



ISSN NO. 2320-5407

Journal homepage: <http://www.journalijar.com>

INTERNATIONAL JOURNAL
OF ADVANCED RESEARCH

RESEARCH ARTICLE

Using Generalized Square Loss Function to Estimate the Shape Parameter of the Burr Type XII Distribution

Tasnim H.K. AlBaldawi, Huda A. Rasheed, Saleemaha H. Jaseim

- Dept. of Math. College of Science University of Baghdad /Iraq
- Dept. of Math. College of Science University of Almustansiria/ Iraq
- Dept. of Astronomy College of Science University of Baghdad/ Iraq

Manuscript Info

Manuscript History:

Received: 12 March 2015
Final Accepted: 22 April 2015
Published Online: May 2015

Key words: Burr type XII distribution, Generalized square error loss function, MinMSE estimator.

*Corresponding Author

Huda A. Rasheed

Abstract

In this paper we compare the performance of Bayes estimators for the shape parameter of the Burr Type XII distribution with the classical minimum mean square error (MinMSE) estimator. We considered the generalized square loss function (GSLF) with Jeffreys prior, as well as the exponential prior. The comparison was made through a Monte Carlo simulation study with respect to the mean square error MSE. The results of comparison show that Bayes estimators of the shape parameter under the GSLF with the exponential prior gives better results among other estimators. Accordingly; if adequate information is available about the parameters it is preferable to use conjugate informative priors.

Copy Right, IJAR, 2015,. All rights reserved

INTRODUCTION

The Burr type XII distribution is one of the most important members in the family of Burr distributions. It is one of the most popular distribution of failure time, life testing and reliability theory. Its capacity of the various shapes often permits a good fit when used to describe biological, clinical or other experimental data [7].

Burr type XII distribution has been studied by many authors some of whom are Soliman *et al* [8], Wu *et al* [11], Jang *et al* [2], Makhdoom and Jafari [3], Panahi and Asadi [4], Sinhu and Aslam [6] and AL-Saiari *et al* [1].

The Bayesian approach has a lot of advantages in comparison with the classical approach; it can utilize the information in a formal way, satisfies the axioms of coherence and utilize decision theory [9].

A comparison of priors made by Taher, M. and Saleem, M. (2011) showed that the informative priors are more advantageous than the non-informative priors [10]. In this research we will consider Bayes estimators under Jeffreys prior information as a non informative prior and the exponential prior as an informative prior with the generalized square loss function.

The Model

Let $t_1, t_2, \dots, t_n$ be independent identically distributed lifetimes from Burr type XII distribution. The probability density function is given by:

$$f(t|\theta, \lambda) = \frac{\theta \lambda t^{\lambda-1}}{(1+t^\lambda)^{\theta+1}} \quad 0 < t < \infty; \theta, \lambda > 0 \quad (1)$$

where θ and λ are the shape and scale parameters, respectively.

The cumulative distribution function and the reliability function are given, respectively, by:

$$F(t) = 1 - (1 + t^\lambda)^{-\theta}, \text{ and, } R(t) = (1 + t^\lambda)^{-\theta}$$

The likelihood function (when λ is known) for the given the sample observation is:

$$L(\theta) \propto \theta^n \left[\prod_{i=1}^n (1 + t_i^\lambda) \right]^{-(\theta-1)} \quad (2)$$

And the log-likelihood function is

$$\ell(\theta) \propto n \ln \theta - (\theta - 1) \sum_{i=1}^n \ln(1 + t_i^\lambda)$$

Yarmohammadi and Pazira [12] obtained three classical estimators and three minimax estimators of the shape parameter θ for the Burr type XII distribution with the corresponding mean squared error. The classical estimators are the maximum likelihood estimator (MLE), the uniformly minimum variance unbiased estimator (UMVUE) and the minimum mean squared error estimator (MinMSE). The results are summarized in the following table:

		$\hat{\theta}_{MinMSE} = \frac{n-2}{\sum_{i=1}^n \ln(1 + t_i^\lambda)}$
$MSE_{\hat{\theta}_{MLE}} = \frac{(n+2)\theta^2}{(n-1)(n-2)}$	$MSE_{\hat{\theta}_{umvue}} = \frac{\theta^2}{n-2}$	$MSE_{\hat{\theta}_{MinMSE}} = \frac{\theta^2}{n-1}$

Further, they showed that

$$MSE_{\hat{\theta}_{MinMSE}} < MSE_{\hat{\theta}_{umvue}} < MSE_{\hat{\theta}_{MLE}}$$

They finally showed the minimax estimator under the weighted balanced squared error (MWBSE) loss function is the best and that for large sample size the MinMSE and the MWBSE are identical.

