

Minimax Estimation of the Scale Parameter in a Class of Life-Time Distributions for Different Loss Functions

Chandan Kumer Podder

Department of Statistics, University of Chittagong
Chittagong – 4331, Bangladesh

*Correspondence should be addressed to Chandan Kumer Podder
(Email: podder_ck@yahoo.com)

[Received June 02 2020; Accepted November 20, 2020]

Abstract

In this paper minimax estimation of the scale parameter in a class of life-time distributions for Squared Log Error, Modified Linear Exponential (MLINEX) and Quadratic loss functions have been studied. A study on risks of the estimators are also given.

Keywords: Minimax estimators, Risk functions, MLINEX loss, Squared log error loss and Quadratic loss functions.

AMS Classification: 62N02.

1. Introduction

The class of life-time distributions introduced by Prakash and Singh [7] is an important life-time distribution in survival analysis. Suppose a random variable X follows the distribution presented by a class of probability density function with the parameter θ and two known positive quantities b and c is given as

$$f(X = x; \theta) = \frac{c}{\Gamma b} \frac{1}{\theta^b} x^{bc-1} e^{-\frac{1}{\theta} x^c}; \quad x \geq 0, \theta > 0, b > 0, c > 0, \quad (1.1)$$

It can be seen that for different values of b and c the model (1.1) reduces to negative exponential distribution, two-parameter gamma distribution, Erlang

distribution, two-parameter Weibull distribution, Rayleigh distribution and Maxwell distribution.

The various properties and estimation of different life-time distributions such as exponential distribution, Weibull distribution, two-parameter gamma distribution and Maxwell's velocity distribution have been studied by Abu-Taleb et. al. [1], Ahmed et. al. [2], Sonand Oh [8] and Tyagi and Bhattacharya [9] etc. Prakash and Singh [7] discussed the Bayesian shrinkage estimation in a class of life testing distribution. Podder [6] studied the risks of the Bayes' estimators for the parameter of the distribution in (1.1) under squared-error and MLINEX loss functions.

The purpose of this paper is to find the minimax estimators of the scale parameter in a class of life-time distributions for squared log error, MLINEX and quadratic loss functions. In addition, a study on their risks have been also done.

2. Preliminary Theory

Suppose X is a random variable whose distribution depends on k parameters $\theta_1, \theta_2, \dots, \theta_k$ and let Ω denotes the parameter space of possible values of θ , the k - dimensional vector $(\theta_1, \theta_2, \dots, \theta_k)$. Now consider the general problem of estimating the unknown parameter θ , from the results of a random sample of n observations, by the methods of Bayes' and minimax estimation.

Denoting the sample results $x_1, x_2, \dots, x_n$ by $\underline{x}$, let $\hat{\theta}$ be an estimator of θ and also let $L(\hat{\theta}, \theta)$ be a loss function, the loss incurred by taking the value of θ to be $\hat{\theta}$. The risk function $R(\hat{\theta}, \theta)$ is the expected value of the loss function with respect to the sample observations.

If $l(\theta | x)$ is the likelihood function of θ given the sample $\underline{x}$ and $\pi(\theta)$ be the prior density of θ , then combining $l(\theta | x)$ and $\pi(\theta)$, it produces the posterior distribution $P(\theta | x)$ though the Bayes' theorem as

$$P(\theta | x) = \frac{l(\theta | x)g(\theta)}{p(x)}, \quad (2.1)$$

where $p(x) = \int_{\Omega} l(\theta | x) g(\theta) d\theta$.

Hence the Bayes' estimator $\hat{\theta}$ of θ will be a solution of the equation

$$\int_{\Omega} \frac{\partial L}{\partial \hat{\theta}} P(\theta | x) d\theta = 0, \quad (2.2)$$

where L stands for loss function and assume that the sufficient regularities conditions prevail to permit differentiation under the integral sign.

