

Bayesian Estimation for the shape parameter of Exponentiated Rayleigh distribution

Ass. Prof .Dr. Nada S. Karam,^a Ahmed K. Jbur^b

^a Department of Mathematics, College of Education, Al-Mostanseriyah University, Baghdad, Iraq
nadaskaram@yahoo.com

^b Master's student of Mathematics, College of Education, Al-Mostanseriyah University, Baghdad, Iraq
ahmedmathematical@yahoo.com

Abstract. The paper is concerned with Bayesian analysis of Exponentiated Rayleigh distribution for complete samples. The Bayes estimators expressions for the shape parameter of the distribution have been derived under different priors and loss functions. The Gamma, Exponential, Chi-Square and Jeffrey prior have been assumed for posterior analysis. The Bayes estimation has been obtained under eight different loss functions (Squared error, Quadratic, Weighted, Linear exponential, Precautionary, Entropy, De Groot and non-Linear exponential loss functions). The study aims to find out a suitable estimator of the unknown shape parameter. The simulation study has been conducted to compare by mean square error (MSE) for the performance of various estimators.

Keywords: Bayesian Estimation, Exponentiated Rayleigh Distribution, Loss function, Prior, Posterior, (Squared error, Quadratic, Weighted, Linear exponential, Precautionary, Entropy, De Groot and non-Linear exponential) loss functions.

1 INTRODUCTION

For modeling data, Burr [Burr] introduced twelve different forms of cumulative distribution function. Among them Burr Type X and Burr Type XII received the maximum attention, where Surles and Padgett [Surles] observed that the Burr Type X distribution (Exponentiated Rayleigh distribution) can be used quite effectively in modeling strength data and also modeling general lifetime data. Several aspects of the ER distribution were studied in literature, see for example Sartawi and Abu-Salih [Sartawi], Jaheen [Jaheen], Ahmed, Fakhry and Jaheen [Ahmed], Karam and Jbur [Karam], Feroze and Aslam [Feroze] and Sindhu and Aslam [Sindhu]. The cumulative distribution function (CDF), and the probability density function (pdf) of the ER distribution with shape parameter ($\lambda > 0$) are respectively as follows: [Jonson]

$$F(x; \lambda) = (1 - e^{-x^2})^\lambda \quad x > 0; \lambda > 0 \quad (1)$$

$$f(x; \lambda) = 2\lambda x e^{-x^2} (1 - e^{-x^2})^{\lambda-1} \quad x > 0; \lambda > 0 \quad (2)$$

The random number X has been generated by inverse function method, which is for uniform random U :

$$X_i = \left(-\ln \left(1 - U_i^{\frac{1}{\lambda}} \right) \right)^{\frac{1}{2}} ; i = 1, 2, \dots, n \quad (3)$$

One of the important problems facing those who are interested in applied statistics is to study a certain phenomenon by The problem of estimating the unknown parameters in statistical distributions. This paper considers the estimations of the unknown shape parameter of the ER distribution, which is an important distribution in statistics and operations research and applied in several areas such as health, agriculture, biology, and other sciences. The main aim of this is to consider the Bayesian analysis of the unknown parameters under different priors (informative and non-informative) and loss functions for complete samples.

2 LIKELIHOOD FUNCTION

Let $X_1, X_2, \dots, X_n$ be a random sample from ER distribution with shape parameter $\lambda > 0$. Therefore the likelihood function of λ , from (1), for the observed sample is:

$$L(\lambda | \underline{x}) = \prod_{i=1}^n f(x_i, \lambda)$$

$$L(\lambda | \underline{x}) = \prod_{i=1}^n \left(2\lambda x e^{-x_i^2 + \ln(1-e^{-x_i^2})^{-1}} e^{-\lambda \ln(1-e^{-x_i^2})^{-1}} \right)$$

$$\therefore L(\lambda | \underline{x}) = Q_1 \lambda^n e^{-\lambda [\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}]}$$
(4)

Such that $Q_1 = 2^n e^{(\sum_{i=1}^n \ln x_i - \sum_{i=1}^n x_i^2 + \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1})}$

3 BAYESIAN ESTIMATORS USING DIFFERENT PRIOR AND LOSS FUNCTIONS

In this section Bayesian Estimators of the shape parameter for four different prior functions and under eight different loss functions has been determined.

• Types of loss functions

If $\hat{\lambda}$ represent of estimator for the shape parameter λ , then for:

1- Squared error loss function(slf): the squared error loss function defined as: [Singh]

$$L(\hat{\lambda}, \lambda) = c(\hat{\lambda} - \lambda)^2 ; \hat{\lambda}_{slf} = E(\lambda) \quad (5)$$

2- Quadratic loss function(qlf): the quadratic loss function defined as: [Singh]

$$L(\hat{\lambda}, \lambda) = \left(\frac{\hat{\lambda} - \lambda}{\lambda} \right)^2 ; \hat{\lambda}_{qlf} = \frac{E(\lambda^{-1})}{E(\lambda^{-2})} \quad (6)$$

3- Weighted loss function(wlf): the weighted loss function defined as: [Feroze]

$$L(\hat{\lambda}, \lambda) = \frac{(\hat{\lambda} - \lambda)^2}{\lambda} \quad ; \quad \hat{\lambda}_{wlf} = \frac{1}{E(\lambda^{-1})} \quad (7)$$

