text{and}$$

$$R_S(\hat{\delta}_{SU}, \theta) = \frac{n+1}{(\sum_{i=1}^n x_i)^2}, \quad R_S(\hat{\delta}_{SJ}, \theta) = \frac{n}{(\sum_{i=1}^n x_i)^2}, \quad R_S(\hat{\delta}_{SC}, \theta) = \frac{n+a}{(b+\sum_{i=1}^n x_i)^2}.$$

3.2 LINEX Loss Function

A symmetric loss function assumes that positive and negative errors of the same magnitudes are equally serious. However, in some estimation problem such an assumption may be inappropriate and several loss functions have been introduced to handle such a problem, one of them is LINEX loss function. It was introduced by Klebnav (1972) and used by Varian (1975) in the context of real estate assessment. Zellner (1986) and Varian (1975) have discussed its several properties. The LINEX loss function is defined as

$$L_L(\delta, \theta) = b_1[e^{c_1(\delta-\theta)} - c_1(\delta - \theta) - 1] \quad (20)$$

Where, $b_1 > 0$ and $c_1 \neq 0$.

The constant c_1 determines the shape of loss function. It is to be noted that when $c_1 > 0$, over estimation causes is more serious than under estimation and when $c_1 < 0$, under estimation more serious than over estimation. Now, the posterior risk corresponding to this loss function is given by

$$R_L(\delta, \theta) = \int_{\theta} L_L(\delta, \theta) h(\theta|\underline{x}) d\theta \quad (21)$$

the Bayes estimator of the parameter θ under this loss function is

$$\hat{\delta}_L = \frac{-1}{c_1} \ln[E(e^{-c_1\theta}|\underline{x})] \quad (22)$$

From (22), we get Bayes estimator under different priors for LINEX loss function for different priors are given below,

$$\hat{\delta}_{LU} = \frac{n+1}{c_1} \ln\left(1 + \frac{c_1}{\sum_{i=1}^n x_i}\right), \hat{\delta}_{LJ} = \frac{n}{c_1} \ln\left(1 + \frac{c_1}{\sum_{i=1}^n x_i}\right) \text{ and } \hat{\delta}_{LC} = \frac{n+a}{c_1} \ln\left(1 + \frac{c_1}{b + \sum_{i=1}^n x_i}\right).$$

Using (22) in (21), we get posterior risk under LINEX loss function for different priors.

3.3 Generalized-Entropy Loss Function

Calabria and Pulcini (1996) defined generalized-entropy loss function (GELF) which is given by

$$L_G(\delta, \theta) = b_2 \left[\left(\frac{\delta}{\theta} \right)^{c_2} - c_2 \ln \left(\frac{\delta}{\theta} \right) - 1 \right] ; c_2 \neq 0, b_2 > 0. \quad (23)$$

The posterior risk under GELF is given by

$$R_G(\delta, \theta) = \int_{\theta} L_G(\delta, \theta) h(\theta|\underline{x}) d\theta \quad (24)$$

The Bayes estimator under GELF is given by

$$\hat{\delta}_G = [E(\theta^{-c_2}|\underline{x})]^{-1/c_2} \quad (25)$$

provided $E(\theta^{-c_2}|\underline{x})$ exists and finite. When $c_2 > 0$, a positive error ($\hat{\delta}_2 > 0$) causes more serious consequences than a negative error.

In particular for $c_2 = 1$, we get entropy loss function

$$L_G(\delta, \theta) = b_2 \left[\left(\frac{\delta}{\theta} \right) - \ln \left(\frac{\delta}{\theta} \right) - 1 \right]$$

The Bayes estimator under the entropy loss function is the posterior harmonic mean.