Here, we consider the following types of the loss functions as

$$\text{i) } L_1(\hat{\theta}, \theta) = (\ln \hat{\theta} - \ln \theta)^2 = \left(\ln \left(\frac{\hat{\theta}}{\theta} \right) \right)^2; \quad (2.3)$$

$$\text{ii) } L_2(\hat{\theta}, \theta) = \varpi \left[\left(\frac{\hat{\theta}}{\theta} \right)^{\gamma} - \gamma \ln \left(\frac{\hat{\theta}}{\theta} \right) - 1 \right]; \gamma \neq 0, \varpi > 0; \quad (2.4)$$

$$\text{iii) } L_3(\hat{\theta}, \theta) = \left(\frac{\hat{\theta} - \theta}{\theta} \right)^2. \quad (2.5)$$

The loss function L_1 is called squared log error proposed by Brown [3], is balanced and $\lim_{\hat{\theta} \rightarrow 0 \text{ or } \infty} L(\hat{\theta}, \theta) = \infty$. This loss is not always convex, it is convex for $\frac{\hat{\theta}}{\theta} \leq e$ and concave otherwise, but its risk function has a unique minimum with respect to $\hat{\theta}$.

Again the loss function L_2 is a modified linear-exponential (MLINEX) which is an asymmetric one. If $\frac{\hat{\theta}}{\theta} = 1$, then $L(\hat{\theta}, \theta) = 0$, writing $R = \frac{\hat{\theta}}{\theta}$, the relative error $L(R)$ is minimum at $R = 1$. If we write $D = \ln R = \ln \left(\frac{\hat{\theta}}{\theta} \right)$, then $L(R)$ can be expressed as the same form of LINEX (Linear-exponential) loss function, introduced by Varian [10].

The loss function L_3 is called quadratic loss function which is asymmetric one. The derivation depends primarily on a theorem which is due to Lehmann [5] and can be stated as follows.

Theorem 2.1: Let $\tau = \{F_\theta; \theta \in \Theta\}$ be a family of distribution functions and D be a class of estimators of θ . Suppose that $\delta^* \in D$ is a Bayes' estimator concerning to a prior distribution $\pi(\theta)$ on the parameter space Θ . If the risk function $R(\delta^*, \theta) = \text{constant on } \Theta$, then δ^* is a minimax estimator for θ .

3. Main Results

Theorem 3.1: Let $X_1, X_2, \dots, X_n$ be a random sample of size n from the distribution in (1.1). If θ has Jeffrey's non-informative prior density $\pi(\theta) \propto \frac{1}{\theta}$; $\theta > 0$, then

(a) $\hat{\theta}_1^* = T e^{-\psi(nb)}$ is the minimax estimator of parameter θ for squared log error loss function of the type $L_1(\hat{\theta}, \theta) = (\ln \hat{\theta} - \ln \theta)^2 = \left(\ln \frac{\hat{\theta}}{\theta} \right)^2$, where

$$\psi(nb) = \frac{\Gamma'(nb)}{\Gamma(nb)}.$$

(b) $\hat{\theta}_2^* = \left(\frac{\Gamma(nb)}{\Gamma(nb + \gamma)} \right)^{\frac{1}{\gamma}} T$ is the minimax estimator of parameter θ for MLINEX

$$\text{loss function of the type } L_2(\hat{\theta}, \theta) = \varpi \left[\left(\frac{\hat{\theta}}{\theta} \right)^\gamma - \gamma \ln \left(\frac{\hat{\theta}}{\theta} \right) - 1 \right]; \gamma \neq 0, \varpi > 0.$$

(c) $\hat{\theta}_3^* = \frac{T}{nb+1}$ is the minimax estimator of parameter θ for quadratic loss

$$\text{function of the type } L_3(\hat{\theta}, \theta) = \left(\frac{\hat{\theta} - \theta}{\theta} \right)^2, \text{ where } T = \sum_{i=1}^n X_i^c.$$

Proof. Part (a): It is enough to show that the estimator $\hat{\theta}_1^* = Te^{-\psi^{(nb)}}$ is the Bayes' estimator for parameter θ , in a class of life-time distributions in (1.1), with constant risk under the prior density $\pi(\theta) \propto \frac{1}{\theta}; \theta > 0$.