4- Linear exponential loss function(LINX): the (LINX) loss function defined as: [Islam]

$$L(\hat{\lambda}, \lambda) = (e^{c(\hat{\lambda} - \lambda)} - c(\hat{\lambda} - \lambda) - 1) \quad ; \quad \hat{\lambda}_{luf} = -\frac{1}{c} \ln E(e^{-c\lambda}) \quad (8)$$

5- Precautionary loss function(plf): the Precautionary loss function defined as: [Feroze]

$$L(\hat{\lambda}, \lambda) = \frac{(\lambda - \hat{\lambda})^2}{\hat{\lambda}} \quad ; \quad \hat{\lambda}_{plf} = \sqrt{E(\lambda^2)} \quad (9)$$

6- Entropy loss function(Elf): the entropy loss function defined as: [Islam]

$$L(\hat{\lambda}, \lambda) = ((\hat{\lambda}/\lambda)^t - t \ln(\hat{\lambda}/\lambda) - 1) \quad ; \quad \hat{\lambda}_{Elf} = (E(\lambda^{-t}))^{-1/t} \quad (10)$$

7- De Groot loss function(Dlf): the De Groot loss function defined as: [Islam]

$$L(\hat{\lambda}, \lambda) = \left(\frac{\lambda - \hat{\lambda}}{\hat{\lambda}}\right)^2 \quad ; \quad \hat{\lambda}_{Dlf} = \frac{E(\lambda^2)}{E(\lambda)} \quad (11)$$

8- Non- Linear exponential loss functions(NLINEX): the (NLINX) loss function defined as: [Sindhun]

$$L(\hat{\lambda}, \lambda) = (e^{c(\hat{\lambda} - \lambda)} + c(\hat{\lambda} - \lambda)^2 - c(\hat{\lambda} - \lambda) - 1) \\ \hat{\lambda}_{Nluf} = -\frac{1}{c+2} (\ln E(e^{-c\lambda}) - 2E(\lambda)) \quad (12)$$

• The Posterior distributions with different priors

For the given random variable X, the posterior density function of the shape parameter λ is well known as:

$$p(\lambda|\underline{x}) = \frac{L(\lambda|\underline{x}) \cdot p(\lambda)}{\int_0^\infty L(\lambda|\underline{x}) \cdot p(\lambda) d\lambda} \quad (13)$$

For Bayesian estimation, we specify four different prior distributions for the shape parameter, and which can be obtained four different posterior distributions under complete samples, as follows:

1- The Non-information prior, for any parameter λ , with pdf as: $p(\lambda) = \frac{1}{\lambda} \quad \lambda > 0$.

Then under the assumption of this prior distribution and by equation (13), the posterior distribution will be as:

$$p_j(\lambda|\underline{x}) = \frac{Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{1}{\lambda}}{\int_0^\infty Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{1}{\lambda} d\lambda}$$

$$\therefore p_J(\lambda/\underline{x}) = \frac{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1})^n}{\Gamma(n)} \lambda^{n-1} e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \quad \lambda > 0 \quad (14)$$

2– **The Chi-Square prior** is assumed to be: $p(\lambda) = \frac{\lambda^{(d/2)-1} e^{-\frac{\lambda}{2}}}{\Gamma(d/2) 2^{(d/2)}} \quad \lambda > 0, d > 0$

By equation (13) the posterior distribution under the assumption Chi-Square prior is:

$$p_{Ch}(\lambda/\underline{x}) = \frac{Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{\lambda^{(d/2)-1} e^{-\frac{\lambda}{2}}}{\Gamma((d/2)) 2^{(d/2)}}}{\int_0^\infty Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{\lambda^{(d/2)-1} e^{-\frac{\lambda}{2}}}{\Gamma((d/2)) 2^{(d/2)}} d\lambda}$$

$$p_{Ch}(\lambda|\underline{x}) = \frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + \frac{1}{2}}{\Gamma(n + \frac{d}{2})} \lambda^{n + (d/2) - 1} e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} - \frac{\lambda}{2}} \quad \lambda > 0 \quad (15)$$

3 – **The Exponential prior**, for any parameter λ , as: $p(\lambda) = b e^{-b\lambda} \quad \lambda > 0, b > 0$

By equation (13) the posterior distribution under the assumption Exponential prior is:

$$p_J(\lambda/\underline{x}) = \frac{Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} b e^{-b\lambda}}{\int_0^\infty Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} b e^{-b\lambda} d\lambda}$$

$$\therefore p_E(\lambda|\underline{x}) = \frac{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + b)^{n+1}}{\Gamma(n+1)} \lambda^n e^{-\lambda (\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + b)} \quad \lambda > 0 \quad (16)$$

4 – **The Gamma prior**, assumed to be: $p(\lambda) = \frac{b^a \lambda^{a-1}}{\Gamma(a)} e^{-\lambda b} \quad \lambda > 0, a, b > 0$