The Bayes estimators of the parameter θ of $\exp(\theta)$ -distribution under GELF $L_G(\delta, \theta)$ for different priors, viz uniform, Jefferys & conjugate priors are given by

$$\hat{\delta}_{GC} = \left[\frac{\Gamma(n+a-c_2)}{\Gamma(n+a)} (b + \sum_{i=1}^n x_i)^{c_2} \right]^{\frac{-1}{c_2}},$$

$$\hat{\delta}_{GU} = \left[\frac{\Gamma(n+1-c_2)}{\Gamma(n+1)} (\sum_{i=1}^n x_i)^{c_2} \right]^{\frac{-1}{c_2}} \text{ and } \hat{\delta}_{GJ} = \left[\frac{\Gamma(n-c_2)}{\Gamma(n)} (\sum_{i=1}^n x_i)^{c_2} \right]^{\frac{-1}{c_2}}$$

Now, using (25) in (24), we get the posterior risk of the Bayes estimator for θ under GELF are given below

$$R_G(\hat{\delta}_G, \theta) = b_2 \left[\hat{\delta}_G^{c_2} E(\theta^{-c_2}|\underline{x}) - c_2 \ln \hat{\delta}_G + c_2 E(\ln \theta|\underline{x}) - 1 \right] \quad (26)$$

3.4 Quadratic Type Square Error Loss Function (QTSELF)

Using (3), the bayes estimator of θ of $\exp(\theta)$ -distribution under the proposed loss function (1), for different considered priors viz., uniform, Jefferys, conjugate priors are the corresponding posterior risks will be given in terms of A, B & C, where

$$A = c \left(1 + \frac{c}{b + \sum_{i=1}^n x_i} \right)^{-(n+a)}$$

$$B = 2 \left(1 - \frac{c(n+a)}{b+c + \sum_{i=1}^n x_i} \right) \left(1 + \frac{c}{b + \sum_{i=1}^n x_i} \right)^{-(n+a)}$$

$$C = \left(\frac{c(n+a+1)}{b+c + \sum_{i=1}^n x_i} - 2 \right) \left(\frac{(n+a)}{b+c + \sum_{i=1}^n x_i} \right) \left(1 + \frac{c}{b + \sum_{i=1}^n x_i} \right)^{-(n+a)}$$

It is to be noted here that the above values of A, B, C are corresponding to conjugate prior only. The value of A, B, C corresponding to uniform prior can be obtained by taking $a=1$, $b=0$ and those for Jefferys prior can be obtained by putting $a=0$, $b=0$.

4 Real Data Application

In this section, a real data set is considered and analyzed to check the suitability of the proposed asymmetric loss function over the other existing loss functions such as SELF, GELF and LINEX loss functions. The considered real data set is the remission time (in months) of 128 bladder cancer patients data extracted from Lee, E. T., & Wang, J. (2003).

Some of the statistical properties such as mean, mode, variance, skewness and kurtosis of the posterior distribution of θ given the real data set $\underline{X}$ corresponding to UP, JP, and CP has been calculated and shown in Table 1.

The Bayes estimators of θ and the corresponding risk (in parenthesis) under SELF, LINEX, GELF and proposed loss function with respect to uniform prior (UP), Jeffreys prior (JP) and Conjugate prior (CP) are shown in the Table 2.

Table 1: Properties of Posterior distributions of θ given the real data set $\underline{X}$ corresponding to UP, JP, and CP

Prior	Mean	Mode	Variance	Skewness	Kurtosis
UP	0.1067174	0.1058901	8.83E-05	0.1760902	0.04651163
JP	0.1058901	0.1050629	8.76E-05	0.1767767	0.046875
CP	0.1097994	0.1089738	9.06E-05	0.173422	0.04511278

Table 2: Bayes estimates and Posterior Risk (in parenthesis) under different Loss Functions and priors for the considered real data

