

RESEARCH ARTICLE

Bayesian analysis of doubly censored lifetime data using two-component mixture of Weibull distribution

Navid Feroze^{1*} and Muhammad Aslam²

¹ Department of Mathematics and Statistics, Allama Iqbal Open University, Islamabad, Pakistan.

² Department of Statistics, Quaid-i-Azam University, Islamabad, Pakistan.

Revised: 27 March 2014; Accepted: 18 July 2014

Abstract: In recent years analysis of the mixture models under Bayesian framework has received considerable attention. However, the Bayesian estimation of the mixture models under doubly censored samples has not yet been reported. This paper proposes a Bayesian estimation procedure for analyzing lifetime data under doubly censored sampling when the failure times belong to a two-component mixture of the Weibull model. An extended version of the likelihood function for doubly censored samples for the analysis of a mixture of lifetime models has been introduced. The posterior estimation has been considered under the assumption of gamma prior using a couple of loss functions. The performance of the different estimators has been investigated and compared through the analysis of simulated data. A real-life example has been included to demonstrate the practical applicability of the results. The results indicated the preference of the estimates under squared logarithmic loss function (SLLF) for the estimation of the mixture model. The proposed method can be extended for more than two component mixtures.

Keywords: Credible intervals, loss functions, mixture models, posterior distributions, Weibull distributions.

INTRODUCTION

As a result of the adaptability in fitting time-to-failure of a very widespread multiplicity to multifaceted mechanisms, the Weibull distribution has assumed centre stage especially in the field of life-testing and reliability in survival analysis. It has shown to be very useful for modelling and analyzing lifetime data in medical, biological and engineering sciences (Lawless, 1982). Much of the attractiveness of the Weibull distribution is due to a variety of shapes based on its parameters. Nevertheless, the exponential, gamma and Weibull distributions are widely used in life-testing and reliability

experiments under different scenarios. However, the exponential distribution is suitable only when the hazard rate is constant. When the hazard rate changes with time, the gamma and Weibull distributions are the most suitable due to the flexibility of the scale and shape parameters of these distributions. The Weibull distribution has an extra edge over the gamma distribution as its distribution function and hazard function can be expressed in closed forms, which is not possible for the gamma distribution when the shape parameter is not an integer (Danish & Aslam, 2012).

Finite mixture distributions consist of a weighted sum of standard distributions and are a useful tool for reliability analysis of a heterogeneous population. They provide the necessary flexibility to model failure distributions of components with multiple failure modes. Furthermore, it is the logical way of modelling a probability distribution of a population with distinct subpopulations. However, the additional modelling capability comes at a cost of additional parameters and analytic difficulties.

Some of the recent developments on the estimation of the Weibull distribution - for the applications of this distribution in different situations - have been discussed by many Authors (Kundu & Raqab, 2009; Upadhyay & Gupta, 2010; Vasile *et al.*, 2010; Syuan-Rong & Shuo-Jye, 2011; Danish & Aslam, 2012; Singh *et al.*, 2013; Teimouri & Gupta, 2013; Yu & Peng, 2013). Several studies have been published on the classical analysis of the mixture of lifetime distributions under complete and censored samples. Some recent contributions include the following: Sultan *et al.* (2007), Nair and Abdul (2010), Afify (2011) and Erisoglu *et al.* (2011). In addition, the analysis of mixture models under Bayesian framework has

* Corresponding author (navidferoz@hotmail.com)

developed a significant interest among the statisticians. The authors dealing with Bayesian analysis of mixture models (Saleem & Aslam, 2008a; b; Saleem *et al.*, 2010; Saleem & Irfan, 2010; Majeed & Aslam, 2012) have restricted their discussions to the Bayes point estimation of the parameters under type I censored data. The real - life situations may demand the Bayesian analysis of the mixture models under some other censoring techniques such as doubly censored samples, which has not been considered for mixture models up till now.

This paper aims to consider the Bayesian analysis of two-component mixture of the Weibull distribution under doubly censored samples. An extended version of the likelihood function under doubly censored samples has been introduced for two-component mixture of lifetime distributions. The Bayes estimators have been derived and evaluated under the assumption of two loss functions using gamma prior.

METHODS AND MATERIALS

The model and likelihood function

The probability density function (pdf) of the Weibull distribution, defined by Weibull (1951) is:

$$f_j(x_{ji}) = \alpha_j \theta_j x_{ji}^{\alpha_j-1} e^{-\theta_j x_{ji}^{\alpha_j}} \quad x_{ji} > 0, \theta_j > 0, j=1,2, i=1,2,\dots,n \quad \dots(1)$$

The cumulative distribution function (CDF) of the distribution is:

$$F_j(x_{ji}) = 1 - e^{-\theta_j x_{ji}^{\alpha_j}} \quad x_{ji} > 0, \theta_j > 0, j=1,2, i=1,2,\dots,n \quad \dots(2)$$

A density function for the mixture of two component densities with mixing weights $(\pi, 1-\pi)$ is:

$$f(x) = \pi f_1(x) + (1-\pi) f_2(x) \quad 0 < \pi < 1 \quad \dots(3)$$

The cumulative distribution function for the mixture model is:

$$F(x) = \pi F_1(x) + (1-\pi) F_2(x) \quad \dots(4)$$

Consider a random sample size of 'n' from Weibull distribution, and let $x_{r_1}, x_{r_1+1}, \dots, x_s$ be the ordered observations. The remaining 'r-1' smallest observations and the 'n-s' largest observations have been assumed to be censored. Now based on the causes of failure,

the failed items are assumed to come either from subpopulation 1 or from subpopulation 2; so that the $x_{1r_1}, \dots, x_{1s_1}$ and $x_{2r_2}, \dots, x_{2s_2}$ failed items come from first and second subpopulations, respectively. The rest of the observations, which are less than x_r and greater than x_s have been assumed to be censored from each component. Where $x_s = \max(x_{1s_1}, x_{2s_2})$ and $x_r = \min(x_{1r_1}, x_{2r_2})$.