By equation (13) the posterior distribution under the assumption Gamma prior is:

$$p_G(\lambda/\underline{x}) = \frac{Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{b^a \lambda^{a-1}}{\Gamma(a)} e^{-\lambda b}}{\int_0^\infty Q_1 \lambda^n e^{-\lambda \sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} \frac{b^a \lambda^{a-1}}{\Gamma(a)} e^{-\lambda b} d\lambda}$$

$$p_G(\lambda|\underline{x}) = \frac{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + b)^{n+a}}{\Gamma(n+a)} \lambda^{n+a-1} e^{-\lambda (\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + b)} \quad \lambda > 0 \dots (17)$$

3.1 Bayesian Estimators under Non-information Prior using the eight different loss functions:

$$\hat{\lambda}_{1S} = \frac{n}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} ; \quad \hat{\lambda}_{1Q} = \frac{(n-2)}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} ;$$

$$\hat{\lambda}_{1W} = \frac{n-1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} ; \quad \hat{\lambda}_{1L} = \frac{-1}{c} \ln \left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + c} \right) ;$$

$$\hat{\lambda}_{1P} = \sqrt{\frac{(n)(n+1)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1})^2}} ; \quad \hat{\lambda}_{1E} = \left(\frac{\Gamma(n)}{\Gamma(n-t)}\right)^{(1/t)} / \left(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}\right) ;$$

$$\hat{\lambda}_{1D} = \frac{(n+1)}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}} ; \quad \hat{\lambda}_{1N} = \frac{-\left(\left(\ln\left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+c}}\right)^n\right)-2\left(\frac{n}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1}}\right)\right)}{c+2}$$

3.2 Bayesian Estimators under Chi-Square Prior using the eight different loss functions:

$$\hat{\lambda}_{2S} = \frac{n+\frac{d}{2}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}} ; \quad \hat{\lambda}_{2Q} = \frac{(n-\frac{d}{2}-2)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}})} ;$$

$$\hat{\lambda}_{2W} = \frac{n+\frac{d}{2}-1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}} ; \quad \hat{\lambda}_{2L} = \frac{-1}{c} \ln \left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}+c}} \right)^{n+\frac{d}{2}} ;$$

$$\hat{\lambda}_{2P} = \sqrt{\frac{(n+\frac{d}{2})(n+\frac{d}{2}+1)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}})^2}} ; \quad \hat{\lambda}_{2E} = \sqrt{\frac{\Gamma(n+\frac{d}{2})}{\Gamma(n+\frac{d}{2}-t)}} / \left(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}} \right) ;$$

$$\hat{\lambda}_{2D} = \frac{n+\frac{d}{2}+1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}} ; \quad \hat{\lambda}_{2N} = \frac{-\left(\left(\ln\left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}+c}}\right)^{n+\frac{d}{2}}\right)-2\left(\frac{n+\frac{d}{2}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+\frac{1}{2}}}\right)\right)}{c+2}$$

3.3 Bayesian Estimators under Exponential Prior using the eight different loss functions:

$$\hat{\lambda}_{3S} = \frac{n+1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} ; \quad \hat{\lambda}_{3Q} = \frac{(n-1)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b})} ;$$

$$\hat{\lambda}_{3W} = \frac{n}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} ; \quad \hat{\lambda}_{3L} = \frac{-1}{c} \ln \left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b+c}} \right)^{n+1} ;$$

$$\hat{\lambda}_{3P} = \sqrt{\frac{(n+1)(n+2)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b})^2}} ; \quad \hat{\lambda}_{3E} = \sqrt{\frac{\Gamma(n+1)}{\Gamma(n+1-t)}} / \left(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b} \right) ;$$

$$\hat{\lambda}_{3Dlf} = \frac{n+2}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} ;$$

$$\hat{\lambda}_{3Nlf} = \frac{-\left(\left(\ln\left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b+c}}\right)^{n+1}\right)-2\left(\frac{n+1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}\right)\right)}{c+2}$$

3-4 Bayesian Estimators under Gamma Prior using the eight different loss functions:

$$\hat{\lambda}_{4S} = \frac{n+a}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} ; \quad \hat{\lambda}_{4Q} = \frac{(n+a-2)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b})} ;$$

$$\hat{\lambda}_{4W} = \frac{n+a-1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} ; \quad \hat{\lambda}_{4L} = \frac{-1}{c} \ln \left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b+c}} \right)^{n+a} ;$$

$$\hat{\lambda}_{4P} = \sqrt{\frac{(n+a)(n+a+1)}{(\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b})^2}} ; \quad \hat{\lambda}_{4E} = \sqrt[t]{\frac{\Gamma(n+a)}{\Gamma(n+a-t)}} / (\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1} + b) ;$$

$$\hat{\lambda}_{4D} = \frac{n+a+1}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}$$

$$\hat{\lambda}_{4N} = \frac{-\left(\ln \left(\frac{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b+c}} \right)^{n+a} \right) - 2 \left(\frac{n+a}{\sum_{i=1}^n \ln(1-e^{-x_i^2})^{-1+b}} \right)}{c+2}$$

4 SIMULATION RESULTS AND CONCLUSIONS

In this section we mainly perform some simulation experiments to observe the behavior of the different Bayes estimators for the shape parameter proposed in the previous section (3), for different sample sizes ($n=10, 20, 30, 50, 75$ and 100) and for two different parameter values ($\lambda=1.5$ and 3), by applying the Monte Carlo simulation to compare the performance of these estimators using the mean squared error (MSE). All results are computed by MATLAB program based on (1000) replications and for two different cases { *case I*: ($a=3, b=0.8, d=3, c=1, t=2$; *case II*: $a=4, b=1.2, d=4, c=1, t=2$ for the formulas in sections 3.1), (3.2), (3.3), (3.4). The mean and MSE values for the Bayesian estimator $\hat{\lambda}$, are recorded in tables (1), (2), (3) and (4), and summarized generally in table (5).