Prior	SELF	LINEX		GELF		QTSELF	
		C1=2	C1=2	C2=2	C2=2	C=0.5	C=-0.5
UP	0.1067174 (0.0000883)	0.1066292 (0.0001764)	0.1068058 (0.0001768)	0.1054757 (0.0164467)	0.1071302 (0.0146828)	0.1066512 (0.0000882)	0.1067837 (0.0000884)
JP	0.1058901 (0.0000876)	0.1058026 (0.0001750)	0.1059778 (0.0001754)	0.1046484 (0.0169891)	0.1063030 (0.0143846)	0.1058245 (0.0000875)	0.1059559 (0.0000877)
CP V=0.1	0.1367188 (0.0001113)	0.1366076 (0.0002223)	0.1368301 (0.0002228)	0.1354974 (0.0116683)	0.1371250 (0.0122128)	0.1366353 (0.0001112)	0.1368022 (0.0001114)
CP V=150	0.1059110 (0.0000876)	0.1058235 (0.0001750)	0.1059987 (0.0001754)	0.1046693 (0.0148241)	0.1063239 (0.0165430)	0.1058454 (0.0000875)	0.1059768 (0.0000877)

5 Simulation Studies:

In this section, we have compared the considered Bayes estimators $\hat{\delta}_S, \hat{\delta}_L, \hat{\delta}_G, \hat{\delta}_{QT}$ of the parameter θ of $\exp(\theta)$ -distribution in terms of the posterior risks under different loss functions for different priors viz informative as well as non-informative priors. For the purpose we have generated 5000 different samples. It is clear that the posterior

risks are the function of n , the no. of observations put on the test, hyper parameters a and b of the informative prior distribution taken and the loss parameters of the loss functions considered. The different values of n taken are 10, 15, 20, 25, 30, 40, 50, 60, 70, 80; the arbitrarily chosen values of θ are 1, 1.5, 2. The different values of loss parameters of different asymmetric loss functions are arbitrarily taken as $c_1 = \pm 3$, $b_1 = 1$, $c_2 = \pm 3$, $b_2 = 1$ and $c = \pm 2$. Tables 3-6 shows the posterior risk of the different estimators of θ for the variation of values of n and θ under Uniform, Jeffreys, Conjugate priors (with high and low confidence level) respectively. Table 5 shows posterior risks under CP for different values of n where the prior mean M is kept 1.5 and prior variance is low ($V=0.1$), whereas Table 6 shows posterior risks when prior variance is high ($V=150$). While prior mean M is again kept as 1.5. To achieve this goal, the values of hyper-parameters were taken as $a=0.015$, $b=0.01$ (so that $M=1.5$, $V=150$). Table 7 shows the posterior risks of different estimators of θ for fixed value of $\theta=1$, whereas guessed value of θ i.e, prior mean M is taken as 0, 1 and 2 and prior variance varies as 0.05, 0.1, 0.5, 1, 2, 5, 10, 80, 250.

From Tables 3 and 4, it is seen that the posterior risks of the Bayes estimator under proposed loss QTSELF is smallest as compared to those of other Bayes estimators when sample size is small. On the other hand, for large sample sizes, among considered asymmetric loss function, our proposed Bayes estimator has the smallest posterior risk but when we compared Bayes estimators under all considered loss functions, Bayes estimator under SELF outperform. Bayes estimators under considered asymmetric loss functions in the sense of having smallest posterior risk when sample size is large.

Further, Tables 5 and 6 shows the posterior risks of the different Bayes estimators for CP of θ under different considered loss functions when prior variance is taken as $V=0.05$ & $V=150$ respectively. The performance of the Bayes estimators are found similar to Tables 3 & 4. It is also noted that posterior risks for $V=150$ is higher than the corresponding posterior risks for $V=0.05$.

Finally, Table 7 shows the posterior risk of different Bayes estimators under different loss functions with respect to conjugate prior (CP) when mean & variance of the CP (i.e, guessed value M & confidence level V of the guessed value of θ) are varying. It is calculated for $M=0.5, 1, 2$; $V=0.05, 0.1, 0.5, 1, 2, 5, 10, 80, 250$ and the fixed value of $\theta=1$. It is observed from this table that the posterior risk of the Bayes estimator of θ under the proposed loss function is smaller as compared to those under other considered loss function for all the considered values of loss parameters of the different asymmetric loss function considered. Whatever may be the situation of seriousness, i.e, either over estimation is more serious than under estimation or vice-versa. It is remarkable point to note here that the posterior risks of the different Bayes estimators increases with increase in the prior variance V , for each considered value of the prior mean M .