Therefore, $m_1 = s_1 - r_1 + 1$ and $m_2 = s_2 - r_2 + 1$ number of failed items can be observed from the first and the second subpopulation, respectively. The remaining $n - (s - r + 2)$ items are assumed to be censored observations and $s - r + 2$ are the uncensored items. Where $r = r_1 + r_2$, $s = s_1 + s_2$ and $m = m_1 + m_2$. Assuming the causes of the failure of the left censored items are identified, the likelihood function for the type II doubly censored sample $x = \{(x_{1r_1}, \dots, x_{1s_1}), (x_{2r_2}, \dots, x_{2s_2})\}$, can be written as:

$$L(\theta_1, \theta_2, \pi | x) \propto \pi^{s_1} (1-\pi)^{s_2} \left\{ F_1(x_{1r_1}) \right\}^{r_1-1} \left\{ F_1(x_{2r_2}) \right\}^{r_2-1} \left\{ 1 - F(x_s) \right\}^{n-s} \left\{ \prod_{i=r_1}^{s_1} f_1(x_{1,i}) \right\} \left\{ \prod_{i=r_2}^{s_2} f_2(x_{2,i}) \right\} \quad \dots(5)$$

Substituting the corresponding entries using equations (1) - (4)

$$L(\theta_1, \theta_2, \pi | x) \propto \pi^{s_1} (1-\pi)^{s_2} \left\{ 1 - e^{-\theta_1 x_{1r_1}^{\alpha_1}} \right\}^{r_1-1} \left\{ 1 - e^{-\theta_2 x_{2r_2}^{\alpha_2}} \right\}^{r_2-1} \left\{ \pi e^{-\theta_1 x_{1r_1}^{\alpha_1}} + (1-\pi) e^{-\theta_2 x_{2r_2}^{\alpha_2}} \right\}^{n-s} \left\{ \prod_{i=r_1}^{s_1} \alpha_1 \theta_1 x_{1i}^{\alpha_1-1} e^{-\theta_1 x_{1i}^{\alpha_1}} \right\} \left\{ \prod_{i=r_2}^{s_2} \alpha_2 \theta_2 x_{2i}^{\alpha_2-1} e^{-\theta_2 x_{2i}^{\alpha_2}} \right\} \quad \dots(6)$$

$$L(\theta_1, \theta_2, \pi | x) \propto \pi^{s_1} (1-\pi)^{s_2} \left\{ \sum_{u_1=0}^{r_1-1} (-1)^{u_1} \binom{r_1-1}{u_1} e^{-\theta_1 x_{1r_1}^{\alpha_1}} \right\} \left\{ \sum_{u_2=0}^{r_2-1} (-1)^{u_2} \binom{r_2-1}{u_2} e^{-\theta_2 x_{2r_2}^{\alpha_2}} \right\} \left\{ \sum_{u_3=0}^{n-s} \binom{n-s}{u_3} \pi^{u_3} e^{-\theta_1 x_{1r_1}^{\alpha_1}} (1-\pi)^{n-s-u_3} e^{-\theta_2 (n-s-u_3) x_{2r_2}^{\alpha_2}} \right\} \left\{ \theta_1^{m_1} e^{-\theta_1 \sum_{i=r_1}^{s_1} x_{1i}^{\alpha_1}} \right\} \left\{ \theta_2^{m_2} e^{-\theta_2 \sum_{i=r_2}^{s_2} x_{2i}^{\alpha_2}} \right\} \quad \dots(7)$$

$$L(\theta_1, \theta_2, \pi | x) \propto \sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} (-1)^{u_1} (-1)^{u_2} \binom{r_1-1}{u_1} \binom{r_2-1}{u_2} \binom{n-s}{u_3} \pi^{u_3+s_1} (1-\pi)^{n-s-u_3+s_2} \theta_1^{m_1} \theta_2^{m_2} e^{-\theta_1 \left(\sum_{i=r_1}^{s_1} x_{1i}^{\alpha_1} + u_1 x_{1r_1}^{\alpha_1} + u_3 x_{1s}^{\alpha_1} \right)} e^{-\theta_2 \left(\sum_{i=r_2}^{s_2} x_{2i}^{\alpha_2} + u_2 x_{2r_2}^{\alpha_2} + (n-s-u_3) x_{2s}^{\alpha_2} \right)}$$

Hence, the final version of the likelihood function can be written as:

$$L(\theta_1, \theta_2, \pi | x) \propto \sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \pi^{u_3+s_1} (1-\pi)^{n-s-u_3+s_2} \theta_w^{m_w} e^{-\theta_w \Omega(x_{wi})} \quad \dots(8)$$

where $\Omega(x_{1i}) = \sum_{l=1}^{s_1} x_{1i}^{\alpha_1} + u_1 x_{r_1}^{\alpha_1} + u_3 x_s^{\alpha_1}$ and

$$\Omega(x_{2i}) = \sum_{l=1}^{s_2} x_{2i}^{\alpha_2} + u_2 x_{r_2}^{\alpha_2} + (n-s-u_3) x_s^{\alpha_2}$$

Prior and posterior distributions

One of the most commonly used priors is the conjugate/gamma prior. The gamma prior for the mixture models can be defined as: let $\theta_1 \sim G(a_1, b_1)$ and $\theta_2 \sim G(a_2, b_2)$ are the gamma priors for each parameter and $\pi \sim B(a_3, b_3)$ is the beta prior for the mixing parameter π . Under the assumption of independence, these priors have been combined to produce a joint improved gamma prior for parameter as:

$$g(\theta_1, \theta_2, \pi) \propto \theta_1^{a_1-1} \theta_2^{a_2-1} e^{-(\theta_1 b_1 + \theta_2 b_2)} \pi^{a_3-1} (1-\pi)^{b_3-1}, \quad \theta_1, \theta_2 > 0, 0 < \pi < 1, a_1, a_2, a_3, b_1, b_2, b_3 > 0 \quad \dots(9)$$

where a_1, a_2, a_3, b_1, b_2 and b_3 are the parameters of prior distribution named as hyper-parameters.