Table (1): The mean value of $\hat{\lambda}$, when $\lambda = 1.5$

n	Jeffery priors							
	BS	BQ	BW	BP	BD	BL	BE	BNL
10	1.6567	1.3254	1.4910	1.7376	1.8224	1.5218	1.4058	1.6117
20	1.5575	1.4018	1.4796	1.5960	1.6354	1.4971	1.4402	1.5374
30	1.5501	1.4468	1.4985	1.5757	1.6018	1.5101	1.4724	1.5368
50	1.5261	1.4651	1.4956	1.5413	1.5567	1.5028	1.4803	1.5184
75	1.5217	1.4811	1.5014	1.5318	1.5420	1.5063	1.4912	1.5166
100	1.5193	1.4889	1.5041	1.5268	1.5345	1.5077	1.4965	1.5154
Chi-Square priors ($d=3$)								
10	1.7457	1.4421	1.5939	1.8200	1.8975	1.6161	1.5161	1.7025
20	1.6085	1.4589	1.5337	1.6455	1.6833	1.5487	1.4958	1.5886
30	1.5852	1.4846	1.5349	1.6102	1.6355	1.5454	1.5095	1.5719
50	1.5478	1.4877	1.5178	1.5628	1.5779	1.5246	1.5027	1.5401
75	1.5364	1.4962	1.5163	1.5464	1.5564	1.5209	1.5062	1.5312
100	1.5303	1.5002	1.5152	1.5378	1.5454	1.5188	1.5077	1.5265
Exponential priors ($b=0.8$)								
10	1.5909	1.3016	1.4463	1.6616	1.7355	1.4786	1.3721	1.5535
20	1.5350	1.3888	1.4619	1.5712	1.6081	1.4794	1.4249	1.5165
30	1.5361	1.4370	1.4865	1.5607	1.5856	1.4981	1.4615	1.5234
50	1.5188	1.4593	1.4890	1.5336	1.5486	1.4962	1.4741	1.5113

75	1.5171	1.4771	1.4971	1.5270	1.5370	1.5019	1.4871	1.5120
100	1.5158	1.4858	1.5008	1.5233	1.5309	1.5045	1.4933	1.5121
Gamma priors ($a = 3, b = 0.8$)								
10	1.8802	1.5909	1.7355	1.9511	2.0248	1.7475	1.6616	1.8359
20	1.6812	1.5350	1.6081	1.7174	1.7543	1.6203	1.5712	1.6609
30	1.6352	1.5361	1.5856	1.6598	1.6847	1.5947	1.5607	1.6217
50	1.5784	1.5188	1.5486	1.5932	1.6082	1.5549	1.5336	1.5705
75	1.5570	1.5171	1.5370	1.5669	1.5769	1.5414	1.5270	1.5518
100	1.5459	1.5158	1.5309	1.5533	1.5609	1.5343	1.5233	1.5420
Chi-Square priors ($d = 4$)								
10	1.8216	1.5180	1.6698	1.8960	1.9734	1.6864	1.5921	1.7765
20	1.6459	1.4963	1.5711	1.6829	1.7207	1.5847	1.5332	1.6255
30	1.6104	1.5097	1.5601	1.6354	1.6607	1.5699	1.5347	1.5969
50	1.5628	1.5027	1.5328	1.5778	1.5929	1.5394	1.5177	1.5550
75	1.5400	1.5000	1.5200	1.5500	1.5600	1.5246	1.5100	1.5349
100	1.5336	1.5035	1.5186	1.5411	1.5486	1.5221	1.5110	1.5298
Exponential priors ($b = 1.2$)								
10	1.4975	1.2252	1.3613	1.5640	1.6336	1.3982	1.2915	1.4644
20	1.4896	1.3477	1.4187	1.5246	1.5605	1.4372	1.3827	1.4721
30	1.5053	1.4082	1.4567	1.5294	1.5538	1.4688	1.4322	1.4931
50	1.5006	1.4417	1.4712	1.5152	1.5300	1.4785	1.4564	1.4932
75	1.4988	1.4593	1.4791	1.5086	1.5185	1.4840	1.4692	1.4939
100	1.5026	1.4728	1.4877	1.5100	1.5175	1.4914	1.4803	1.4989
Gamma priors ($a = 4, b = 1.2$)								
10	1.9059	1.6336	1.7697	1.9727	2.0420	1.7795	1.7003	1.8637
20	1.7024	1.5605	1.6315	1.7375	1.7733	1.6425	1.5956	1.6824
30	1.6509	1.5538	1.6024	1.6750	1.6995	1.6109	1.5779	1.6376
50	1.5889	1.5300	1.5594	1.6035	1.6183	1.5655	1.5446	1.5811
75	1.5579	1.5185	1.5382	1.5678	1.5777	1.5426	1.5283	1.5528
100	1.5472	1.5175	1.5324	1.5546	1.5621	1.5357	1.5249	1.5434