6 Conclusion

In the present piece of work, we have proposed a new asymmetric loss function QTSELF as a generalization of SELF and the proposed loss function can be used in the situation when over estimation is more serious than under estimation and vice-versa unlike weighted squared error loss function (WSELF). To exhibits its application for long run use, we have performed simulation study and the study that the Bayes estimator of the parameter θ of $\exp(\theta)$ -distribution has smallest posterior risk, as compared to other asymmetric loss function LINEX & GELF as well as SELF in some situations. Thus, we recommended its application in Bayesian paradigm in order to have estimation closed to the reality.

Acknowledgment

The authors are very thankful to the Editor and the anonymous referees for their valuable, fruitful and constructive suggestions that led to improvement of the manuscript.

Table 3: Posterior Risk for Different Loss Functions under Uniform Prior

n	θ	SELF	LINEX		GELF		QTSELF	
			$c_1 = 3$	$c_1 = -3$	$c_2 = 3$	$c_2 = -3$	$c = 2$	$c = -2$
10	1	0.12735	0.47332	0.73754	0.47605	0.39245	0.06563	0.11406
	1.5	0.12973	0.48141	0.75337	0.47603	0.39247	0.07298	0.11508
	2	0.13151	0.48747	0.76529	0.47598	0.39252	0.07154	0.11580
15	1	0.10232	0.39703	0.55303	0.31102	0.27310	0.06248	0.09364
	1.5	0.11736	0.45124	0.64313	0.31098	0.27314	0.06595	0.10236
	2	0.13095	0.49988	0.72565	0.31103	0.27309	0.06656	0.10933
20	1	0.07446	0.29930	0.38241	0.23098	0.20953	0.05565	0.07181
	1.5	0.08804	0.35084	0.45740	0.23115	0.20936	0.06400	0.08216
	2	0.09083	0.36164	0.47245	0.23095	0.20957	0.05943	0.08463
25	1	0.05631	0.23176	0.28022	0.18379	0.16992	0.04541	0.05576
	1.5	0.06359	0.26049	0.31837	0.18381	0.16990	0.06117	0.06212
	2	0.07234	0.29471	0.36476	0.18381	0.16990	0.05589	0.06935
30	1	0.04546	0.18998	0.22198	0.15266	0.14288	0.03945	0.04544
	1.5	0.05349	0.22227	0.26295	0.15261	0.14292	0.04407	0.05268
	2	0.05674	0.23534	0.27961	0.15265	0.14289	0.04432	0.05559
40	1	0.03294	0.14024	0.15731	0.11396	0.10846	0.03353	0.03312
	1.5	0.04009	0.16979	0.19261	0.11397	0.10844	0.04200	0.03983
	2	0.04153	0.17574	0.19971	0.11393	0.10849	0.03429	0.04119
50	1	0.02540	0.10940	0.11971	0.09089	0.08743	0.02679	0.02560
	1.5	0.03061	0.13134	0.14494	0.09091	0.08741	0.03346	0.03061
	2	0.03275	0.14030	0.15531	0.09089	0.08742	0.02877	0.03264
60	1	0.02093	0.09083	0.09788	0.07573	0.07309	0.02395	0.02110
	1.5	0.02690	0.11616	0.12640	0.07557	0.07325	0.04305	0.02685
	2	0.02703	0.11975	0.12703	0.07565	0.07317	0.02433	0.02700
70	1	0.01771	0.07725	0.08234	0.06467	0.06303	0.02280	0.01784
	1.5	0.02194	0.09537	0.10234	0.06480	0.06289	0.02168	0.02197
	2	0.01812	0.11751	0.10242	0.06477	0.06292	0.01952	0.01911
80	1	0.01527	0.06687	0.07067	0.05661	0.05521	0.01997	0.01538
	1.5	0.01925	0.08405	0.08941	0.05659	0.05524	0.01982	0.01928
	2	0.01163	0.10840	0.09707	0.05655	0.05528	0.01866	0.01435