The posterior distribution under gamma prior has been derived as:

$$h(\theta_1, \theta_2, \pi | x) = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \pi^{u_3+s_1} (1-\pi)^{n-s-u_3+s_2} \theta_w^{m_w} e^{-\theta_w \Omega(x_{wi})}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \int_0^\infty \int_0^\infty \int_0^1 \pi^{u_3+s_1} (1-\pi)^{n-s-u_3+s_2} \theta_w^{m_w} e^{-\theta_w \Omega(x_{wi})} d\pi d\theta_1 d\theta_2} \quad \dots(10)$$

$$h(\theta_1, \theta_2, \pi | x) = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \pi^{A_1-1} (1-\pi)^{A_2-1} \theta_w^{m_w+a_w-1} e^{-\theta_w \{\Omega(x_{wi})+b_w\}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(11)$$

where $A_1 = u_3 + s_1 + a_3$, $A_2 = n - u_3 + s_2 + b_3$ and $B(x, y)$ is a standard beta function.

Loss functions and Bayes estimators

The performance of different Bayes estimators can be compared in terms of posterior risks associated with each estimator. The posterior risk is defined to be the expected value of a loss function. We have assumed a couple of loss functions for posterior estimation and the description of these loss functions is as following:

Squared error loss function (SELF): The squared error loss function proposed by Legendre (1805) and Gauss (1810) is defined as: $L(\theta, \theta_{SELF}) = (\theta - \theta_{SELF})^2$. The Bayes estimator under this loss function is: $\theta_{SELF} = E(\theta)$.

Squared logarithmic loss function (SLLF): Another loss function Brown (1968) is called the squared based on logarithmic loss function. It can be defined as:

$L(\theta_{SLLF}, \theta) = (\log \theta_{SLLF} - \log \theta)^2$. The Bayes estimate under SLLF is: $\theta_{SLLF} = \exp \{E(\log \theta)\}$.

The expressions for Bayes estimators and posterior risks under these loss functions have been presented in the following. The Bayes estimators and posterior risks based on SELF are given as:

$$(B.E)_{SELF} = E(\theta_1) = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \int_0^1 \int_0^1 \pi^{A_1-1} (1-\pi)^{A_2-1} \theta_1^{m_1+a_1} e^{-\theta_1\{\Omega(x_{1i})+b_1\}} \theta_2^{m_2+a_2-1} e^{-\theta_2\{\Omega(x_{2i})+b_2\}} d\theta_1 d\theta_2 d\pi}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(12)$$

$$(B.E)_{SELF} = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_1 + a_1 + 1) \Gamma(m_2 + a_2)}{\{\Omega(x_{1i}) + b_1\}^{m_1+a_1+1} \{\Omega(x_{2i}) + b_2\}^{m_2+a_2}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(13)$$

Similarly, the Bayes estimators for other parameters can be derived and the generalized version of the Bayes estimator can be written as:

$$(B.E)_{SELF} = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1 + k_1, A_2) \Gamma(m_w + a_w + k_{1+w})}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w+k_{1+w}}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(14)$$

The generalized version of the posterior risk is as following:

$$\rho(B.E)_{SELF} = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1 + 2k_1, A_2) \Gamma(m_w + a_w + 2k_{1+w})}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w+2k_{1+w}}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} - \left[\frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1 + k_1, A_2) \Gamma(m_w + a_w + k_{1+w})}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w+k_{1+w}}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \right]^2 \quad \dots(15)$$

where $B(x, y)$ is the standard beta function; $(B, E)_{SELF}$ and $\rho(B, E)_{SELF}$ are the Bayes estimator and posterior risk under SELF respectively, and $\zeta(x_{wir})$ has been defined in (6). 'K' is a helping constant used to derive the generalized Bayes estimates.

Now, the Bayes estimators and posterior risks for θ_1 , θ_2 and π under SELF can be derived by putting $k_1 = 0$, $k_2 = 1$, $k_3 = 0$, $k_1 = 0$, $k_2 = 0$, $k_3 = 1$ and $k_1 = 1$, $k_2 = 0$, $k_3 = 0$, respectively in the above equations.

Finally, the Bayes estimators and posterior risks based on SLLF are given as:

$$E(\log \theta_1) = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \int_0^1 \int_0^1 \pi^{A_1-1} (1-\pi)^{A_2-1} \log(\theta_1) \theta_1^{m_1+a_1-1} e^{-\theta_1\{\Omega(x_{1i})+b_1\}} \theta_2^{m_2+a_2-1} e^{-\theta_2\{\Omega(x_{2i})+b_2\}} d\theta_1 d\theta_2 d\pi}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(16)$$

$$E(\log \theta_1) = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_1 + a_1) \{\psi(m_1 + a_1) - \ln\{\Omega(x_{1i}) + b_1\}\} \Gamma(m_2 + a_2)}{\{\Omega(x_{1i}) + b_1\}^{m_1+a_1} \{\Omega(x_{2i}) + b_2\}^{m_2+a_2}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w + a_w)}{\{\Omega(x_{wi}) + b_w\}^{m_w+a_w}}} \quad \dots(17)$$

where
$$\int_0^{\infty} \log(\theta_1) \theta_1^{m_1+a_1-1} e^{-\theta_1\{\Omega(x_{1i})+b_1\}} d\theta_1 = \Gamma(m_1+a_1) \left\{ \psi(m_1+a_1) - \ln\{\Omega(x_{1i})+b_1\} \right\}$$