Table (2): The mean value of $\hat{\lambda}$, when $\lambda = 3$

n	Jeffery priors							
	BS	BQ	BW	BP	BD	BL	BE	BNL
10	3.3134	2.6507	2.9821	3.4751	3.6447	2.8284	2.8115	3.1517
20	3.1150	2.8035	2.9593	3.1919	3.2708	2.8857	2.8803	3.0386
30	3.0768	2.8717	2.9743	3.1277	3.1794	2.9249	2.9225	3.0262
50	3.0627	2.9401	3.0014	3.0931	3.1239	2.9708	2.9706	3.0320
75	3.0479	2.9666	3.0073	3.0682	3.0886	2.9868	2.9869	3.0276
100	3.0352	2.9745	3.0048	3.0503	3.0655	2.9896	2.9896	3.0200
Chi-Square priors ($d = 3$)								
10	3.2255	2.6646	2.9451	3.3629	3.5060	2.8222	2.8013	3.0911
20	3.0959	2.8079	2.9519	3.1671	3.2399	2.8851	2.8790	3.0256
30	3.0685	2.8737	2.9711	3.1168	3.1659	2.9245	2.9220	3.0205
50	3.0589	2.9401	2.9995	3.0885	3.1183	2.9699	2.9697	3.0293
75	3.0461	2.9665	3.0063	3.0660	3.0860	2.9863	2.9864	3.0262
100	3.0342	2.9744	3.0043	3.0491	3.0641	2.9893	2.9893	3.0192
Exponential priors ($b = 0.8$)								
10	2.8301	2.3155	2.5728	2.9559	3.0874	2.5037	2.4408	2.7213
20	2.8938	2.6182	2.7560	2.9618	3.0315	2.7051	2.6862	2.8309
30	2.9317	2.7426	2.8371	2.9786	3.0263	2.7981	2.7895	2.8872
50	2.9752	2.8585	2.9169	3.0042	3.0335	2.8902	2.8875	2.9469
75	2.9901	2.9114	2.9507	3.0097	3.0294	2.9320	2.9310	2.9707
100	2.9921	2.9329	2.9625	3.0069	3.0218	2.9483	2.9477	2.9775
Gamma priors ($a = 3, b = 0.8$)								

10	3.3447	2.8301	3.0874	3.4709	3.6019	2.9590	2.9559	3.2161
20	3.1693	2.8938	3.0315	3.2375	3.3071	2.9627	2.9618	3.1005
30	3.1209	2.9317	3.0263	3.1678	3.2154	2.9786	2.9786	3.0735
50	3.0919	2.9752	3.0335	3.1209	3.1502	3.0035	3.0042	3.0624
75	3.0687	2.9901	3.0294	3.0883	3.1081	3.0092	3.0097	3.0489
100	3.0514	2.9921	3.0218	3.0662	3.0810	3.0067	3.0069	3.0365
Chi-Square priors ($d = 4$)								
10	3.3658	2.8048	3.0853	3.5032	3.6463	2.9449	2.9417	3.2255
20	3.1679	2.8799	3.0239	3.2391	3.3119	2.9522	2.9510	3.0960
30	3.1391	2.9429	3.0410	3.1878	3.2372	2.9904	2.9916	3.0895
50	3.0785	2.9601	3.0193	3.1079	3.1377	2.9892	2.9895	3.0487
75	3.0616	2.9821	3.0219	3.0815	3.1014	3.0016	3.0019	3.0416
100	3.0525	2.9926	3.0225	3.0674	3.0824	3.0072	3.0075	3.0374
Exponential priors ($b = 1.2$)								
10	2.5521	2.0880	2.3200	2.6655	2.7841	2.2849	2.2010	2.4630
20	2.7373	2.4766	2.6070	2.8018	2.8677	2.5684	2.5410	2.6810
30	2.8401	2.6568	2.7485	2.8855	2.9317	2.7143	2.7023	2.7982
50	2.8970	2.7833	2.8402	2.9252	2.9538	2.8163	2.8116	2.8701
75	2.9390	2.8617	2.9004	2.9583	2.9777	2.8830	2.8810	2.9204
100	2.9599	2.9013	2.9306	2.9745	2.9892	2.9170	2.9159	2.9456
Gamma priors ($a = 4$, $b = 1.2$)								
10	2.9599	2.9013	2.9306	2.9745	2.9892	2.9170	2.9159	2.9456
20	3.1284	2.8677	2.9980	3.1929	3.2587	2.9353	2.9321	3.0640
30	3.1149	2.9317	3.0233	3.1604	3.2065	2.9770	2.9771	3.0690
50	3.0674	2.9538	3.0106	3.0956	3.1242	2.9820	2.9820	3.0389
75	3.0551	2.9777	3.0164	3.0743	3.0937	2.9968	2.9970	3.0356
100	3.0478	2.9892	3.0185	3.0624	3.0771	3.0036	3.0038	3.0331

