



The Bayesian Estimation for The Shape Parameter of The Power Function Distribution (PFD-I) to Use Hyper Prior Functions

Jinan Abbas Naser Al-obedy
Technical College of Management-
Baghdad, Baghdad ,IRAQ
Jinan.abbas@mtu.edu.iq

Received:6/12/2020

Accepted:24/1/2021

Published: March/ 2021



This work is licensed under a [Creative Commons Attribution-NonCommercial 4.0 International \(CC BY-NC 4.0\)](https://creativecommons.org/licenses/by-nc/4.0/)

Abstract

The objective of this study is to examine the properties of Bayes estimators of the shape parameter of the Power Function Distribution (PFD-I), by using two different prior distributions for the parameter θ and different loss functions that were compared with the Maximum likelihood estimators. In many practical applications, we may have two different prior information about the prior distribution for the shape parameter of the Power Function Distribution, which influences to the parameter estimation. So, we used two different kind of the conjugate priors of shape parameter θ of the Power Function Distribution (PFD-I) to estimate it. The conjugate prior function of the shape parameter θ was considered as combination of two different prior distributions such gamma distribution with Erlang distribution and Erlang distribution with exponential distribution and Erlang distribution with non-informative distribution and exponential distribution with non-informative distribution. We derived Bayes estimators for shape parameter θ of the Power Function Distribution (PFD-I) according to different loss functions such as the squared error loss function (SELF), the weighted error loss function (WSELF) and modified linear exponential (MLINEX) loss function (MLF), with two different double priors. In addition to the classical estimation (maximum likelihood estimation). We used simulation to get results of this study, for different cases of the shape parameter (θ) of the Power Function Distribution used to generate data for different samples sizes.

Paper type: Research paper.

Keywords: The power function distribution (PFD-I), MLE, Bayes Estimation, SELF, WSELF, MLINEX.

1. Introduction

The power function distribution (PFD-I) is a member of continuous probability distributions. The power function distribution used in a wide range of fields such as physics, earth science, economics, social science and the electrical component reliability [4]. The power function distribution used in the analysis of lifetime data and in problems related to the modeling of failure processes. Also, the power function distribution is a flexible life time distribution model that may offer a good fit to some sets of failure data. Theoretically, Power function distribution is a special case of Pareto distribution. The power function distribution is the best distribution to check the reliability of any electrical component. We mention some of studies in a brief manner:

Rahman et.al(2012) [5] estimated the shape parameter θ of the power function distribution (PFD-I) using Bayes estimation under different loss functions such as squared error loss function, quadratic loss function, modified linear exponential (MLINEX) loss function and non-linear exponential (NLINEX) loss function and along with maximum likelihood to identify the best estimation among all methods .They concluded that except for few cases Bayes estimator under NLINEX loss function and squared error loss function are better than other estimators in their study.

Kifayat et.al (2012) [3] used Bayesian analysis of The power function distribution under different priors, which are informative (gamma and Rayleigh) priors and non-informative (Jeffreys and uniform) priors. They derived the posterior distribution for the unknown parameter θ of the power distribution. Also they derived prior predictive distribution under informative priors, which is used for the elicitation of hyper parameters.

Zaka and Akhtar(2013) [10] used various methods to estimate the shape parameter θ of the power function distribution , such as the least squares method and relative least squares method and ridge regression method .They obtain the results by using simulation .They used total deviation (T.D) and mean square error (M.S.E) to identify the best estimation among all methods.

Sultan et.al (2014) [7] derived the posterior distribution of power function distribution under three double priors (gamma-exponential distribution, chi-square-exponential distribution, gamma-chi-square distribution) and three type of single priors. Also they developed posterior predictive distributions under double priors. From the empirical results they determine the best method of estimation according to the smallest value of the posterior standard error and AIC and BIC values. They observed that in most cases, Bayesian estimator under the double prior gamma distribution with exponential distribution has the less posterior standard error and less AIC and BIC values.

Hanif et.al (2015) [2] estimated the shape parameter θ of the power function distribution (PFD-I) using Bayes estimation assuming Weibull and Generalized Gamma distributions as priors for the unknown parameters .They derived posterior distribution for parameter θ under different priors. Then they derived Bayes estimator of the shape parameter θ under the squared error loss function In addition to the classical estimation (maximum likelihood estimation) to identify the best estimation among all used methods. From the empirical results, they concluded for small sample sizes the Bayes estimator with weibull prior performed better as compared to other estimators.