Hence, in this research we will consider the MinMSE estimator since it has the smallest MSE along with the Bayes estimators under Jeffreys prior information as a non informative prior and the exponential prior as an informative prior with the generalized square loss function.

Bayes Estimators under the generalized square loss function

It is well known that, the performance of Bayes estimators depend on the form of the prior distribution and the loss function assumed. We will consider the generalized square loss function (GSLF) to obtain our Bayes estimators. It is introduced as follows [5]

$$l(\hat{\theta}, \theta) = (\sum_{j=0}^k a_j \theta^j)(\hat{\theta} - \theta)^2, k = 0, 1, 2, 3 \dots \quad (3)$$

where $a_j, (j = 1, 2, 3, \dots, k)$ is a constant.

The risk function is the posterior expectation of the loss function $l(\hat{\theta}, \theta)$ with respect to $h(\theta|t)$. That is

$$\begin{aligned} R(\hat{\theta}, \theta) &= \int_0^\infty l(\hat{\theta}, \theta) h(\theta|t) d\theta \\ &= \int_0^\infty (a_0 + a_1 \theta + a_2 \theta^2 + \dots + a_k \theta^k)(\hat{\theta}^2 - 2\hat{\theta}\theta + \theta^2) h(\theta|t) d\theta \end{aligned}$$

Thus,

$$\begin{aligned} R(\hat{\theta}, \theta) &= a_0 \hat{\theta}^2 - 2a_0 \hat{\theta} E(\theta|t) + a_0 E(\theta^2|t) + a_1 \hat{\theta}^2 E(\theta|t) - 2a_1 \hat{\theta} E(\theta^2|t) + a_1 E(\theta^3|t) + \dots + a_k \hat{\theta}^2 E(\theta^k|t) \\ &\quad - 2a_k \hat{\theta} E(\theta^{k+1}|t) + a_k E(\theta^{k+2}|t) \end{aligned}$$

The value of $\hat{\theta}$ that minimizes the posterior risk $R(\hat{\theta}, \theta)$ is obtained by setting its first partial derivative with respect to $\hat{\theta}$ equal to zero. That is $\hat{\theta}$ is the solution for the equation

1. Posterior distribution under non-informative prior

The Jeffreys prior information for the parameter θ is given by

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

where $I(\theta)$ is the fisher information.

For the Burr type XII distribution, the fisher information is

$$I(\theta) = \frac{n}{\theta^2}$$

Hence

$$g_1(\theta) = c \frac{\sqrt{n}}{\theta}, \text{ where } c \text{ is a constant} \quad (4)$$

Combining the prior distribution (4) with the likelihood function (2), the posterior distribution of θ is

$$g_1(\theta|t) = \frac{\theta^{n-1} \left[\sum_{i=1}^n \ln(1+t_i^\lambda) \right]^n e^{-\theta \sum_{i=1}^n \ln(1+t_i^\lambda)}}{\Gamma(n)} \quad (5)$$

Thus, the posterior density has a gamma distribution with parameters

$\left(n, \sum_{i=1}^n \ln(1+t_i^\lambda) \right)$, with the following r^{th} moment

$$\begin{aligned} E(\theta^r|t) &= \int_0^\infty \theta^r h(\theta|t) d\theta = \int_0^\infty \theta^r \frac{\theta^{n-1} \left[\sum_{i=1}^n \ln(1+t_i^\lambda) \right]^n e^{-\theta \sum_{i=1}^n \ln(1+t_i^\lambda)}}{\Gamma(n)} d\theta \\ &= \frac{\Gamma(n+r)}{\Gamma(n) \left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^r} \end{aligned}$$

Hence, Bayes estimator under the generalized square loss function is given by

$$\hat{\theta}_1 = \frac{a_0 \frac{n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)} + a_1 \frac{(n+1)n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^2} + a_2 \frac{(n+2)(n+1)n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^3} + \dots + a_k \frac{(n+k)(n+k-1)\dots n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^{k+1}}}{a_0 + a_1 \frac{n}{\sum_{i=1}^n \ln(1+t_i^\lambda)} + a_2 \frac{(n+1)n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^2} + \dots + a_k \frac{(n+k-1)\dots n}{\left(\sum_{i=1}^n \ln(1+t_i^\lambda) \right)^k}} \quad (6)$$