Let us consider the case of estimating the single parameter θ in a class of life-time distributions in (1.1). The likelihood function is given by

$$\begin{aligned} l(\theta | \underline{x}) &= \left\{ \left(\frac{c}{\Gamma b} \right)^n \prod_{i=1}^n x_i^{bc-1} \right\} \frac{1}{\theta^{nb}} e^{-\frac{1}{\theta} \sum_{i=1}^n x_i^c} \\ &= \kappa \frac{1}{\theta^{nb}} e^{-\frac{1}{\theta} T}, \end{aligned} \quad (3.1)$$

$$\text{where } \kappa = \left(\frac{c}{\Gamma b} \right)^n \prod_{i=1}^n x_i^{bc-1} \text{ and } T = \sum_{i=1}^n x_i^c.$$

The maximum likelihood estimator of θ is $\frac{T}{nb}$, where T is defined above. T is also a complete sufficient statistic for θ . It is to be noted that the part of the likelihood function which is relevant to Bayesian inference on the unknown parameter θ is $\frac{1}{\theta^{nb}} e^{-\frac{1}{\theta} T}$.

Since the parametric range in (1.1) is 0 to ∞ , therefore according to the Jeffrey's rule of thumb, the Jeffrey's prior becomes

$$\pi(\theta) \propto \frac{1}{\theta}; \quad \theta \geq 0, \quad (3.2)$$

By combining equations (3.1) and (3.2), the posterior distribution of θ given $\underline{x}$ becomes as

$$\pi(\theta | \underline{x}) = \frac{T^{nb}}{\Gamma(nb)} \frac{1}{\theta^{nb+1}} e^{-\frac{1}{\theta} T}; \quad \theta \geq 0, T > 0, \quad (3.3)$$

The mean and variance of the posterior distribution in (3.3) are $\frac{T}{(nb-1)}$ and $\frac{T^2}{(nb-1)^2(nb-2)}$ respectively. Also the distribution in (3.3) gives

$$\begin{aligned} E(\theta^{-\gamma} | x) &= \int_0^{\infty} \theta^{-\gamma} \pi(\theta | x) d\theta \\ &= \frac{\Gamma(nb + \gamma)}{\Gamma(nb)} \frac{1}{T^{\gamma}}. \end{aligned} \quad (3.4)$$

$$\begin{aligned} \text{And } E(\ln \theta | x) &= \int_0^{\infty} \ln \theta \pi(\theta | x) d\theta \\ &= \frac{T^{nb}}{\Gamma(nb)} \int_0^{\infty} \ln \theta \frac{1}{\theta^{nb+1}} e^{-\frac{T}{\theta}} d\theta \end{aligned}$$

Using a transformation, $y = \frac{1}{\theta} T$, then

$$E(\ln \theta | x) = \ln T - \psi(nb), \quad (3.5)$$

where $\psi(nb) = \frac{\Gamma'(nb)}{\Gamma(nb)}$ and $\Gamma'(nb) = \int_0^{\infty} \ln y e^{-y} y^{nb-1} dy$, the differentiation of $\Gamma(nb)$ with respect to n .

The Bayes' estimator of θ for squared log error loss function in (2.3) using (2.2) is

$$\hat{\theta}_1^* = \exp[E_{\theta}(\ln \theta | x)], \quad (3.6)$$

where $E_{\theta}(\cdot | x)$ stands for posterior expectation.

Using (3.5), the relation (3.6) gives

$$\hat{\theta}_1^* = T e^{-\psi(nb)}.$$

Therefore, it is enough to show that the risk of $\hat{\theta}_1^*$ is constant.

$$\text{Again, } R_1(\hat{\theta}_1^*, \theta) = E[\ln \hat{\theta}_1^* - \ln \theta]^2$$

$$\begin{aligned}
&= V(\ln \hat{\theta}_1^*) - [E(\ln \hat{\theta}_1^*) - \ln \theta]^2 \\
&= V(\ln T) + [E(\ln T) - \ln \theta - \psi(nb)]^2
\end{aligned} \tag{3.7}$$

Since x follows a class of life-time distributions in (1.1) with parameter θ , then $T = \sum_{i=1}^n x_i^c$ is distributed as a gamma distribution with parameters nb and $\frac{1}{\theta}$, i. e., $T \sim \text{Gamma}\left(nb, \frac{1}{\theta}\right)$.