Ronak and Achyut (2016) [6] estimated the shape parameter θ of the power function distribution (PFD-I) using Bayes estimation assuming gamma and uniform priors, gamma and jeffrey's priors, gamma and priors and only gamma prior distributions as priors for the unknown parameters. They derived posterior distribution for parameter θ under different priors. Then they derived Bayes estimator of the shape parameter θ under the squared error loss function. In addition to the reliability at time t , and they constructed of equal tail credible interval for future observation by using simulation to compare the performance of the estimators under different double priors. According to the type-II censored sample from the power function distribution.

A few studies have examined the Bayes estimator of the parameter by considering a combination of two prior distribution, so we try in this study to use Bayes estimator for the shape parameter θ of the power function distribution by using the conjugate prior of the parameter θ is considered as combination of two prior distribution and by classical estimation (Maximum Likelihood Estimation).

So the aim of this study is examine the properties of Bayes estimators of the shape parameter of the Power Function Distribution (PFD-I), by using two different prior distributions for the parameter θ according to each of the posterior distributions for the parameter θ , and different loss functions, and compared these estimators with the Maximum likelihood estimators.

Bayes estimation make under different double prior selection for continuous case and under different loss functions. We have assumed gamma with Erlang distribution and Erlang with exponential distribution and Erlang with non-information distribution and exponential with non-information distribution as double priors. And we derive Bayes estimator of shape parameter θ of the Power Function Distribution under different loss function such as the Squared Errors Loss Function (SELF) and Weighted Squared Errors Loss Function (WSELF) and Modified Linear Exponential (MLINEX) Loss Function .

2. The Power Function Distribution (PFD-I)

Let us consider $t_1, t_2, \dots, t_n$ is a random sample of n independent observations from a Power Function Distribution (PFD-I) having the probability density function (pdf) with the shape parameter θ as "Eqn (1)" [2]:

$$f(t; \theta) = \theta t^{(\theta-1)}, \quad 0 < t < 1, \quad \theta > 0 \quad (1)$$

and the cumulative distribution function (cdf) is ;

$$(t; \theta) = t^\theta, \quad 0 \leq t \leq 1, \quad \theta > 0 \quad (2)$$

And the r^{th} moment about origin is

$$E(t^r) = \left(\frac{\theta}{\theta + r} \right) \quad (3)$$

Also, the mean and the variance are as follow

$$\text{Mean} = E(t) = \left(\frac{\theta}{\theta + 1} \right) \quad (4)$$

$$\text{Variance}(t) = \left(\frac{\theta}{(\theta + 2)(\theta + 1)^2} \right) \quad (5)$$

3. Estimation Methods

In this section, we used several methods to estimate of shape parameter θ of the Power Function Distribution (PFD-I), such classical estimation (Maximum Likelihood Estimation) and Bayes Estimation Methods as shown below.

3.1 Maximum Likelihood Estimation (MLE)

Here we obtain the MLE for the of θ based on the density as given in “Eqn (1)” and we can define the likelihood function as follows[1]:

$$L(\theta) = \prod_{i=1}^n f(t_i; \theta) = \prod_{i=1}^n \theta t_i^{(\theta-1)} = \theta^n \prod_{i=1}^n t_i^{(\theta-1)} \quad (6)$$

We can rewrite it as follow:

$$L(\theta) = \theta^n \exp \left[\theta \sum_{i=1}^n \ln(t_i) \right] \exp \left[- \sum_{i=1}^n \ln(t_i) \right] \quad (7)$$

By taking the log likelihood function on both side in “Eqn (7)” as follows

$$\log L(\theta) = n \log(\theta) + \theta \sum_{i=1}^n \ln(t_i) - \sum_{i=1}^n \ln(t_i) \quad (8)$$