2. Posterior distribution under exponential prior distribution.

This is a conjugate prior distribution. It is given as

$$g_2(\theta) = \delta e^{-\delta\theta}; \quad \theta > 0, \delta > 0 \quad (7)$$

Combining the prior distribution (7) with the likelihood function (2), the posterior distribution of θ under the exponential prior is

$$g_2(\theta|t) = \frac{\theta^{n+1-1} \left[\delta + \sum_{i=1}^n \ln(1+t_i^\lambda) \right]^{n+1} e^{-\theta(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))}}{\Gamma(n+1)}$$

Again, the posterior density has a gamma distribution with parameters

$$\left((n+1), \left(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda) \right) \right)$$

Hence, Bayes estimator under the generalized square loss function is given by

$$\hat{\theta}_2 = \frac{a_0 \frac{(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))} + a_1 \frac{(n+2)(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))^2} + a_2 \frac{(n+3)(n+2)(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))^3} + \dots + a_k \frac{(n+k+1)\dots(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))^k}}{a_0 + a_1 \frac{(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))} + a_2 \frac{(n+2)(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))^2} + \dots + a_k \frac{(n+k)\dots(n+1)}{(\delta + \sum_{i=1}^n \ln(1+t_i^\lambda))^k}} \quad (8)$$

Simulation and Results

In order to investigate the performance of the estimators obtained in the above section, a simulation study is conducted. Samples of size $n = 5, 10, 30, 50$, and 100 are generated from the Burr type XII distribution with two values of the shape parameter, $\theta=1$ and 3 . Two values for the parameter of the exponential prior are chosen, $\delta=1$ and 1.8 . For the GSLF with $k=1$, values of the constants were chosen as $a_0=10$ and 100 , $a_1=50$, and with $k=2$, the same values of a_0 and a_1 in addition to $a_2=10$ and 100 were applied.

The number of replication taken was 5000 , and the mean of the estimated values for the parameter θ is obtained along with its mean square error (MSE) to compare the efficiency of the estimators, where

$$\mu(\hat{\theta}) = \frac{\sum_{i=1}^{5000} \hat{\theta}_i}{5000}, \quad \text{and} \quad MSE(\hat{\theta}) = \frac{\sum_{i=1}^{5000} (\hat{\theta}_i - \theta)^2}{5000}$$

The results of the simulation study are summarized and tabulated in tables 1- 8 for each estimator and for all sample sizes.

Table 1: Estimated mean of the shape parameter with $\theta=1$, $\delta=1$, $a_0=100$, $a_1=50$, $a_2=10$

n	Criteria	MinMSE	Jeffreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	0.7539	1.3623	1.4179	1.2269	1.2603
	MSE	0.2629	0.8730	1.0810	0.3060	0.3558
10	Mean	0.8864	1.1494	1.1662	1.1218	1.1358
	MSE	0.1090	0.1922	0.2102	0.1365	0.1477
30	Mean	0.9648	1.0456	1.0450	1.0428	1.0468
	MSE	0.0343	0.0416	0.0427	0.0383	0.0393
50	Mean	0.9770	1.0247	1.0269	1.0237	1.0259
	MSE	0.0220	0.0224	0.0228	0.0214	0.0218
100	Mean	0.9891	1.0127	1.0114	1.0124	1.0135
	MSE	0.0104	0.0110	0.0111	0.0108	0.0109

Table 2: Estimated mean of the shape parameter with $\theta=1$, $\delta=1.8$, $a_0=100$, $a_1=50$, $a_2=10$

n	Criteria	MinMSE	Jeffreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	0.7539	1.3624	1.4180	1.0385	1.0607
	MSE	0.2629	0.8730	1.0810	0.1273	0.1425
10	Mean	0.8865	1.1494	1.1662	1.0314	1.0429
	MSE	0.1090	0.1912	0.2110	0.8625	0.0918
30	Mean	0.9648	1.0456	1.0497	1.0475	1.0184
	MSE	0.0343	0.0416	0.0427	0.0328	0.0335
50	Mean	0.9770	1.0246	1.0269	1.0072	1.0094
	MSE	0.0212	0.0243	0.0279	0.0196	0.0198
100	Mean	0.9891	1.0122	1.0137	1.0042	1.0054
	MSE	0.0104	0.0110	0.0111	0.0103	0.0104