The probability density function of T is

$$f(T; \theta) = \frac{1}{\Gamma(nb)} \frac{1}{\theta^{nb}} e^{-\frac{1}{\theta}T} T^{nb-1}; \quad T \geq 0, \theta > 0 \tag{3.8}$$

The mean and variance of the distribution in (3.8) are $nb\theta$ and $nb\theta^2$ respectively. The distribution in (3.8) also gives

$$\begin{aligned}
E(T^\gamma) &= \int_0^\infty T^\gamma f(T; \theta) dT \\
&= \frac{\Gamma(nb + \gamma)}{\Gamma(nb)} \theta^\gamma
\end{aligned} \tag{3.9}$$

$$\begin{aligned}
E(\ln T) &= \int_0^\infty \ln T f(T; \theta) dT \\
&= \frac{1}{\Gamma(nb)} \frac{1}{\theta^{nb}} \int_0^\infty \ln T e^{-\frac{1}{\theta}T} T^{nb-1} dT,
\end{aligned}$$

Further using a transformation $y = \frac{1}{\theta}T$, we have

$$\begin{aligned}
E(\ln T) &= \frac{1}{\Gamma(nb)} \int_0^\infty (\ln \theta + \ln y) e^{-y} y^{nb-1} dy \\
&= \ln \theta + \psi(nb),
\end{aligned} \tag{3.10}$$

Similarly,

$$\begin{aligned}
E(\ln T)^2 &= \int_0^{\infty} (\ln T)^2 f(T; \theta) dT \\
&= \frac{1}{\Gamma(nb)} \int_0^{\infty} (\ln \theta + \ln y)^2 e^{-y} y^{nb-1} dy \\
&= (\ln \theta)^2 + 2 \ln \theta \psi(nb) + \frac{1}{\Gamma(nb)} \int_0^{\infty} (\ln y)^2 e^{-y} y^{nb-1} dy.
\end{aligned} \tag{3.11}$$

Again

$$\begin{aligned}
\psi(nb) &= \frac{\Gamma'(nb)}{\Gamma(nb)} \\
\Rightarrow \psi(nb) \Gamma(nb) &= \Gamma'(nb) \\
\Rightarrow \psi(nb) \Gamma(nb) &= \int_0^{\infty} \ln y e^{-y} y^{nb-1} dy.
\end{aligned} \tag{3.12}$$

Differentiating the equation in (3.12) with respect to n , we have

$$\psi'(nb) \Gamma(nb) + \psi(nb) \Gamma'(nb) = \int_0^{\infty} (\ln y)^2 e^{-y} y^{nb-1} dy.$$

Therefore,

$$\begin{aligned}
\psi'(nb) \Gamma(nb) + \psi(nb) \Gamma'(nb) &= \int_0^{\infty} (\ln y)^2 e^{-y} y^{nb-1} dy \\
\Rightarrow \frac{1}{\Gamma(nb)} \int_0^{\infty} (\ln y)^2 e^{-y} y^{nb-1} dy &= \psi'(nb) + \psi^2(nb)
\end{aligned} \tag{3.13}$$

Hence, we have from (3.11), using (3.13)

$$E(\ln T)^2 = (\ln \theta)^2 + 2 \ln \theta \psi(nb) + \psi'(nb) + \psi^2(nb). \tag{3.14}$$

And

$$\begin{aligned}
V(\ln T) &= E(\ln T)^2 - [E(\ln T)]^2 \\
&= \psi'(nb)
\end{aligned} \tag{3.15}$$

Substituting the relation (3.10) and (3.15) in (3.7), we have

$$R_1(\hat{\theta}_1^*, \theta) = \psi'(nb); \quad (3.16)$$

which is a constant with respect to θ , as b and n are known and independent of θ .

So, from the Lehmann's theorem it follows that $\hat{\theta}_1^* = Te^{-\psi(nb)}$, is the minimax estimator of the scale parameter θ in a class of life-time distributions for squared log error loss function in (2.3).