Similarly, the Bayes estimators for other parameters can be derived and the generalized version of the Bayes estimator can be written as:

$$(B.E)_{SLLF} = \text{Exp} \left[\frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \{W_1\}^{k_1} \Gamma(m_w+a_w) \{T_w\}^{k_{1+w}}}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w+a_w)}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}} \right] \quad \dots(18)$$

The generalized version of the posterior risk is as following:

$$\rho(B.E)_{SLLF} = \frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \{W_1^2 + W_2\}^{k_1} \{T_w^2 + \psi_1(m_w+a_w)\}^{k_{1+w}}}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2)}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}} - \left[\frac{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \{W_1\}^{k_1} \Gamma(m_w+a_w) \{T_w\}^{k_{1+w}}}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}}{\sum_{u_1=0}^{r_1-1} \sum_{u_2=0}^{r_2-1} \sum_{u_3=0}^{n-s} \prod_{w=1}^2 (-1)^{u_w} \binom{r_w-1}{u_w} \binom{n-s}{u_3} \frac{B(A_1, A_2) \Gamma(m_w+a_w)}{\{\Omega(x_{wi})+b_w\}^{m_w+a_w}}} \right] \quad \dots(19)$$

where $T_w = \psi(m_w+a_w) - \ln\{\Omega(x_{wi})+b_w\}$, $W_1 = \psi(A_1) \psi(A_1+A_2)$, $W_2 = \psi_1(A_1) - \psi_1(A_1+A_2)$; $B(x,y)$ is standard beta function; $\psi(x)$ is digamma function; $\psi_1(x)$ is trigamma function; $(B.E)_{SLLF}$ and $\rho(B.E)_{SLLF}$ are Bayes estimator and posterior risk under SLLF respectively, and $\Omega(x_{wi})$ has been defined in (8).

Prior elicitation

Elicitation is a method to formulate the prior beliefs regarding some quantities into a probabilistic model. Under Bayesian inference it can be regarded as a technique to specify the values of hyper-parameters in a prior distribution for one or more parameters of the sampling distribution. It is not easy to have an accurate elicitation, because in many real-life situations the experts are often not familiar with the concept of probabilities. Even when the expert is familiar with the probabilities and their concept, it is by no means straightforward to evaluate

exactly a probability value for an event exactly. In such cases, elicitation encourages the expert and the facilitator to consider the meaning of the parameters being elicited. This has two helpful consequences. First, it brings the analysis closer to the application by demanding attention to what is being modelled, and what is reasonable to believe about it. Second, it helps to make the posterior distributions, once calculated, into meaningful quantities. Some methods of prior elicitation are reported by Geisser (1980); Chaloner and Duncan (1983); Garthwaite and Dickey (1992); Chaloner *et al.* (1993); Aslam (2003) and Gelman *et al.* (2004).

We have used the method suggested by Aslam (2003) for the prior elicitation. This method uses the prior predictive probabilities for elicitation. It compares the prior predictive distribution with expert's assessment about this distribution and permits the choice of the hyper-parameters that make the assessment agree closely with the member of the family. The prior predictive

distribution can be defined as:

$$g(y) = \int_0^{\infty} \int_0^{\infty} \int_0^1 g(\theta_1, \theta_2, \pi) f(y|\theta_1, \theta_2, \pi) d\pi d\theta_1 d\theta_2 \quad \dots(20)$$

where, $g(\theta_1, \theta_2, \pi|x)$ is the prior distribution and $f(y|\theta_1, \theta_2, \pi)$ is the mixture density for future observation from the model (3). According to equation (20), the prior predictive distribution under gamma prior; presented in equation (9) is:

$$g(y) = \frac{b_1^{a_1} b_2^{a_2}}{\Gamma(a_1) \Gamma(a_2) B(a_3, b_3)} \left[\frac{\Gamma(a_1 + 1) \Gamma(a_2) B(a_3 + 1, b_3) \alpha_1 y^{a_1 - 1}}{\{b_1 + y^{a_1}\}^{a_1 + 1} b_2^{a_2}} + \frac{\Gamma(a_1) \Gamma(a_2 + 1) B(a_3, b_3 + 1) \alpha_2 y^{a_2 - 1}}{b_1^{a_1} \{b_2 + y^{a_2}\}^{a_2 + 1}} \right] \quad \dots(21)$$

As we have to elicit six hyper-parameters ($a_1, a_2, a_3, b_1, b_2, b_3$) we have to consider six integrals. The set of hyper-parameters with minimum values has been chosen to be the elicited values of the hyper-parameters. By considering the prior predictive distribution in equation (21), we have assumed the expert's probabilities to be 0.10 for each integral. We considered the following integrals:

$$\int_0^{10} g(y) dy = 0.10, \quad \int_{10}^{20} g(y) dy = 0.10, \quad \int_{20}^{30} g(y) dy = 0.10, \\ \int_{30}^{40} g(y) dy = 0.10, \quad \int_{40}^{50} g(y) dy = 0.10, \quad \text{and} \quad \int_{50}^{60} g(y) dy = 0.$$

For convenience in the numerical calculations we have taken $\alpha_1 = \alpha_2 = 1.5$. Now, in order to solve these integrals simultaneously, a programme has been framed in SAS package using the "PROC SYSLIN" command and the elicited values of the hyper-parameters have been found to be: $(a_1, b_1, a_2, b_2, a_3, b_3) = (0.872167, 0.376822, 0.746821, 0.487262, 0.028352, 0.037613)$. (Aslam, 2003). For the elicitation purpose the expert probability assumed to be so that the sum of probabilities for all the integrals is less than or equal to one. Generally these probabilities are considered the same for each integral Kazmi *et al.* (2012).