The MLE for θ is

$$\hat{\theta}_{MLE} = - \frac{n}{\sum_{i=1}^n \ln(t_i)} \quad (9)$$

3.2 Bayes Estimation Method

We used different estimation methods to estimate of shape parameter θ of the Power Function Distribution (PFD-I) . By assuming t_i , $i = 1, 2, \dots, n$ are (iid) from the Power Function Distribution (PFD-I) as in “Eqn (1)” and likelihood function in “Eqn (7)” from the Power Function Distribution pdf given in “Eqn (1)” can be written as follows[1]:

$$L(\theta) = \theta^n \exp \left[\theta \sum_{i=1}^n \ln(t_i) \right] \exp \left[- \sum_{i=1}^n \ln(t_i) \right] \quad (7)$$

To derive the posterior distributions for the parameter θ , using a combination of two prior distributions such gamma distribution [8] with Erlang distribution [7] and Erlang distribution with exponential distribution [1], and Erlang distribution with non-informative distribution and exponential distribution with non-informative distribution.

3.2.1 The Conjugate Priors and posterior distributions

To derive the posterior distributions for the parameter θ , we need to determine the prior distributions for θ with pdf, as given below: By assuming a gamma prior for θ having pdf [9]

$$h_1(\theta) = \frac{b^a}{\Gamma a} \theta^{a-1} \exp(-b\theta) \quad \text{with } \theta, a, b \geq 0 \quad (10)$$

And a Erlang prior for θ having pdf [7],

$$h_2(\theta) = \lambda^2 \theta \exp(-\lambda\theta) \quad \text{with } \theta, \lambda \geq 0 \quad (11)$$

And a exponential prior for θ having pdf [1],

$$h_3(\theta) = \lambda_1 \exp(-\lambda_1 \theta) \quad \text{ith } \theta, \lambda_1 \geq 0 \quad (12)$$

And a non-informative prior for θ having pdf,

$$h_4(\theta) \propto \frac{1}{\theta^{c_1}} \quad \text{with } \theta, c_1 > 0 \quad (13)$$

Here we define their conjugate prior of the parameter θ by combining two priors as follows:

- If $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$, it means $P_1(\theta) \propto h_1(\theta)h_2(\theta)$ then we have

$$P_1(\theta) \propto \left[\frac{b^a}{\Gamma a} \lambda^2 \right] \theta^a \exp(-\theta(b + \lambda)) \quad \text{for } \theta \geq 0, a, b, \lambda > 0 \quad (14)$$

- If $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$, it means $P_2(\theta) \propto h_2(\theta)h_3(\theta)$ then we have

$$P_2(\theta) \propto [\lambda^2 \lambda_1] \theta \exp(-\theta(\lambda + \lambda_1)) \quad \text{for } \theta \geq 0, \lambda, \lambda_1 > 0 \quad (15)$$

- If $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$, it means $P_3(\theta) \propto h_2(\theta)h_4(\theta)$ then we have

$$P_3(\theta) \propto \lambda^2 \theta^{1-c_1} \exp(-\theta\lambda) \quad \text{for } \theta \geq 0, \lambda, c_1 > 0 \quad (16)$$

- If $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$, it means $P_4(\theta) \propto h_3(\theta)h_4(\theta)$ then we have

$$P_4(\theta) \propto \lambda_1 \theta^{-c_1} \exp(-\theta\lambda_1) \quad \text{for } \theta \geq 0, \lambda_1, c_1 > 0 \quad (17)$$

We obtain posterior distribution of the parameter θ for the given random sample t is given by [1]:

$$\pi(\theta \setminus t) = \frac{L(\theta) P(\theta)}{\int_{\theta} L(\theta) P(\theta) d\theta} \quad (18)$$

using “Eqn (7)” and for each $P_i(\theta)$, $i = 1, 2, 3, 4$ as shown above in “Eqn (14)” to “Eqn (15)” in “Eqn (18)”, after simplified steps, we get the posterior distributions for the parameter θ as follows:

- for $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$, we have the posterior distribution is

$$\pi_1(\theta \setminus t) = \frac{(b + \lambda - \sum_{i=1}^n \ln t_i)^{(a+n+1)}}{\Gamma(a+n+1)} \theta^{(a+n+1)-1} \exp(-\theta(b + \lambda - \sum_{i=1}^n \ln t_i))$$

with $\theta \geq 0, a, b, \lambda > 0$ (19)

$$\pi_1(\theta \setminus t) \sim \text{gamma dist}^n. (a_{(new)} = (a + n + 1), b_{(new)} = (b + \lambda - \sum_{i=1}^n \ln t_i))$$