Part (b): The Bayes' estimator for θ under the MLINEX loss function in (2.4) is

$$\hat{\theta}_2^* = \left[E_\theta \left(\theta^{-\gamma} \mid x \right) \right]^{-\frac{1}{\gamma}} \quad (3.17)$$

Using (3.4), we have from the relation (3.17)

$$\hat{\theta}_2^* = \left[\frac{\Gamma(nb)}{\Gamma(nb+\gamma)} \right]^{\frac{1}{\gamma}} T = KT, \quad (3.18)$$

$$\text{where } K = \left[\frac{\Gamma(nb)}{\Gamma(nb+\gamma)} \right]^{\frac{1}{\gamma}} \text{ and } T = \sum_{i=1}^n x_i^c.$$

Now the risk function of $\hat{\theta}_2^*$ under the MLINEX loss function is given by

$$\begin{aligned} R_2(\hat{\theta}_2^*, \theta) &= E[L(\hat{\theta}_2^*, \theta)] \\ &= \varpi \left[\frac{1}{\theta^\gamma} E(\hat{\theta}_2^*)^\gamma - \gamma E(\ln \hat{\theta}_2^*) + \gamma \ln \theta - 1 \right] \\ &= \varpi \left[\frac{K^\gamma}{\theta^\gamma} E(T^\gamma) - \gamma E(\ln T) - \gamma \ln K + \gamma \ln \theta - 1 \right] \end{aligned} \quad (3.19)$$

Using (3.9) and (3.10), the relation (3.19) gives

$$R_2(\hat{\theta}_2^*, \theta) = \varpi \left[\ln \frac{\Gamma(nb+\gamma)}{\Gamma(nb)} - \gamma \psi(nb) \right]; \quad (3.20)$$

which is a constant with respect to θ , as b and n are known and independent of θ .

So, from the Lehmann's theorem it follows that $\hat{\theta}_2^* = \left[\frac{\Gamma(nb)}{\Gamma(nb+\gamma)} \right]^{\frac{1}{\gamma}} T$, is the minimax estimator of the scale parameter θ in a class of life-time distributions under MLINEX loss function.

It has seen that when $\gamma = 1$, then $\hat{\theta}_2^* = \frac{T}{nb}$ is the maximum likelihood estimator of

θ and when $\gamma = -1$, then $\hat{\theta}^* = \frac{T}{(nb-1)}$ is the mean of the posterior distribution in (3.3)

Part (c): The Bayes' estimator for θ under quadratic loss function in (2.5) is

$$\hat{\theta}_3^* = \frac{E_{\theta}(\theta^{-1} | x)}{E_{\theta}(\theta^{-2} | x)}, \quad (3.21)$$

where $E_{\theta}(\cdot | x)$ stands for posterior expectation.

Using (3.4), the relation (3.21) gives

$$\hat{\theta}_3^* = \frac{1}{(nb+1)} T \quad (3.22)$$

Now the risk function of $\hat{\theta}_3^*$ under the quadratic loss function is given by

$$\begin{aligned} R_3(\hat{\theta}_3^*, \theta) &= E[L_3(\hat{\theta}_3^*, \theta)] \\ &= \frac{1}{\theta^2} E(\hat{\theta}_3^* - \theta)^2 \\ &= \frac{1}{\theta^2} \left[V(\hat{\theta}_3^*) + \{E(\hat{\theta}_3^*) - \theta\}^2 \right] \\ &= \frac{1}{\theta^2} \left[\frac{1}{(nb+1)^2} V(T) + \left\{ \frac{1}{(nb+1)} E(T) - \theta \right\}^2 \right] \\ &= \frac{1}{nb+1}; \end{aligned} \quad (3.23)$$

which is constant with respect to θ , as b and n are known and independent of θ .

So from the Lehmann's theorem, it follows that $\hat{\theta}_3^* = \frac{1}{(nb+1)}T$, is the minimax estimator of the scale parameter θ in a class of life-time distributions under the quadratic loss function.

The following tables give the risks of the estimators $R_1(\hat{\theta}_1^*)$, $R_2(\hat{\theta}_2^*)$ and $R_3(\hat{\theta}_3^*)$ under Squared log error, Modified linear-exponential (MLINEX) and Quadratic loss functions for different values of sample size n , b , γ and ϖ .