Table 4: Posterior Risk for Different Loss Functions under Jeffreys Prior

n	θ	SELF	LINEX		GELF		QTSELF	
			$c_1 = 3$	$c_1 = -3$	$c_2 = 3$	$c_2 = -3$	$c = 2$	$c = -2$
10	1	0.13332	0.48957	0.78878	0.53257	0.43024	0.06572	0.11826
	1.5	0.13419	0.49247	0.79477	0.53287	0.42994	0.06594	0.11856
	2	0.13457	0.49374	0.79739	0.53265	0.43016	0.07021	0.11869
15	1	0.10535	0.40634	0.57457	0.33413	0.29081	0.06302	0.09660
	1.5	0.11777	0.45071	0.65015	0.33412	0.29082	0.07349	0.10356
	2	0.12915	0.49112	0.71999	0.33414	0.29079	0.06615	0.10942
20	1	0.07737	0.30952	0.39995	0.24357	0.21970	0.05587	0.07438
	1.5	0.09799	0.38671	0.51630	0.24362	0.21966	0.07549	0.08874
	2	0.09575	0.37892	0.50229	0.24347	0.21981	0.06130	0.08833
25	1	0.05745	0.23586	0.28677	0.19169	0.17653	0.04674	0.05696
	1.5	0.06637	0.27084	0.33393	0.19165	0.17657	0.05357	0.06460
	2	0.07265	0.29533	0.36731	0.19164	0.17658	0.05222	0.06983
30	1	0.04738	0.19743	0.23213	0.15793	0.14765	0.04510	0.04725
	1.5	0.05672	0.23484	0.28015	0.15805	0.14754	0.05259	0.05557
	2	0.05910	0.24443	0.29233	0.15792	0.14767	0.04468	0.05775
40	1	0.03318	0.14115	0.15860	0.11688	0.11117	0.03170	0.03340
	1.5	0.04080	0.17257	0.19631	0.11696	0.11109	0.04218	0.04053
	2	0.04263	0.18012	0.20537	0.11694	0.11111	0.03498	0.04225
50	1	0.02564	0.11037	0.12094	0.09274	0.08919	0.03074	0.02586
	1.5	0.03174	0.13597	0.15048	0.09277	0.08915	0.03201	0.03169
	2	0.03416	0.14612	0.16227	0.09282	0.08911	0.03397	0.03399
60	1	0.02096	0.09093	0.09805	0.07697	0.07435	0.02605	0.02114
	1.5	0.02592	0.11200	0.12176	0.07693	0.07440	0.02519	0.02595
	2	0.02068	0.12115	0.14867	0.07695	0.07437	0.02401	0.03099
70	1	0.01775	0.07738	0.08251	0.06564	0.06390	0.02013	0.01789
	1.5	0.02221	0.09650	0.10365	0.06567	0.06387	0.02232	0.02224
	2	0.01601	0.10969	0.11548	0.06568	0.06386	0.02131	0.02780
80	1	0.01534	0.06715	0.07100	0.05735	0.05588	0.01731	0.01545
	1.5	0.01926	0.08405	0.08945	0.05732	0.05592	0.02211	0.01929
	2	0.00887	0.07075	0.08352	0.05728	0.05595	0.01955	0.02569