- for $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$, we have the posterior distribution is

$$\pi_2(\theta \setminus t) = \frac{(\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i)^{(n+2)}}{\Gamma(n+2)} \theta^{(n+2)-1} \exp(-\theta(\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i))$$

for $\theta \geq 0$ and $\lambda, \lambda_1, n > 0$ (20)

$$\pi_2(\theta \setminus t) \sim \text{gamma dist}^n. (a_{(new)} = (n + 2), b_{(new)} = (\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i))$$

- for $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$, we have the posterior distribution

$$\pi_3(\theta \setminus t) = \frac{(\lambda - \sum_{i=1}^n \ln t_i)^{(n+2-c_1)}}{\Gamma(n+2-c_1)} \theta^{(n+2-c_1)-1} \exp(-\theta(\lambda - \sum_{i=1}^n \ln t_i))$$

$$, \theta \geq 0, \lambda, n, c_1 > 0 \quad (21)$$

$$\pi_3(\theta \setminus t) \sim \text{gamma dist}^n. (a_{(\text{new})} = (n+2-c_1), b_{(\text{new})} = (\lambda - \sum_{i=1}^n \ln t_i))$$

- for $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$, we have the posterior distribution is

$$\pi_4(\theta \setminus t) = \frac{(\lambda_1 - \sum_{i=1}^n \ln t_i)^{(n-c_1+1)}}{\Gamma(n-c_1+1)} \theta^{(n-c_1+1)-1} \exp(-\theta(\lambda_1 - \sum_{i=1}^n \ln t_i))$$

$$\theta \geq 0, n, \lambda_1, c_1 \quad (22)$$

$$\pi_4(\theta \setminus t) \sim \text{gamma dist}^n. (a_{(\text{new})} = (n-c_1+1), b_{(\text{new})} = (\lambda_1 - \sum_{i=1}^n \ln t_i))$$

3.2.2 Bayes' Estimators

Here we derive Bayes' estimators ($\hat{\theta}$) for the parameter θ according to different loss functions such as the squared error loss function (SELF), the weighted error loss function (WSELF) and modified linear exponential (MLINEX) loss function (MLF), with two different double priors as follows :

we derive Bayes' estimators ($\hat{\theta}$) for the parameter θ according the squared error loss function (SELF) , by minimize the posterior expected loss, as follows :

$$L_1(\hat{\theta}, \theta) = (\hat{\theta} - \theta)^2, \text{ the risk function is } R(\hat{\theta} - \theta) = \hat{\theta}^2 - 2\hat{\theta}E(\theta \setminus t) + E(\theta^2 \setminus t).$$

Let $\frac{\partial}{\partial \hat{\theta}} R(\hat{\theta} - \theta) = 0$, we get Bayes estimator of θ denoted by $\hat{\theta}_{\text{Bayes}}$ for the

above prior as follows

$$\hat{\theta}_{\text{SE}} = E(\theta \setminus t) = \int_0^\infty \theta \pi(\theta \setminus t) d\theta \quad (23)$$

So, we derive Bayes' estimators ($\hat{\theta}$) for θ according to the squared error loss function (SELF) with different conjugate prior of the parameter as the mean of the posterior distribution as follows:

- For $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$, we have

$$\hat{\theta}_{\text{SE1}} = \frac{(a+n+1)}{(b+\lambda - \sum_{i=1}^n \ln t_i)} \quad a, n, b, \lambda > 0 \quad (24)$$

- For $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$, we have

$$\hat{\theta}_{SE2} = \frac{(n+2)}{(\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i)} \quad n, \lambda, \lambda_1 > 0 \quad (25)$$

- For $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{SE3} = \frac{(n+2-c_1)}{(\lambda - \sum_{i=1}^n \ln t_i)} \quad n, c_1, \lambda > 0 \quad (26)$$

- For $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{SE4} = \frac{(n-c_1+1)}{(\lambda_1 - \sum_{i=1}^n \ln t_i)} \quad n, c_1, \lambda_1 > 0 \quad (27)$$