Table 3: Estimated mean of the shape parameter with $\theta=3$, $\delta=1$, $a_0=100$, $a_1=50$, $a_2=10$

n	Criteria	MinMSE	Jeffreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	2.2617	4.2757	4.6770	2.6525	2.7976
	MSE	2.3663	8.7246	12.245	0.6685	0.7038
10	Mean	2.659	3.5349	3.3684	2.8289	2.9193
	MSE	0.9814	1.8862	2.3135	0.5347	0.5735
30	Mean	2.8940	3.1641	3.2050	2.9509	2.9650
	MSE	0.3093	0.3889	0.4213	0.2566	0.2657
50	Mean	2.9309	3.0909	3.1136	2.9659	2.9873
	MSE	0.1818	0.2067	0.2170	0.1627	0.1661
100	Mean	2.9672	3.0460	3.0574	2.9849	2.9958
	MSE	0.0940	0.1005	0.1030	0.0890	0.0900

Table 4: Estimated mean of the shape parameter with $\theta=3$, $\delta=1.8$, $a_0=100$, $a_1=50$, $a_2=10$

n	Criteria	MinMSE	Jeffreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	2.2617	4.2758	4.6777	1.9550	2.0346
	MSE	2.3663	8.7247	12.2452	1.2586	1.1264
10	Mean	2.6594	3.5349	3.6844	2.3377	2.4060
	MSE	0.9814	1.8862	2.3135	0.6731	0.6191
30	Mean	2.8942	3.1641	3.2050	2.7391	2.7694
	MSE	0.3093	0.3889	0.4214	0.2559	0.2490
50	Mean	2.9310	3.090	3.1136	2.8327	2.8520
	MSE	0.1817	0.2067	0.2171	0.1622	0.1599
100	Mean	2.9672	3.0460	3.0574	2.9155	2.9260
	MSE	0.0936	0.1005	0.0103	0.0880	0.0870

Table 5: Estimated mean of the shape parameter with $\theta=1$, $\delta=1$, $a_0=10$, $a_1=50$, $a_2=100$

n	Criteria	MinMSE	Jefreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	0.7539	1.4740	1.6810	1.3169	1.4709
	MSE	0.2629	1.0310	1.5620	0.3810	0.5882
10	Mean	0.8865	1.2022	1.2886	1.1693	1.1246
	MSE	0.1090	0.2220	0.2983	0.1587	0.2118
30	Mean	0.9477	1.0626	1.0882	1.0593	1.0840
	MSE	0.0304	0.0414	0.0501	0.0409	0.0467
50	Mean	0.9769	1.0347	1.0498	1.0336	1.0483
	MSE	0.0220	0.0233	0.0255	0.0223	0.0243
100	Mean	0.9891	1.0177	1.0251	1.0174	1.0247
	MSE	0.0104	0.0113	0.0118	0.0110	0.0154

Table 6: Estimated mean of the shape parameter with $\theta=1$, $\delta=1.8$, $a_0=10$, $a_1=50$, $a_2=100$

n	Criteria	MinMSE	Jefreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	0.7539	1.4740	1.6814	1.1186	1.2445
	MSE	0.2629	1.0310	1.5603	0.1554	0.2437
10	Mean	0.8864	1.2022	1.2886	1.0761	1.1448
	MSE	0.1090	0.2220	0.2983	0.0969	0.1275
30	Mean	0.9648	1.0626	1.0882	1.0308	1.0547
	MSE	0.0344	0.0445	0.0519	0.0344	0.0385
50	Mean	0.9787	1.6387	1.6497	1.0170	1.0314
	MSE	0.0202	0.0234	0.0255	0.0200	0.0216
100	Mean	0.9890	1.0176	1.0250	1.0090	1.0160
	MSE	0.0104	0.0113	0.0179	0.0105	0.0109