Table 5: Posterior Risk for Different Loss Functions under Conjugate Prior with high confidence level $V=0.05$ & $M=1.5$

n	θ	SELF	LINEX		GELF		QTSELF	
			c1=3	c1=-3	c2=3	c2=-3	c=2	c=-2
10	1	0.04825	0.20160	0.23564	0.14526	0.13638	0.04211	0.04786
	1.5	0.05414	0.22537	0.26566	0.14531	0.13633	0.04581	0.05314
	2	0.05505	0.22903	0.27026	0.14519	0.13645	0.04228	0.05396
15	1	0.03984	0.16832	0.19203	0.12513	0.11837	0.03850	0.03980
	1.5	0.04712	0.19806	0.22845	0.12499	0.11851	0.05270	0.04645
	2	0.04654	0.19572	0.22547	0.12499	0.11851	0.03716	0.04597
20	1	0.03540	0.15057	0.16925	0.10975	0.10470	0.04428	0.03537
	1.5	0.03851	0.16348	0.18452	0.10981	0.10465	0.03540	0.03832
	2	0.04041	0.17135	0.19392	0.10982	0.10463	0.03308	0.04009
25	1	0.02962	0.12694	0.14040	0.09782	0.09380	0.03658	0.02975
	1.5	0.03430	0.14652	0.16323	0.09783	0.09378	0.03035	0.03420
	2	0.03570	0.15232	0.17003	0.09784	0.09377	0.03002	0.03551
30	1	0.02628	0.11319	0.12389	0.08825	0.08491	0.04619	0.02642
	1.5	0.03051	0.13101	0.14433	0.08816	0.08501	0.03023	0.03048
	2	0.03203	0.13737	0.15166	0.08825	0.08492	0.02734	0.03191
40	1	0.02174	0.09431	0.10168	0.07372	0.07150	0.02602	0.02186
	1.5	0.02545	0.11009	0.11938	0.07383	0.07139	0.02383	0.02546
	2	0.02649	0.11449	0.12434	0.07375	0.07147	0.02376	0.02645
50	1	0.01817	0.07927	0.08449	0.06335	0.06169	0.02053	0.01829
	1.5	0.02212	0.09620	0.10318	0.06336	0.06167	0.02226	0.02213
	2	0.02268	0.09857	0.10580	0.06339	0.06165	0.02090	0.02267
60	1	0.01580	0.06918	0.07315	0.05557	0.05421	0.02267	0.01589
	1.5	0.01909	0.08337	0.08862	0.05554	0.05423	0.01883	0.01912
	2	0.01971	0.08605	0.09155	0.05551	0.05427	0.01832	0.01972
70	1	0.01372	0.06028	0.06331	0.04948	0.04836	0.04368	0.01381
	1.5	0.01690	0.07407	0.07820	0.04945	0.04839	0.01661	0.01693
	2	0.01774	0.07633	0.07668	0.04949	0.04835	0.01740	0.01762
80	1	0.01232	0.05424	0.05668	0.04454	0.04371	0.01629	0.01239
	1.5	0.01520	0.06679	0.07013	0.04454	0.04370	0.01604	0.01523
	2	0.01665	0.06904	0.07220	0.04461	0.04374	0.01676	0.01634

Table 6: Posterior Risk for Different Loss Functions under Conjugate Prior with low confidence level $V=150$ & $M=1.5$.