Table 1: The risks of Squared log error, MLINEX and Quadratic loss functions for different sample sizes when $b=1$, $\gamma \neq 0$ and $\varpi=1$

n	$R_1(\hat{\theta}_1^*)$	$R_2(\hat{\theta}_2^*)$						$R_3(\hat{\theta}_3^*)$
		$\gamma = -3$	$\gamma = -2$	$\gamma = -1$	$\gamma = 1$	$\gamma = 2$	$\gamma = 3$	
5	0.21875	1.35583	0.53768	0.12500	0.09814	0.37861	0.81322	0.16667
10	0.10494	0.53576	0.22889	0.05556	0.04980	0.19492	0.42705	0.09091
15	0.06888	0.33540	0.14554	0.03571	0.03328	0.13110	0.28954	0.06250
20	0.05125	0.24424	0.10670	0.02632	0.02498	0.09875	0.21903	0.04762
25	0.04080	0.19207	0.08423	0.02083	0.01999	0.07920	0.17615	0.03846

Table 2: The risks of Squared log error, MLINEX and Quadratic loss functions for different sample sizes when $b=\frac{3}{2}$, $\gamma \neq 0$ and $\varpi=2$.

n	$R_1(\hat{\theta}_1^*)$	$R_2(\hat{\theta}_2^*)$						$R_3(\hat{\theta}_3^*)$
		$\gamma = -3$	$\gamma = -2$	$\gamma = -1$	$\gamma = 1$	$\gamma = 2$	$\gamma = 3$	
5	0.14201	1.53110	0.64180	0.15385	0.13236	0.51504	1.12017	0.11765
10	0.06888	0.67080	0.29107	0.07143	0.06656	0.26219	0.57907	0.06250
15	0.04543	0.43007	0.18828	0.04651	0.04441	0.17580	0.39053	0.04255
20	0.03389	0.31655	0.13915	0.03448	0.03332	0.13222	0.29462	0.03226
25	0.02702	0.25046	0.11035	0.02740	0.02666	0.10595	0.23653	0.02597

Table 3: The risks of Squared log error, MLINEX and Quadratic loss functions for different sample sizes when $b=2$, $\gamma \neq 0$ and $\varpi = 3$.

n	$R_1(\hat{\theta}_1^*)$	$R_2(\hat{\theta}_2^*)$						$R_3(\hat{\theta}_3^*)$
		$\gamma = -3$	$\gamma = -2$	$\gamma = -1$	$\gamma = 1$	$\gamma = 2$	$\gamma = 3$	
5	0.10494	1.60729	0.68668	0.16667	0.14941	0.58476	1.28114	0.09091
10	0.05125	0.73272	0.32010	0.07895	0.07493	0.29624	0.65710	0.04762
15	0.03389	0.47482	0.20872	0.05172	0.04998	0.19833	0.44193	0.03226
20	0.02531	0.35124	0.15485	0.03846	0.03749	0.14906	0.33292	0.02439
25	0.02020	0.27871	0.12308	0.03061	0.03000	0.11940	0.26706	0.01961

Table 4: The risks of Squared log error, MLINEX and Quadratic loss functions for different sample sizes when $b=3$, $\gamma \neq 0$ and $\varpi = 4$.

n	$R_1(\hat{\theta}_1^*)$	$R_2(\hat{\theta}_2^*)$						$R_3(\hat{\theta}_3^*)$
		$\gamma = -3$	$\gamma = -2$	$\gamma = -1$	$\gamma = 1$	$\gamma = 2$	$\gamma = 3$	
5	0.06888	1.34161	0.58215	0.14286	0.13311	0.52438	1.15815	0.06250
10	0.03389	0.63310	0.27830	0.06897	0.06664	0.26444	0.58924	0.03226
15	0.02247	0.41440	0.18287	0.04545	0.04444	0.17679	0.39517	0.02174
20	0.01681	0.30802	0.13617	0.03390	0.03333	0.13278	0.29727	0.01639
25	0.01342	0.24510	0.10848	0.02703	0.02667	0.10631	0.23825	0.01316

4. Discussion

(a) The estimators $\hat{\theta}_1^*$, $\hat{\theta}_2^*$ and $\hat{\theta}_3^*$ are the minimax estimators of scale parameter θ in the following distributions for loss functions L_1 in (2.3), L_2 in (2.4) and L_3 in (2.5) respectively :

(i) negative exponential distribution when $b = 1$, $c = 1$; (ii) two-parameter gamma distribution when $b = b$, $c = 1$; (iii) Erlang distribution when $b = \text{positive integer}$, $c = c$; (iv) two-parameter Weibull distribution when $b = 1$, $c = c$; (v) Rayleigh distribution when $b = 1$, $c = 2$ and (vi) Maxwell distribution when $b = 3/2$, $c = 2$.