Also, we derive Bayes' estimators ($\hat{\theta}$) for θ according to the weighted error loss function (WSELF) by minimize the posterior expected loss, as follows

$$L_2(\hat{\theta}, \theta) = \frac{(\hat{\theta} - \theta)^2}{\theta}, \text{ the risk function is}$$

$$R_2(\frac{\hat{\theta} - \theta}{\theta}) = \hat{\theta}^2 E(\frac{1}{\theta} \setminus t) - 2\hat{\theta} + E(\theta \setminus t)$$

Let $\frac{\partial}{\partial \hat{\theta}} R_2(\frac{\hat{\theta} - \theta}{\theta}) = 0$, we get Bayes estimator of θ denoted by $\hat{\theta}_{WSE}$ for the

above prior as follows:

$$\hat{\theta}_{WSE} = \frac{1}{E(\frac{1}{\theta} \setminus t)} = \frac{1}{\int_0^\infty \frac{1}{\theta} \pi(\theta \setminus t) d\theta} \quad (28)$$

So, we derive Bayes' estimators ($\hat{\theta}$) for θ according to the weighted squared error loss function (WSELF) with different conjugate prior of the parameter as follows:

- For $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$, we have

$$\hat{\theta}_{WSE1} = \frac{(a+n)}{(b+\lambda - \sum_{i=1}^n \ln t_i)} \quad a, n, b, \lambda > 0 \quad (29)$$

- For $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$, we have

$$\hat{\theta}_{WSE2} = \frac{(n+1)}{(\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i)} \quad n, \lambda, \lambda_1 > 0 \quad (30)$$

- For $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{\text{WSE3}} = \frac{(n+1-c_1)}{(\lambda - \sum_{i=1}^n \ln t_i)} \quad n, c_1, \lambda > 0 \quad (31)$$

- For $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{\text{WSE4}} = \frac{(n-c_1)}{(\lambda_1 - \sum_{i=1}^n \ln t_i)} \quad n, c_1, \lambda_1 > 0 \quad (32)$$

And, we derive Bayes' estimators ($\hat{\theta}$) for the parameter θ according to the Modified linear exponential (MLINEX) loss function (MLF), by minimize the posterior expected loss, as follows [5]:

$$L_3(\hat{\theta}, \theta) = w \left[\left(\frac{\hat{\theta}}{\theta} \right)^c - c \log \left(\frac{\hat{\theta}}{\theta} \right) - 1 \right], \quad w > 0, c \neq 0, \text{ the risk function is}$$

$$R_3(\hat{\theta}, \theta) = \hat{\theta}^{-c} E(\theta^{-c} \setminus t) - c \log(\hat{\theta}) + c E(\log(\theta \setminus t)) - 1$$

$$\text{Let } \frac{\partial}{\partial \hat{\theta}} R_3(\hat{\theta}, \theta) = 0 \Rightarrow c \hat{\theta}^{-c-1} E(\theta^{-c} \setminus t) - \frac{c}{\hat{\theta}} = 0 \quad \text{Then we have the Bayes}$$

estimator of θ denoted by $\hat{\theta}_{\text{MLF}}$ for the above prior as follows

$$\hat{\theta}_{\text{MLF}} = [E(\theta^{-c} \setminus t)]^{-\frac{1}{c}} = \left[\int_0^\infty \theta^{-c} \pi(\theta \setminus t) d\theta \right]^{-\frac{1}{c}}, \quad c \neq 0 \quad (33)$$

So, we derive Bayes' estimators ($\hat{\theta}$) for θ according the MLINEX loss function (MLF) with different conjugate prior of the parameter as follows :

- For $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$, we have

$$\hat{\theta}_{\text{MLF1}} = \left[\frac{\Gamma(a+n-c+1)}{\Gamma(a+n+1)} \right]^{-\frac{1}{c}} (b + \lambda - \sum_{i=1}^n \ln t_i)^{-1}, \quad a, n, b, \lambda > 0, c \neq 0 \quad (34)$$

- For $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$, we have

$$\hat{\theta}_{\text{MLF2}} = \left[\frac{\Gamma(n-c+2)}{\Gamma(n+2)} \right]^{-\frac{1}{c}} (\lambda + \lambda_1 - \sum_{i=1}^n \ln t_i)^{-1}, \quad n, \lambda, \lambda_1 > 0, c \neq 0 \quad (35)$$