Table (3): The MSE value of $\hat{\lambda}$, when $\lambda = 1.5$

n	Jeffery								Best
	BS	BQ	BW	BP	BD	BL	BE	BNL	
10	0.3375	0.2308	0.2536	0.4007	0.4826	0.2197	0.2342	0.2923	BL
20	0.1335	0.1151	0.1179	0.1459	0.1618	0.1082	0.1149	0.1240	BL
30	0.0903	0.0793	0.0820	0.0964	0.1041	0.0789	0.0799	0.0861	BL
50	0.0498	0.0465	0.0472	0.0518	0.0543	0.0461	0.0466	0.0484	BL
75	0.0317	0.0300	0.0304	0.0327	0.0339	0.0301	0.0301	0.0311	BQ
100	0.0243	0.0231	0.0235	0.0249	0.0256	0.0233	0.0233	0.0239	BQ
Chi-Square priors ($d = 3$)									
10	0.3454	0.1979	0.2465	0.4123	0.4948	0.2208	0.2153	0.2987	BQ
20	0.1362	0.1040	0.1143	0.1514	0.1699	0.1072	0.1076	0.1255	BQ
30	0.0940	0.0763	0.0826	0.1017	0.1107	0.0802	0.0788	0.0890	BQ
50	0.0512	0.0453	0.0473	0.0538	0.0569	0.0466	0.0461	0.0495	BQ
75	0.0325	0.0296	0.0307	0.0338	0.0352	0.0304	0.0300	0.0318	BQ
100	0.0248	0.0230	0.0237	0.0256	0.0265	0.0236	0.0233	0.0244	BQ
Exponential priors ($b = 0.8$)									
10	0.2211	0.1818	0.1788	0.2583	0.3088	0.1581	0.1747	0.1963	BL
20	0.1079	0.0997	0.0982	0.1168	0.1288	0.0909	0.0976	0.1014	BL
30	0.0801	0.0730	0.0740	0.0851	0.0913	0.0711	0.0728	0.0768	BL
50	0.0466	0.0443	0.0445	0.0483	0.0504	0.0435	0.0442	0.0454	BL
75	0.0304	0.0290	0.0293	0.0312	0.0322	0.0289	0.0291	0.0298	BL
100	0.0235	0.0225	0.0228	0.0240	0.0247	0.0226	0.0226	0.0232	BQ
Gamma priors ($a = 3$, $b = 0.8$)									
10	0.4418	0.2211	0.3088	0.5236	0.6202	0.2814	0.2583	0.3831	BQ
20	0.1608	0.1079	0.1288	0.1808	0.2040	0.1230	0.1168	0.1472	BQ
30	0.1076	0.0801	0.0913	0.1176	0.1290	0.0896	0.0851	0.1012	BQ

50	0.0561	0.0466	0.0504	0.0595	0.0635	0.0500	0.0483	0.0539	BQ
75	0.0349	0.0304	0.0322	0.0366	0.0384	0.0321	0.0312	0.0339	BQ
100	0.0263	0.0235	0.0247	0.0273	0.0284	0.0246	0.0240	0.0257	BQ
Chi-Square priors ($d = 4$)									
10	0.4138	0.2159	0.2897	0.4931	0.5884	0.2605	0.2456	0.3571	BQ
20	0.1516	0.1077	0.1238	0.1697	0.1911	0.1169	0.1142	0.1390	BQ
30	0.1017	0.0788	0.0876	0.1107	0.1210	0.0855	0.0825	0.0959	BQ
50	0.0538	0.0461	0.0490	0.0569	0.0604	0.0485	0.0473	0.0519	BQ
75	0.0335	0.0302	0.0315	0.0348	0.0363	0.0312	0.0308	0.0327	BQ
100	0.0243	0.0223	0.0231	0.0251	0.0260	0.0230	0.0226	0.0238	BQ
Exponential priors ($b = 1.2$)									
10	0.1654	0.1863	0.1559	0.1846	0.2147	0.1355	0.1665	0.1526	BL
20	0.0932	0.0994	0.0911	0.0982	0.1059	0.0835	0.0940	0.0893	BL
30	0.0725	0.0719	0.0698	0.0757	0.0802	0.0666	0.0702	0.0702	BL
50	0.0440	0.0440	0.0431	0.0451	0.0466	0.0419	0.0434	0.0432	BL
75	0.0293	0.0295	0.0290	0.0298	0.0305	0.0284	0.0291	0.0290	BL
100	0.0218	0.0217	0.0215	0.0221	0.0225	0.0212	0.0215	0.0216	BL
Gamma priors ($a = 4, b = 1.2$)									
10	0.4327	0.2147	0.3038	0.5106	0.6013	0.2808	0.2534	0.3774	BQ
20	0.1626	0.1059	0.1290	0.1831	0.2067	0.1243	0.1160	0.1489	BQ
30	0.1100	0.0802	0.0927	0.1204	0.1322	0.0912	0.0858	0.1033	BQ
50	0.0572	0.0466	0.0511	0.0610	0.0652	0.0508	0.0486	0.0549	BQ
75	0.0351	0.0305	0.0324	0.0367	0.0385	0.0323	0.0313	0.0341	BQ
100	0.0253	0.0225	0.0237	0.0263	0.0274	0.0237	0.0231	0.0247	BQ