SIMULATION STUDY AND RESULTS

As the analytical comparisons among the performance of different estimators are not possible, a simulation study has been conducted to serve this purpose. The performance of various estimators has been investigated and compared with respect to different priors, loss functions, parametric values, mixing weights and sample sizes. The parametric space used is: $(\theta_1, \theta_2, \pi) \in \{(0.10, 0.13, 0.45), (10, 13, 0.45), (0.10, 13, 0.45), (10, 0.13, 0.45)\}$. The samples of sizes $n = 20, 50$ and 100 have been generated by the inverse transformation method from two components mixture

of Weibull distribution under 10000 replications. The probabilistic mixing has been used to generate the mixture data. A random number u from Uniform (0,1) has been generated for each observation. If $u \leq \pi$, the observation has been randomly taken from the first subpopulation and if $u > \pi$, the observation is selected from the second subpopulation. The observations in the respective samples have been assumed to be 20 % censored. The amounts of posterior risks associated with each Bayes estimate have been given in parenthesis in the tables. The simulated samples have been drawn by following the listed steps:

- Step 1: Draw samples of size 'n' from the mixture model
- Step 2: Generate a uniform random numbers u for each observation
- Step 3: If $u \leq \pi$, take the observation from the first subpopulation otherwise from the second subpopulation
- Step 4: Determine the test termination points on left and right, that is, determine the values of x_r and x_s
- Step 5: The observations, which are less than x_r and greater than x_s have been considered to be censored from each component
- Step 6: Use the remaining observations from each component for the analysis

We have used the elicited values of the hyper-parameters obtained in the subsection elicitation for the numerical estimation. The amount of posterior risks associated with each Bayes estimate have been given in the parenthesis in Table 1.

Table 1 contains the Bayes estimates and posterior risks under gamma prior using a couple of loss functions. The Bayes estimates tend to converge to the true parametric values by increasing the sample size. The parameters have been over estimated in majority of the cases with few exceptions. The degree of over estimation

Table 1: Bayes estimates and posterior risks under gamma prior

θ_1, θ_2, π	n	Using SELF			Using SLLF		
		$\hat{\theta}_1$	$\hat{\theta}_2$	$\hat{\pi}$	$\hat{\theta}_1$	$\hat{\theta}_2$	$\hat{\pi}$
0.10, 0.13, 0.45	20	0.11348	0.14144	0.46673	0.11268	0.14108	0.46482
		(0.09732)	(0.11914)	(0.10709)	(0.04689)	(0.04988)	(0.01378)
	50	0.10516	0.13444	0.45448	0.10442	0.13410	0.45263
		(0.07937)	(0.08774)	(0.07569)	(0.03251)	(0.03644)	(0.01006)
	100	0.10324	0.13079	0.45080	0.10251	0.13046	0.45067
		(0.05871)	(0.05761)	(0.05930)	(0.02590)	(0.02881)	(0.00782)
10, 13, 0.45	20	10.40110	13.61047	0.47139	10.29146	13.48383	0.46943
		(0.42651)	(0.42886)	(0.17894)	(0.19818)	(0.21323)	(0.02280)
	50	10.19427	13.45415	0.45867	10.12410	13.38009	0.45692
		(0.34654)	(0.31235)	(0.12594)	(0.13449)	(0.15144)	(0.01640)
	100	10.12308	13.11796	0.45493	10.03259	13.09231	0.45493
		(0.25218)	(0.20283)	(0.09807)	(0.10705)	(0.11957)	(0.01275)
0.10, 13, 0.45	20	0.11492	13.30797	0.46375	0.11434	14.08057	0.46053
		(0.09854)	(0.42668)	(0.10642)	(0.04745)	(0.21326)	(0.01371)
	50	0.10640	13.22346	0.45268	0.10564	13.38727	0.45048
		(0.08031)	(0.30993)	(0.07518)	(0.03287)	(0.15559)	(0.01000)
	100	0.10433	13.13292	0.45028	0.10359	13.05950	0.45058
		(0.05930)	(0.20106)	(0.05909)	(0.02618)	(0.12275)	(0.00779)
10, 0.13, 0.45	20	10.37538	0.14613	0.46965	10.27118	0.14606	0.46644
		(0.42435)	(0.12063)	(0.17820)	(0.19719)	(0.05047)	(0.02266)
	50	10.17931	0.14015	0.45411	10.11064	0.13978	0.45510
		(0.34385)	(0.08876)	(0.12546)	(0.13363)	(0.03685)	(0.01629)
	100	10.08450	0.13218	0.45342	10.02360	0.13183	0.45081
		(0.24998)	(0.05819)	(0.09797)	(0.10615)	(0.02911)	(0.01270)

is directly proportional to true parametric values. In addition, squared logarithmic loss function (SLLF) provides better convergence irrespective of choice of true parametric values.

The posterior risk is the well known criterion for the comparison of the performance of the different Bayes estimates. The amount of the posterior risk is directly proportional to true parametric values and is inversely proportional to the sample size. This suggests that the estimates are consistent. The estimates based on SLLF correspond to the least amounts of risks.