- For $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{\text{MLF3}} = \left[\frac{\Gamma(n+2-c_1-c)}{\Gamma(n+2-c_1)} \right]^{-\frac{1}{c}} (\lambda - \sum_{i=1}^n \ln t_i)^{-1}, \quad n, c_1, \lambda > 0, c \neq 0 \quad (36)$$

- For $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$, we have

$$\hat{\theta}_{\text{MLF4}} = \left[\frac{\Gamma(n-c_1-c+1)}{\Gamma(n-c_1+1)} \right]^{-\frac{1}{c}} (\lambda_1 - \sum_{i=1}^n \ln t_i)^{-1}, \quad n, c_1, \lambda_1 > 0, c \neq 0 \quad (37)$$

4. Simulation Study

For compare the quality of the Bayes estimators with the Maximum likelihood estimator for θ according to the assumed loss functions and the values for the parameters of each of the posterior distributions for the parameter θ , to determine the best estimation among all these estimators, according to smallest value our criterion which are listed below. We used empirical study. So we have considered several steps to perform simulation study by taking

1. The samples of sizes $n = 25, 50, 75$ and 100 which represented small, moderate and large sample size.
2. Different values of the shape parameters were chosen as $\theta = 0.5, 1, 1.5$ for the Power Function Distribution (PFD-I).
3. The generated random samples for all sizes u_i (from uniform (0,1), then $t_i = u_i^{(1/\theta)}$ is random sample from the Power Function Distribution (PFD-I).
4. assuming values of the parameters of each the posterior distributions of the parameter θ as different combinations to be compare listed below:

$\pi_1(\theta \setminus t)$			$\pi_2(\theta \setminus t)$		$\pi_3(\theta \setminus t)$		$\pi_4(\theta \setminus t)$	
a	b	λ	λ	λ_1	λ	c_1	λ_1	c_1
3	2	4	4	1	4	1	1	2
3	2	6	4	2	4	2	2	2
1	2	6	4	4	3	3	2	1
2	1	6	6	4	2	3	3	1

5. The experimental results were repeated ($r = 1000$) times for each sample size (n).

We obtained Bayes' estimators and the Maximum likelihood estimator (MLE) of shape parameter θ of the Power Function Distribution (PFD-I), by using MATLAB-R2018a program. Then computed the Mean Square Errors (MSE) and Mean Weighted Square Errors (MWSE) and Mean Modified Linear Exponential (MLINEX) to determine the best estimation among all used methods, according to smallest value our criterion which are

- The squared error loss function (SELF)

$$MSE(\hat{\theta}) = \frac{1}{r} \sum_{r=1}^{1000} (\hat{\theta}(r) - \theta)^2 \quad (38)$$

- The weighted error loss function (WSELF)

$$MWSE(\hat{\theta}) = \frac{1}{r} \sum_{r=1}^{1000} \frac{(\hat{\theta}(r) - \theta)^2}{\theta} \quad (39)$$

- Modified linear exponential (MLINEX) loss function (MLF) calculated with $w=1$ and two value for $c=1,2$.

$$MLINEX(\hat{\theta}) = \frac{1}{r} \sum_{r=1}^{1000} (w [(\frac{\hat{\theta}(r)}{\theta})^c - c \log(\frac{\hat{\theta}(r)}{\theta}) - 1]), \quad w > 0, c \neq 0 \quad (40)$$

The experimental results for this study were summarized and presented in tables, according to our aim to the study, to determine the best estimation for of

the parameter θ , according to smallest value our criterion under different loss functions, under different conjugate prior of the parameter for all the true value of $\theta = 0.5, 1, 1.5$ that are assumed.

So the experimental results under the squared error loss function (SELF) are listed in tables (4-1) to (4-3). And the experimental results under the weighted error loss function (WSELF) are listed in tables (4-4) to (4-7). Also, the experimental results under Modified linear exponential (MLINEX) loss function (MLF) calculated with $w=1$ and two value for $c=1, 2$ are listed in tables (4-1) to (4-3).