 Table (4): The mean value of $\hat{\lambda}$, when $\lambda = 3$

1Xn	Jeffery								Best
	BS	BQ	BW	BP	BD	BL	BE	BNL	
10	1.3500	0.9231	1.0143	1.6027	1.9304	0.6826	0.9368	1.0531	BL
20	0.5339	0.4604	0.4716	0.5836	0.6474	0.3804	0.4595	0.4678	BL
30	0.3118	0.2830	0.2865	0.3324	0.3589	0.2547	0.2820	0.2870	BL
50	0.2010	0.1852	0.1893	0.2097	0.2204	0.1751	0.1863	0.1903	BL
75	0.1301	0.1222	0.1245	0.1342	0.1391	0.1180	0.1229	0.1252	BL
100	0.0944	0.0902	0.0914	0.0967	0.0994	0.0878	0.0905	0.0917	BL
Chi-Square priors ($d = 3$)									
10	0.8691	0.6709	0.6851	1.0211	1.2228	0.5093	0.6566	0.7027	BL
20	0.4267	0.3804	0.3819	0.4649	0.5148	0.3203	0.3757	0.3794	BL
30	0.2779	0.2555	0.2569	0.2955	0.3183	0.2306	0.2538	0.2570	BL
50	0.1879	0.1740	0.1774	0.1959	0.2057	0.1646	0.1748	0.1783	BL
75	0.1246	0.1173	0.1193	0.1284	0.1331	0.1132	0.1179	0.1200	BL
100	0.0914	0.0874	0.0885	0.0936	0.0962	0.0851	0.0877	0.0889	BL
Exponential priors ($b = 0.8$)									
10	0.5527	0.8191	0.6154	0.5734	0.6310	0.5672	0.7023	0.5282	BNL
20	0.3387	0.4138	0.3565	0.3445	0.3603	0.3323	0.3806	0.3273	BNL
30	0.2391	0.2714	0.2461	0.2425	0.2505	0.2350	0.2566	0.2333	BNL
50	0.1688	0.1752	0.1685	0.1715	0.1759	0.1616	0.1710	0.1647	BL
75	0.1152	0.1170	0.1146	0.1167	0.1190	0.1110	0.1154	0.1130	BL
100	0.0862	0.0873	0.0859	0.0871	0.0884	0.0838	0.0864	0.0850	BL
Gamma priors ($a = 3, b = 0.8$)									
10	0.8504	0.5527	0.6310	1.0096	1.2108	0.4499	0.5734	0.6760	BL
20	0.4214	0.3387	0.3603	0.4662	0.5220	0.2957	0.3445	0.3684	BL
30	0.2803	0.2391	0.2505	0.3019	0.3284	0.2205	0.2425	0.2554	BL
50	0.1900	0.1688	0.1759	0.1996	0.2111	0.1615	0.1715	0.1787	BL
75	0.1260	0.1152	0.1190	0.1306	0.1361	0.1122	0.1167	0.1206	BL
100	0.0923	0.0862	0.0884	0.0949	0.0979	0.0845	0.0871	0.0892	BL

	Chi-Square priors ($d = 4$)								
10	1.0247	0.6568	0.7559	1.2184	1.4633	0.5232	0.6840	0.8069	BL
20	0.4654	0.3757	0.3989	0.5142	0.5751	0.3239	0.3818	0.4058	BL
30	0.3415	0.2864	0.3040	0.3675	0.3989	0.2645	0.2927	0.3103	BL
50	0.1937	0.1750	0.1807	0.2028	0.2138	0.1666	0.1769	0.1828	BL
75	0.1253	0.1156	0.1188	0.1297	0.1349	0.1122	0.1168	0.1201	BL
100	0.0965	0.0901	0.0924	0.0992	0.1023	0.0883	0.0910	0.0933	BL
	Exponential priors ($b = 1.2$)								
10	0.5439	1.0614	0.7460	0.4863	0.4551	0.7326	0.8937	0.5879	BD
20	0.3260	0.4843	0.3876	0.3086	0.2996	0.3829	0.4322	0.3377	BD
30	0.2540	0.3176	0.2772	0.2488	0.2480	0.2717	0.2954	0.2560	BD
50	0.1631	0.1877	0.1722	0.1611	0.1607	0.1699	0.1792	0.1638	BD
75	0.1095	0.1194	0.1130	0.1089	0.1091	0.1116	0.1158	0.1095	BP
100	0.0860	0.0908	0.0876	0.0859	0.0862	0.0865	0.0890	0.0857	BP
	Gamma priors ($a = 4$, $b = 1.2$)								
10	0.6175	0.4551	0.4796	0.7268	0.8687	0.3668	0.4530	0.5034	BL
20	0.3522	0.2996	0.3083	0.3869	0.4312	0.2609	0.2995	0.3123	BL
30	0.2879	0.2480	0.2593	0.3085	0.3338	0.2292	0.2515	0.2636	BL
50	0.1755	0.1607	0.1648	0.1833	0.1928	0.1529	0.1619	0.1663	BL
75	0.1174	0.1091	0.1117	0.1213	0.1260	0.1058	0.1100	0.1127	BL
100	0.0918	0.0862	0.0881	0.0943	0.0972	0.0844	0.0870	0.0889	BL