Table 7: Estimated mean of the shape parameter with $\theta=3$, $\delta=1$, $a_0=10$, $a_1=50$, $a_2=100$

n	Criteria	MinMSE	Jefreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	2.2617	4.4859	5.1951	2.8014	3.1670
	MSE	2.3663	9.4887	14.731	0.6186	0.7846
10	Mean	2.5940	3.6380	3.9448	2.9148	3.1355
	MSE	0.9813	2.0423	2.8388	0.5273	0.6267
30	Mean	2.8940	3.1981	3.2921	2.9828	3.0672
	MSE	0.3093	0.4046	0.4746	0.2571	0.2776
50	Mean	2.9369	3.1103	3.1657	2.9855	3.0374
	MSE	0.1818	0.2120	0.2355	0.1628	0.1703
100	Mean	2.9672	3.0562	3.0855	2.9948	3.0210
	MSE	0.0939	0.1018	0.1075	0.0891	0.9127

Table 8: Estimated mean of the shape parameter with $\theta=3$, $\delta=1.8$, $a_0=10$, $a_1=50$, $a_2=100$

n	Criteria	MinMSE	Jefreys prior		Exponential prior	
			k=1	k=2	k=1	k=2
5	Mean	2.2617	4.4486	5.5915	2.0819	2.3446
	MSE	2.3663	9.4887	14.7312	1.0224	0.6646
10	Mean	2.6590	3.6379	3.9948	2.4159	2.5952
	MSE	0.9813	2.0423	2.8388	0.5842	0.4489
30	Mean	2.8940	3.1981	3.2921	2.7699	2.8476
	MSE	0.3093	0.4046	0.4746	0.2429	0.2252
50	Mean	2.9309	3.1103	3.1657	2.8519	2.9012
	MSE	0.1817	0.2120	0.2351	0.1571	0.1501
100	Mean	2.9672	3.0561	3.0835	2.9253	2.9512
	MSE	0.0940	0.1019	0.1075	0.0866	0.0850

Discussion

From simulation results it is noted that Bayes estimators based on exponential prior information with GSLF were generally better than Jeffreys prior information especially when the exponential parameter δ is large. On the other hand the MSE of the estimates of the parameter based on the classical MinMSE estimator performs better exceptionally when θ is small.

Finally, comparison with respect to GSLF showed that results are better when $k=1$ and a_0 is large enough.

References:

1. AL-Saiari, A.Y., Baharith, L.A., and Mousa S. A., (2014): Marshal Olkin extended Burr type XII distribution. *International Journal of Statistics and probability*, 3(1):78-84
2. Jang, D. H., et al (2014): Baysian estimation of Burr type XII distribution based on general progressive type II censoring. *Applied mathematical Sciences*, 8(69): 3435-3448.
3. Makhdoom, I. and Jafari, A. (2011): Bayesian estimations on the Burr type XII distribution using grouped and ungrouped data. *Australian Journal of Basic and Applied Sciences*, 5(6), 1525–1531.
4. Panahi, H. and Asadi, S. (2011): Analysis of the type-II hybrid censored Burr type XII distribution under linex loss function. *Applied Mathematical Sciences*, 5(79):3929–3942.
5. Rasheed, H. A. and Al Gazi, N.A.(2014): Bayes estimators for the shape parameter of Pareto type I distribution under generalized square error loss function. *Mathematical Theory and Modeling*, 4(11): 20-32.
6. Sindhu, T. N. and Aslam, M. (2013): Estimation of the Burr type VIII distribution through Baysian framework, *Advances in Arts, Social Sciences and Education Research*, 3(2): 393-402.
7. Soliman, A. A. (2002): Reliability estimation in a generalized life-model with application to the Burr – XII model, *IEEE Transactions on reliability*, 51(3): 337-343.
8. Soliman A.A., Abd Ellah A.H. and Abou-Alheggag N.A. et al (2011): Baysian inference and prediction of Burr type XII distribution for progressive first failure censored sampling. *Intelligent Information Management*, 3:175-185.
9. Tahir, M. and Saleem, M.(2010): On relationship between some Bayesian and classical estimators, *Pakistan Journal of Life and Social Science*, 8 (2): 159-161.
10. Tahir, M. and Saleem, M. (2011): A comparison of priors for the parameter of the time-to-failure model. *Pakistan Journal of science*, 63(1): 49-52.
11. Wu, S.J., Chen, Y. J. and Chang, C. T. (2007): Statistical inference based on progressively censored samples with random removals from the Burr type XII distribution. *Journal of Statistical Computation and Simulation*, (7): 19-27.
12. Yarmohammadi, M. and Pazira, H. (2010): Minimax estimation of the parameter of the Burr type XII distribution. *Australian Journal of Basic and Applied Sciences*, 4(12), 6611–6622.