Table (4-1): Estimated value ($\hat{\theta}_{SE}$) and $MSE(\hat{\theta})$ of PFD-I, under the SELF, under different conjugate prior of the parameter for $\theta = 0.5$.

Method	parameter			Estimated value ($\hat{\theta}_{SE}$)				$MSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	0.52269	0.51209	0.5082	0.50692	0.01202	0.00531	0.00355	0.00259
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.53653	0.52049	0.51413	0.51149	0.01070	0.00512	0.00347	0.00257
	3	2	6	0.51681	0.51048	0.50744	0.50647	0.00832	0.00445	0.00316	0.00238
	1	2	6	0.48116	0.49157	0.49459	0.49673	0.00732	0.00409	0.00298	0.00226
	2	1	6	0.50832	0.50589	0.5043	0.50407	0.00815	0.00438	0.00313	0.00236
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.50926	0.50617	0.50443	0.50416	0.00887	0.00457	0.00322	0.00241
	-	4	2	0.49953	0.50121	0.50111	0.50166	0.00811	0.00436	0.00311	0.00235
	-	4	4	0.48116	0.49157	0.49459	0.49673	0.00732	0.00409	0.00298	0.00226
	-	6	4	0.46414	0.48230	0.48825	0.49189	0.00730	0.00403	0.00294	0.00223
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.50016	0.50141	0.50121	0.50172	0.00883	0.00454	0.00320	0.00240
	-	4	2	0.48092	0.49157	0.49461	0.49675	0.00853	0.00444	0.00314	0.00236
	-	3	3	0.47107	0.48662	0.49130	0.49427	0.00901	0.00455	0.00319	0.00238
	-	2	3	0.48087	0.49159	0.49463	0.49677	0.00926	0.00463	0.00323	0.00241
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.49109	0.49667	0.49800	0.49930	0.00978	0.00476	0.00330	0.00245
	-	2	2	0.48087	0.49159	0.49463	0.49677	0.00926	0.00463	0.00323	0.00241
	-	2	1	0.5009	0.50162	0.50131	0.50179	0.00965	0.00475	0.00329	0.00245
	-	3	1	0.4907	0.49655	0.49794	0.49926	0.00895	0.00456	0.00321	0.00239

Table (4-2): Estimated value ($\hat{\theta}_{SE}$) and $MSE(\hat{\theta})$ of PFD-I, under the SELF, under different conjugate prior of the parameter for $\theta = 1$.

Method	parameter			Estimated value ($\hat{\theta}_{SE}$)				$MSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.03580	1.01840	1.01550	1.01410	0.04866	0.02182	0.01446	0.01129
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.95577	0.97816	0.98831	0.99361	0.02667	0.01603	0.01161	0.00946
	3	2	6	0.89527	0.94344	0.96392	0.97481	0.02990	0.01663	0.01168	0.00935
	1	2	6	0.83353	0.90850	0.93952	0.95606	0.04412	0.02083	0.01351	0.01032
	2	1	6	0.89263	0.94270	0.96361	0.97466	0.03164	0.01720	0.01196	0.00953
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.92107	0.95960	0.97564	0.98400	0.03090	0.01718	0.01206	0.00968
	-	4	2	0.88986	0.94193	0.96329	0.97450	0.03356	0.01779	0.01225	0.00971
	-	4	4	0.83353	0.90850	0.93952	0.95606	0.04412	0.02083	0.01351	0.01032
	-	6	4	0.78404	0.87740	0.91690	0.93832	0.05943	0.02585	0.01584	0.01158
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.91926	0.95915	0.97548	0.98393	0.03302	0.01782	0.01237	0.00987
	-	4	2	0.88390	0.94035	0.96265	0.97419	0.03798	0.01909	0.01286	0.01009
	-	3	3	0.88070	0.93952	0.96232	0.97403	0.04056	0.01979	0.01318	0.01029
	-	2	3	0.91548	0.95824	0.97516	0.98381	0.03805	0.01922	0.01303	0.01027
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.95325	0.97773	0.98836	0.99379	0.03876	0.01947	0.01324	0.01047
	-	2	2	0.91548	0.95824	0.97516	0.98381	0.03805	0.01922	0.01303	0.01027
	-	2	1	0.95363	0.97780	0.98834	0.99375	0.03569	0.01869	0.01288	0.01025
	-	3	1	0.91740	0.95870	0.94932	0.98387	0.03539	0.01850	0.01402	0.01007

Table (4-3): Estimated value ($\hat{\theta}_{SE}$) and $MSE(\hat{\theta})$ of PFD-I, under the SELF, under different conjugate prior of the parameter for $\theta = 1.5$.