Table (5) the best performances for loss function and prior distribution

Test	prior	(Doubly); MSE					
		10	20	30	50	75	100
1	Jeffery	BL	BL	BL	BL	BQ	BQ
	Chi-Sq.	BQ	BQ	BQ	BQ	BQ	BQ
	Exp.	BL	BL	BL	BL	BL	BQ
	Gamma	BQ	BQ	BQ	BQ	BQ	BQ
	Best prior	Exp.	Exp.	Exp.	Exp.	Exp.	Exp.
2	Jeffery	BL	BL	BL	BL	BQ	BQ
	Chi-Sq.	BQ	BQ	BQ	BQ	BQ	BQ
	Exp.	BL	BL	BL	BL	BL	BL
	Gamma	BQ	BQ	BQ	BQ	BQ	BQ
	Best prior	Exp.	Exp.	Exp.	Exp.	Exp.	Exp.
3	Jeffery	BL	BL	BL	BL	BL	BL
	Chi-Sq.	BL	BL	BL	BL	BL	BL
	Exp.	BNL	BNL	BNL	BL	BL	BL
	Gamma	BL	BL	BL	BL	BL	BL
	Best prior	Gamma	Gamma	Gamma	Gamma	Exp.	Exp.
4	Jeffery	BL	BL	BL	BL	BL	BL
	Chi-Sq.	BL	BL	BL	BL	BL	BL
	Exp.	BD	BD	BD	BD	BP	BP
	Gamma	BL	BL	BL	BL	BL	BL
	Best prior	Gamma	Gamma	Gamma	Gamma	Gamma	Gamma

From the tables (3) and (4), the rate of converges of the estimates towards the true value of the shape parameter increases with increase of sample size. From the best performance table (5), the estimates with ($\lambda = 1.5$) in the experiments (1 and 2), the performance of Bayes estimator under Exponential prior using linear Exponential loss function is better than other prior different loss

function used and for each sample sizes in every experiment. But when comparing the three experiments the Exponential prior ($b = 1.2$) using linear Exponential loss function is better than other prior using different loss function and for each sample sizes. And the estimates with ($\lambda = 3$) in experiments (3 and 4), the performance of Bayes estimator under Gamma prior ($a = 4, b = 1.2$) using linear Exponential loss function is better than other prior using different loss functions and for each sample sizes.

5 CONCLUSIONS

The above study suggests that in order to estimate the parameter of Burr type X distribution under a Bayesian framework, when $\lambda = 1.5$ that the performance of Bayes estimator under Exponential prior ($b = 1.2$) with linear loss function, records full appearance "for all sample sizes", as the best prior distribution and using different loss functions and when $\lambda = 3$, that the performance of Bayes estimator under Gamma prior ($a = 4, b = 1.2$) with linear loss function, records full appearance "for all sample sizes", as the best prior distribution and using different loss functions for the complete data.

References

- Ahmad, K. E. Fakhry, Jaheen, Z. F., (1997). *Empirical Bayes estimation of $p(y < x)$ and characterization of Burr Type X model*. Journal of Statistical Planning and Inference, 46, 297-308.
- Burr, I. W. (1942). *Cumulative frequency distribution*. Annals of Mathematical Statistics, Vol. 13, 215-232.
- Feroze, N. and Aslam, M., (2012). *Bayesian Analysis of Burr Type X distribution under Complete and Censored Samples*. International Journal of Pure and Applied Sciences and Technology, 11(2), pp. 16-28.
- Islam, s. (2011). *Loss functions, Utility functions and Bayesian Sample Size Determination*. PH.D. Thesis, University of London.
- Jaheen, Z. F., (1995). *Bayesian approach to prediction with outliers from the Burr Type X model*. Microelectronic Reliability. 35, 45-47.
- Jonson, N. L., Kotz, S. and Balakrishnan, N. (1994). *Continuous univariate distribution*. Volume 1, Second Edition, John and Wiley Sons, Inc., New York.
- Karam, N. S. and Jbur A.K., (2014) "Bayesian Analyses of the Burr Type X Distribution under Doubly Type II censored samples using different Priors and Loss functions" Mathematical Theory and Modeling, Vol. 4, No. 14, pp 9-18.
- Sartawi, H. A. and Abu-Salih M.S., (1991). *Bayes prediction bounds for the Burr Type X model*. Comm., Statist. Theory Methods 20, 2307-2330.
- Sindhu, T. N. and Aslam, M., (2014). *On the Parameter of the Burr Type X under Bayesian Principles*. International Journal of Mathematical, Computational, Physical and Quantum Engineering Vol. 8 no. 1, 2014
- Sindhu, T. N., Feroze, N. and Aslam M. (2013). *Bayesian Analysis of the Kumaraswamy Distribution under Failure Censoring Sampling Scheme*. International Journal of Advanced Science and Technology, vol. 51, 39-58.

- Singh, S. K., Singh, U. and Kumar, D. (2011). *Bayesian Estimation of the Exponentiated Gamma Parameter and Reliability function under asymmetric loss function*. Statistical Journal Volume 9, Number 3, 247–260.
- Surles, J. G. and Padgett, W. J., (2001). *Inference for reliability and stress -strength for a scale Burr Type X distribution*. Lifetime Data Analysis. 7, 187-200.