n	θ	SELF	LINEX		GELF		QTSELF	
			$c_1 = 3$	$c_1 = -3$	$c_2 = 3$	$c_2 = -3$	$c = 2$	$c = -2$
10	1	0.13330	0.48959	0.78852	0.53186	0.42938	0.06594	0.11823
	1.5	0.13388	0.49152	0.79247	0.53184	0.42940	0.06894	0.11843
	2	0.13453	0.49367	0.79692	0.53165	0.42959	0.07601	0.11865
15	1	0.10446	0.40317	0.56910	0.33362	0.29066	0.06177	0.09606
	1.5	0.11787	0.45110	0.65063	0.33381	0.29047	0.06765	0.10363
	2	0.12893	0.49040	0.71850	0.33370	0.29058	0.07077	0.10934
20	1	0.07779	0.31108	0.40229	0.24341	0.21950	0.05630	0.07465
	1.5	0.09152	0.36282	0.47906	0.24342	0.21949	0.06400	0.08483
	2	0.09520	0.37691	0.49909	0.24336	0.21955	0.06256	0.08799
25	1	0.05676	0.23322	0.28311	0.19138	0.17661	0.05052	0.05640
	1.5	0.07092	0.28846	0.35844	0.19158	0.17642	0.06057	0.06815
	2	0.07228	0.29393	0.36532	0.19147	0.17652	0.05439	0.06955
30	1	0.04624	0.19291	0.22624	0.15792	0.14751	0.04052	0.04628
	1.5	0.05684	0.23532	0.28073	0.15789	0.14753	0.05222	0.05568
	2	0.05928	0.24513	0.29325	0.15789	0.14754	0.04518	0.05789
40	1	0.03343	0.14219	0.15986	0.11692	0.11105	0.03417	0.03363
	1.5	0.03933	0.16656	0.18900	0.11691	0.11106	0.03872	0.03919
	2	0.04286	0.18109	0.20655	0.11692	0.11105	0.03557	0.04246
50	1	0.02587	0.11134	0.12206	0.09278	0.08909	0.02601	0.02608
	1.5	0.03159	0.13537	0.14978	0.09281	0.08906	0.03102	0.03156
	2	0.03365	0.14400	0.15979	0.09281	0.08905	0.02943	0.03351
60	1	0.02101	0.09112	0.09827	0.07692	0.07437	0.02356	0.02118
	1.5	0.02610	0.11275	0.12260	0.07699	0.07430	0.02617	0.02611
	2	0.07324	0.30197	0.36452	0.07689	0.07439	0.02442	0.05243
70	1	0.01781	0.07766	0.08282	0.06568	0.06383	0.02230	0.01795
	1.5	0.02214	0.09623	0.10335	0.06570	0.06381	0.02186	0.02218
	2	0.07280	0.31026	0.36917	0.06575	0.06376	0.02031	0.05382
80	1	0.01530	0.06700	0.07083	0.05735	0.05586	0.01715	0.01542
	1.5	0.01936	0.08448	0.08992	0.05734	0.05587	0.01930	0.01939
	2	0.07002	0.29283	0.34196	0.05732	0.05589	0.01858	0.05240

Table 7: Posterior Risk for Different Loss Functions under Conjugate Prior with varying confidence level and guess values for $\theta=1$ & $n=20$.

V	M	SELF	LINEX		GELF		QTSELF	
			$c_1 = 3$	$c_1 = -3$	$c_2 = 3$	$c_2 = -3$	$c = 2$	$c = -2$
0.05	0.5	0.06238	0.25273	0.31679	0.24353	0.21974	0.04445	0.06278
	1	0.03549	0.15016	0.17072	0.13443	0.12676	0.03056	0.03590
	2	0.01689	0.07405	0.07812	0.04812	0.04713	0.01622	0.01691
0.1	0.5	0.08080	0.32024	0.42317	0.28176	0.25028	0.05355	0.07845
	1	0.05481	0.22555	0.27284	0.19165	0.17657	0.04525	0.05471
	2	0.02944	0.12662	0.13895	0.08413	0.08109	0.02913	0.02940
0.5	0.5	0.10073	0.39054	0.54494	0.32215	0.28168	0.06135	0.09330
	1	0.09361	0.36673	0.49876	0.29090	0.25743	0.06086	0.08729
	2	0.06803	0.27612	0.34476	0.20957	0.19156	0.05406	0.06635
1	0.5	0.10245	0.39631	0.55616	0.32805	0.28616	0.06226	0.09456
	1	0.09895	0.38483	0.53288	0.31103	0.27309	0.06189	0.09162
	2	0.08585	0.34044	0.44936	0.25747	0.23105	0.05877	0.08103
2	0.5	0.10370	0.40063	0.56406	0.33108	0.28845	0.06230	0.09551
	1	0.10224	0.39596	0.55417	0.32221	0.28162	0.06202	0.09410
	2	0.09713	0.37956	0.51947	0.29102	0.25731	0.06160	0.08945
5	0.5	0.10395	0.40139	0.56586	0.33299	0.28978	0.06273	0.09573
	1	0.10373	0.40086	0.56399	0.32927	0.28705	0.06237	0.09544
	2	0.10146	0.39363	0.54844	0.31528	0.27657	0.06302	0.09338
10	0.5	0.10402	0.40377	0.57006	0.33354	0.29031	0.06278	0.09614
	1	0.10429	0.40271	0.56771	0.33167	0.28893	0.06241	0.09589
	2	0.10324	0.39943	0.56043	0.32447	0.28347	0.06324	0.09488
80	0.5	0.10441	0.40396	0.56884	0.33406	0.29074	0.06421	0.09613
	1	0.10481	0.40443	0.57124	0.33380	0.29060	0.06294	0.09627
	2	0.10523	0.40598	0.57369	0.33289	0.28987	0.06325	0.09644
250	0.5	0.10488	0.40466	0.57171	0.33408	0.29082	0.06462	0.09639
	1	0.10483	0.40449	0.57140	0.33409	0.29067	0.06388	0.09629
	2	0.10502	0.40517	0.57253	0.33382	0.29042	0.06430	0.09636