Real-life example

This section covers the analysis of a real-life dataset regarding the breaking strengths of 64 single carbon fibers of length 10 presented by Lawless (2003). The idea has been to determine whether the results and properties of the Bayes estimators explored by a simulation study, have the same behaviour under a real-life situation. We have used the Kolmogorov-Smirnov and chi square tests to see whether the data follow the Weibull distribution. These tests show that the data follow the Weibull distribution at 5 % level of significance with p values 0.38473 and

Table 2: Bayes estimates and posterior risks under real-life data with 20 % censoring using SELF and SLLF

π	Using SELF			Using SLLF		
	$\hat{\theta}_1$	$\hat{\theta}_2$	$\hat{\pi}$	$\hat{\theta}_1$	$\hat{\theta}_2$	$\hat{\pi}$
0.45	0.32116	0.33254	0.43600	0.33611	0.33719	0.44134
	(0.26384)	(0.21810)	(0.01624)	(0.01257)	(0.01208)	(0.00574)
0.60	0.31650	0.33660	0.53411	0.32145	0.34135	0.54064
	(0.23889)	(0.22800)	(0.01861)	(0.01139)	(0.01267)	(0.00663)

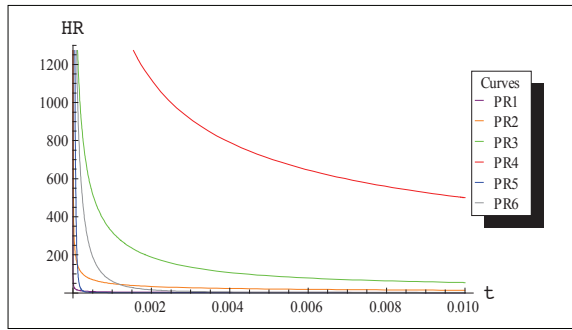


Figure 1: Graph of hazard rates for mixture of model using $\alpha_1 = \alpha_2 = 0.50$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

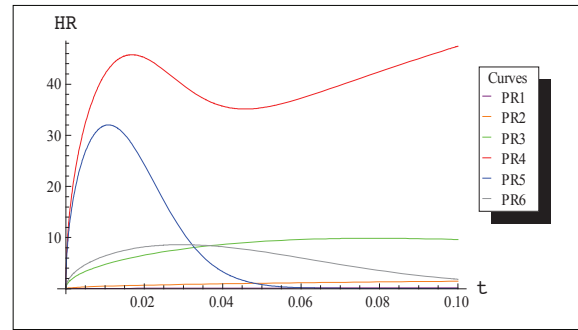


Figure 2: Graph of hazard rates for mixture of model using $\alpha_1 = \alpha_2 = 1.50$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

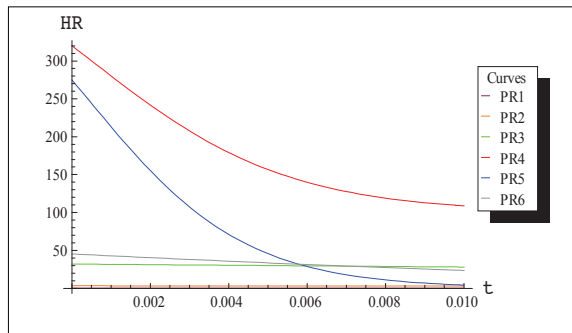


Figure 3: Graph of hazard rates for mixture of model using $\alpha_1 = \alpha_2 = 1.00$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

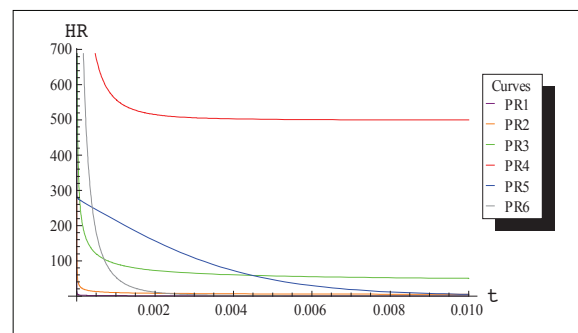


Figure 4: Graph of hazard rates for mixture of model using $\alpha_1 = 0.50$ and $\alpha_2 = 1.00$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

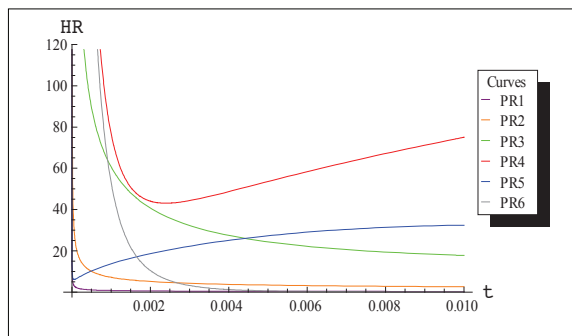


Figure 5: Graph of hazard rates for mixture of model using $\alpha_1 = 0.50$ and $\alpha_2 = 1.50$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

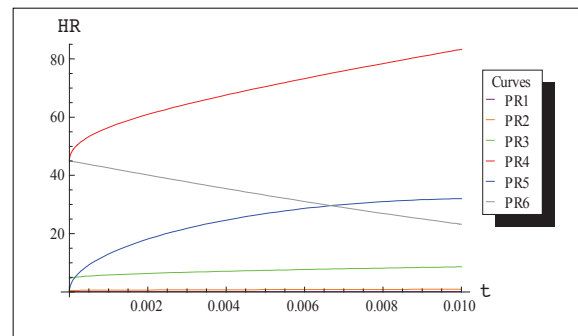


Figure 6: Graph of hazard rates for mixture of model using $\alpha_1 = 1.00$ and $\alpha_2 = 1.50$. HR: hazard rate; PR1: $\theta_1 = 0.1, \theta_2 = 0.5$; PR2: $\theta_1 = 1.0, \theta_2 = 5.0$; PR3: $\theta_1 = 10.0, \theta_2 = 50.0$; PR4: $\theta_1 = 100, \theta_2 = 500$; PR5: $\theta_1 = 0.1, \theta_2 = 500$; PR6: $\theta_1 = 100, \theta_2 = 0.5$