(b) When $b = 1$, the risks of the estimators corresponding to their respective loss functions such as Squared log error, MLINEX and Quadratic, are the same for scale parameter θ of negative exponential, two-parameter Weibull and Rayleigh distributions.

- (c) The risks of the estimators involved $\hat{\theta}_1^*$, $\hat{\theta}_2^*$ and $\hat{\theta}_3^*$ are known quantities n and b only but free from c . These risks are also independent of scale parameter θ and hence constant.
- (d) The risks of the estimators decrease among themselves when sample size n increases, provided $b > 0$, $\gamma \neq 0$ and $\varpi > 0$.
- (e) When $b > 0$, the risk $R_3(\hat{\theta}_3^*)$ is smaller than $R_1(\hat{\theta}_1^*)$ for any sample size n . The risks of the estimators $\hat{\theta}_1^*$ and $\hat{\theta}_3^*$ are also free from $\gamma \neq 0$ and $\varpi > 0$.
- (f) For fixed sample size, $b > 0$ and $\varpi > 0$ the risk $R_2(\hat{\theta}_2^*)$ increases at $\gamma \leq -1$ and $\gamma \geq 1$ only.
- (g) If sample size n increases, $b = 1$ and $\varpi = 1$, $R_2(\hat{\theta}_2^*)$ is minimum at $\gamma = -1$ and $\gamma = 1$ but maximum at $\gamma < -1$ and $\gamma > 1$ among the risks of the estimators.
- (h) For fixed sample size, $b = \frac{3}{2}$ and $\varpi = 2$ the relation $R_3(\hat{\theta}_3^*) < R_2(\hat{\theta}_2^*) < R_1(\hat{\theta}_1^*)$ holds when $\gamma = 1$ but $R_3(\hat{\theta}_3^*) < R_1(\hat{\theta}_1^*) < R_2(\hat{\theta}_2^*)$ holds only when $\gamma \leq -1$ and $\gamma > 1$.
- (i) When $b \geq 2$, $\varpi > 0$ and sample sizes are fixed, then the relation $R_3(\hat{\theta}_3^*) < R_1(\hat{\theta}_1^*) < R_2(\hat{\theta}_2^*)$ holds only at $\gamma \leq -1$ and $\gamma \geq 1$ respectively.

References

- [1] Abu-Taleb, A. A., Samadi, M. M. and Alawneh, J. A. (2007). Bayes estimation of the life time parameters for the exponential distribution, Journal of Math. and Statist., 3(3), 106-108.
- [2] Ahmed, A. O. and M. Ibrahim, N. A. (2011). Bayesian survival estimation for the Weibull distribution with censored data, Applied Science, 11(2), 393-396.
- [3] Brown L. D. (1968). Inadmissibility of the usual estimators of scale parameters in problems with unknown location and scale parameter. Ann. Math. Statist., 39, 28-48.
- [4] Jeffreys, H (1961). Theory of Probability, 3rd ed. Oxford; Clarendon Press.

- [5] Lehmann E. L, (1983). Theory of Point Estimation, New York, John Wiley.
- [6] Podder, C. K. (2019). A comparative study on two risks using a class of life-time distribution, International Journal of Statistical Sciences (IJSS), Vol. 17, 83-95.
- [7] Prakash G. and Singh D. C. (2010). Bayesian shrinkage estimation in a class of life testing distribution, Data Sciences Journal, 8(30), 243-258.
- [8] Son, Y. S. and Oh, M. (2006). Bayesian estimation of the two parameter gamma distribution, Commun.in Statist. Simul. and Comp., 35, 285-293.
- [9] Tyagi, R. K. and Bhattacharya, S. K. (1989). Bayes estimation of the Maxwell's velocity distribution function, Statistica, 29(4), 563-567.
- [10] Varian, H. R. (1975). A Bayesian Approach to Real Estate Assessment in studies in Bayesian Econometrics and Statistics in honour of L. J. Savage, eds, S E. Fienberg and Zellner, Amsterdam, North Holland publishing Co. 195-208.
- [11] Zellner, A. (1986). Bayesian estimation and Prediction using asymmetric loss function, Journal of American Statistical Association, 81, 446-451.