References

- [1] Asgharzadeh, A., & Farsipour, N. S. (2008). Estimation of the exponential mean time to failure under a weighted balanced loss function. *Statistical Papers*, 49(1), 121-131.
- [2] Dey, S., & Maiti, S. S. (2010). Bayesian estimation of the parameter of Maxwell distribution under different loss functions. *Journal of Statistical Theory and Practice*, 4(2), 279-287.
- [3] Kumar, D., Singh, U., & Singh, S. K. (2015). A new distribution using sine function-its application to bladder cancer patients data. *Journal of Statistics Applications & Probability*, 4(3), 417.
- [4] Kazmi, A., Mohsin, S., Aslam, M., & Ali, S. (2012). Preference of Prior for the class of Lifetime distributions under different loss functions. *Pakistan Journal of Statistics*, 28(4).
- [5] Kundu, D. (2008). Bayesian inference and life testing plan for the Weibull distribution in presence of progressive censoring. *Technometrics*, 50(2), 144-154.
- [6] Lee, E. T., & Wang, J. (2003). *Statistical methods for survival data analysis* (Vol. 476). John Wiley & Sons.
- [7] Ren, C., Sun, D., & Dey, D. K. (2004). Comparison of Bayesian and frequentist estimation and prediction for a normal population. *Sankhya: The Indian Journal of Statistics*, 678-706.
- [8] Singh, S. K., Singh, U., & Kumar, D. (2013). Bayesian estimation of parameters of inverse Weibull distribution. *Journal of Applied statistics*, 40(7), 1597-1607.
- [9] Singh, H. P., & Saxena, S. (2005). Bayesian and shrinkage estimation of process capability index C_p . *Communications in Statistics-Theory and Methods*, 34(1), 205-228.
- [10] Singh, S. K., Singh, U., & Kumar, D. (2013). Bayes estimators of the reliability function and parameter of inverted exponential distribution using informative and non-informative priors. *Journal of Statistical computation and simulation*, 83(12), 2258-2269.
- [11] Singh, S. K., Singh, U., & Sharma, V. K. (2014). Bayesian estimation and prediction for the generalized Lindley distribution under asymmetric loss function. *Haceteppe Journal of Mathematics and Statistics*, 43(4), 661-678.
- [12] Sinha, S. K. (1998). *Bayesian estimation*. New Age International (P) Limited, New Delhi.
- [13] Srivastava, U. (2012). Bayesian estimation of shift point in Poisson model under asymmetric loss functions. *Pakistan Journal of Statistics and Operation Research*, 8(1), 31-42.
- [14] Varian, H. R. (1975). A third remark on the number of equilibria of an economy. *Econometrica* (pre-1986), 43(5-6), 985.
- [15] Zellner, A. (1986). Bayesian estimation and prediction using asymmetric loss functions. *Journal of the American Statistical Association*, 81(394), 446-451.