0.82786, respectively. We have taken $n = 64$, $\pi = 0.45$ and the data has been classified into two populations using probabilistic mixing (as discussed in the previous section), which produced: $r_1 = 4$, $r_2 = 3$, $s = 58$, $m_1 = 18$, $m_2 = 33$, so that the censoring rate is close to 20 % (that has been used in simulation study). The details of the censored mixture data are: Population I: 2.397, 2.522, 2.532, 2.614, 2.659, 2.740, 2.856, 2.917, 2.928, 2.937, 2.937, 3.139, 3.235, 3.377, 3.501, 3.537, 3.554, and 3.562. Population II: 2.396, 2.445, 2.454, 2.454, 2.474, 2.518, 2.525, 2.575, 2.616, 2.618, 2.624, 2.675, 2.738, 2.977, 2.996, 3.030, 3.125, 3.145, 3.220, 3.223, 3.243, 3.264, 3.272, 3.294, 3.332, 3.346, 3.408, 3.435, 3.493, 3.628, 3.852, 3.871 and 3.886. The results of the analysis are given in Table 2. The amounts of posterior risks associated with each estimate have been presented in parenthesis in the table.

The analysis under real-life data replicated the findings explored in the simulation study. The posterior risks under SLLF are the minimum for all the cases. The increase in the value of the mixing parameter (π) imposes a positive impact on the performance of the estimates for the first component of the mixture. This is simply due to the reason that the increase in the values of the mixing parameter will incorporate a larger proportion of the sample values for the analysis of the first component of the mixture. Hence, the results under real-life data gave us more confidence to suggest the use of SLLF for the estimation of the parameters of the mixed Weibull distribution under doubly censored samples.

Hazard rate for the mixture of Weibull distribution

The hazard rate is a useful way of describing the distribution of 'time to event' because it has a natural interpretation that relates to the ageing of a population. The hazard function is the risk of failure in a small time interval, given survival at the beginning of the time interval. As a function of time, a hazard function may be increasing, i.e. as time increases the rate for failure increases. For example, when a patient is untreated for a disease such as cancer or the medication do not work properly; or when a person is recovering from severe trauma like a surgery, The hazard function may be constant, meaning the rate of failure is the same regardless of how much time has passed. The constant hazard rate is mostly unrealistic. The hazard rate for the mixture of Weibull distribution has been compared under a range of parametric values.

The hazard rate function for the mixture of Weibull distribution is:

$$H(t) = \frac{\pi \alpha_1 \theta_1 t^{\alpha_1 - 1} e^{-\theta_1 t^{\alpha_1}} + (1 - \pi) \alpha_2 \theta_2 t^{\alpha_2 - 1} e^{-\theta_2 t^{\alpha_2}}}{1 - \left\{ \pi (1 - e^{-\theta_1 t^{\alpha_1}}) + (1 - \pi) (1 - e^{-\theta_2 t^{\alpha_2}}) \right\}} \quad \dots(15)$$

The graphs for the hazard rate of the mixture model for different parametric values and for the various ranges of the variable, are presented in the following.

The graphs suggest that the hazard rate for the mixture model is monotonically decreasing over time for $\alpha_1, \alpha_2 < 1$. If $\alpha_1 = \alpha_2 = 1$, the hazard rate is a constant function except for PR4 and PR5. Similarly for $\alpha_1, \alpha_2 > 1$, the hazard rate has different patterns for various combinations of the parameters. On the other hand, if we mix the above situations, that is, if we take $(\alpha_1, \alpha_2) = \{(0.50, 1.00), (0.50, 1.50), (1.00, 1.50)\}$, the behaviour of the hazard rate becomes different. Using $\alpha_1 = 0.50$ and $\alpha_2 = 1.00$, the hazard rate is decreasing but the tendency is different. In case where $\alpha_1 = 0.50$ and $\alpha_2 = 1.50$, the hazard rate is decreasing except under PR4 and PR5. For $\alpha_1 = 1.00$ and $\alpha_2 = 1.50$, the hazard function of the mixture density is decreasing, increasing and constant for different choices of the parametric values.

CONCLUSION

The Bayesian analysis of the Weibull distribution has been discussed by many authors under different censoring techniques. However, the Bayesian estimation of the mixture of Weibull distribution under doubly censored samples has not been reported in literature to date. This issue has been addressed in this paper. The paper proposes the Bayesian analysis of the doubly censored lifetime data using a two-component mixture of Weibull distributions under different loss functions using gamma prior. From the detailed analysis it can be concluded that the estimates under squared logarithmic loss function (SLLF) can be preferred for the estimation of the mixture model. The proposed estimators are consistent and capable of providing stable results from moderate to large samples. The findings of the study are useful for analysts from different fields dealing with the analysis of lifetime models, when causes of failures are more than one and the data is doubly censored. The study can further be extended for more than two component mixtures of the Weibull distribution.

REFERENCES

1. Afify W.M. (2011). Classical estimation of mixed Rayleigh distribution in type I progressive censored. *Journal of Statistical Theory and Applications* **10**(4): 619 – 632.

2. Aslam M. (2003). An application of prior predictive distribution to elicit the prior density. *Journal of Statistical Theory and Applications* **2**(1): 183 – 197.
3. Brown L.D. (1968). Inadmissibility of the usual estimators of scale parameters in problems with unknown location and scale parameters. *Annals of Mathematical Statistics* **39**: 29 – 48.
DOI: <http://dx.doi.org/10.1214/aoms/1177698503>
4. Chaloner K.M. & Duncan G.T. (1983). Assessment of a beta prior distribution: PM elicitation. *Statistician* **32**: 174 – 180.
DOI: <http://dx.doi.org/10.2307/2987609>
5. Chaloner K.M., Church T., Louis T.A. & Matts J.P. (1993). Graphical elicitation of a prior distribution for a clinical trial. *Statistician* **42**: 341 – 353.
DOI: <http://dx.doi.org/10.2307/2348469>
6. Danish M.Y. & Aslam M. (2012). Bayesian estimation in the proportional hazards model of random censorship under asymmetric loss functions. *Data Science Journal* **11**(6): 72 – 88.
DOI: <http://dx.doi.org/10.2481/dsj.012-004>
7. Erisoglu U., Erisoglu M. & Erol H. (2011). A mixture model of two different distributions approach to the analysis of heterogeneous survival data. *International Journal of Computational and Mathematical Sciences* **5**(2): 75.
8. Garthwaite P.H. & Dickey J.M. (1992). Elicitation of prior distributions for variable-selection problems in regression. *Annals of Statistics* **20**: 1697 – 1719.
DOI: <http://dx.doi.org/10.1214/aos/1176348886>
9. Gauss C.F. (1810). Least squares method for the combinations of observations (translated by J. Bertrand 1955). Mallet-Bachelier, Paris, France.
10. Geisser S. (1980). A Predictive Primer. *Bayesian Analysis in Econometrics and Statistics Essays in Honor of Harold Jeffreys* (ed. A. Zellner), pp. 363 – 381. Amsterdam: North-Holland.
11. Gelman A., Carlin J. B., Stern H. S. & Rubin D.B. (2004). *Bayesian Data Analysis*, 2nd edition. Chapman and Hall/CRC, New York, USA.
12. Kazmi S.M.A., Aslam M. & Ali S. (2012). On the Bayesian estimation for two component mixture of maxwell distribution, assuming type i censored data. *International Journal of Applied Science and Technology (IJAST)* **2**(1): 197 – 218.
13. Kundu D. & Raqab M.Z. (2009). Estimation of $R = P(Y < X)$ for three- parameter Weibull distribution. *Statistics and Probability Letters* **79**: 1839 – 1846.
DOI: <http://dx.doi.org/10.1016/j.spl.2009.05.026>
14. Lawless J.F. (1982). *Statistical Models and Methods for Lifetime Data*. John Wiley & Sons, New York, USA.
15. Lawless J.F. (2003). *Statistical Models Time Data*, 2nd edition. John Wiley & Sons, New York, USA.
16. Legendre A. (1805). *New Methods for the Determination of Orbits of Comets*, Courcier, Paris, France.
17. Majeed M.Y. & Aslam M. (2012). Bayesian analysis of the two component mixture of inverted exponential distribution under quadratic loss functions. *International Journal of Physical Sciences* **7**(9): 1424 – 1434.
18. Nair M.T. & Abdul E.S. (2010). Finite mixture of exponential model and its applications to renewal and reliability theory. *Journal of Statistical Theory and Practice* **4**(3): 367 – 373.
DOI: <http://dx.doi.org/10.1080/15598608.2010.10411992>
19. Saleem M. & Aslam M. (2008a). Bayesian analysis of the two component mixture of the Rayleigh distribution with the uniform and the Jeffreys priors. *Journal of Applied Statistical Science* **16**(4): 105 – 113.
20. Saleem M. & Aslam M. (2008b). On prior selection for the mixture of Rayleigh distribution using predictive intervals. *Pakistan Journal of Statistics* **24**(1): 21 – 35.
21. Saleem M., Aslam M. & Economou P. (2010). On the Bayesian analysis of the mixture of power function distribution using the complete and the censored sample. *Journal of Applied Statistics* **37**(1): 25 – 40.
DOI: <http://dx.doi.org/10.1080/02664760902914557>
22. Saleem M. & Irfan M. (2010). On properties of the Bayes estimates of the Rayleigh mixture parameters: a simulation study. *Pakistan Journal of Statistics* **26**(3): 547 – 555.
23. Singh S.K., Singh U. & Kumar D. (2013). Bayesian estimation of parameters of inverse Weibull distribution. *Journal of Applied Statistics* **40**(7): 1597 – 1607.
DOI: <http://dx.doi.org/10.1080/02664763.2013.789492>
24. Sultan K.S., Ismail M.A. & Al-Moisheer A.S. (2007). Mixture of two inverse Weibull distributions: properties and estimation. *Computational Statistics and Data Analysis* **51**: 5377 – 5387.
DOI: <http://dx.doi.org/10.1016/j.csda.2006.09.016>
25. Syuan-Rong H. & Shuo-Jye W. (2011). Bayesian estimation and prediction for Weibull model with progressive censoring. *Journal of Statistical Computation and Simulation* **81**: 1 – 14.
26. Teimouri M. & Gupta A.K. (2013). On the three-parameter Weibull distribution shape parameter estimation. *Journal of Data Science* **11**: 403 – 414.
27. Upadhyay S.K. & Gupta A. (2010). A Bayes analysis of modified Weibull distribution via Markov chain Monte Carlo simulation. *Journal of Statistical Computation and Simulation* **80**(3): 241 – 254.
DOI: <http://dx.doi.org/10.1080/00949650802600730>
28. Vasile P., Eugene P. & Alina C. (2010). Bayes estimators of modified-Weibull distribution parameters using Lindley's approximation. *WSEAS Transactions on Mathematics* **9**(7): 539 – 549.
29. Weibull W. (1951). A statistical distribution function of wide applicability. *Journal of Applied Mechanics* **18**(3): 293 – 297.
30. Yu H. & Peng C. (2013). Estimation for Weibull distribution with type II highly censored data. *Quality Technology and Quantitative Management* **10**(2): 193 – 202.