Method	parameter			Estimated value ($\hat{\theta}_{SE}$)				$MSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.55380	1.52760	1.52320	1.52120	0.10950	0.04910	0.03254	0.02540
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	1.30180	1.39050	1.42830	1.44850	0.07683	0.04015	0.02737	0.02154
	3	2	6	1.19260	1.32150	1.37790	1.40890	0.12085	0.05478	0.03414	0.02519
	1	2	6	1.11030	1.27260	1.34300	1.38180	0.17468	0.07299	0.04292	0.03021
	2	1	6	1.20180	1.33000	1.38490	1.41470	0.11810	0.05333	0.03340	0.02479
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	1.27040	1.37490	1.41800	1.44090	0.09213	0.04471	0.02946	0.02273
	-	4	2	1.21200	1.33900	1.39210	1.42070	0.11548	0.05203	0.03275	0.02447
	-	4	4	1.11030	1.27260	1.34300	1.38180	0.17468	0.07299	0.04292	0.03021
	-	6	4	1.02460	1.21250	1.2973	1.34510	0.24250	0.10016	0.05698	0.03857
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	1.28540	1.38570	1.42620	1.44740	0.09078	0.04428	0.02934	0.02276
	-	4	2	1.23590	1.35860	1.40740	1.43310	0.11108	0.05002	0.03184	0.02408

	-	3	3	1.25000	1.36920	1.41550	1.43960	0.10966	0.04939	0.03161	0.02403
	-	2	3	1.32110	1.4093	1.44350	1.46110	0.09126	0.04453	0.02967	0.02315
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	1.40100	1.45180	1.47260	1.48320	0.08539	0.04329	0.02945	0.02327
	-	2	2	1.32110	1.40930	1.44350	1.46110	0.09126	0.04453	0.02967	0.02315
	-	2	1	1.37610	1.43800	1.46300	1.47580	0.07963	0.04163	0.02857	0.02266
	-	3	1	1.30210	1.39720	1.43470	1.45410	0.09034	0.04420	0.02941	0.02289

Table (4-4): Estimated value ($\hat{\theta}_{WSE}$) and $MWSE(\hat{\theta})$ of PFD-I, under the WSELF, under different conjugate prior of the parameter for $\theta = 0.5$.

Method	parameter			Estimated value ($\hat{\theta}_{WSE}$)				$MWSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	0.52269	0.51209	0.5082	0.50692	0.02405	0.01062	0.00710	0.00519
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.51803	0.51085	0.50762	0.50658	0.01811	0.00929	0.00651	0.00487
	3	2	6	0.49898	0.50102	0.50102	0.50160	0.01499	0.00836	0.00606	0.00460
	1	2	6	0.46334	0.48212	0.48817	0.49186	0.01561	0.00838	0.00603	0.00455
	2	1	6	0.49017	0.49634	0.49783	0.49918	0.01522	0.00840	0.00607	0.00460
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.49040	0.49644	0.49788	0.49922	0.01648	0.00875	0.00624	0.00470
	-	4	2	0.48102	0.49157	0.49460	0.49674	0.01577	0.00852	0.00613	0.00462
	-	4	4	0.46334	0.48212	0.48817	0.49186	0.01561	0.00838	0.00603	0.00455
	-	6	4	0.44695	0.47303	0.48191	0.48707	0.01679	0.00862	0.00611	0.00458
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.48092	0.49157	0.49461	0.49675	0.01706	0.00888	0.00629	0.00472
	-	4	2	0.46168	0.48174	0.48802	0.49179	0.01799	0.00906	0.00636	0.00474
	-	3	3	0.45145	0.47668	0.48466	0.48927	0.01973	0.00948	0.00654	0.00483
	-	2	3	0.46083	0.48156	0.48794	0.49175	0.01941	0.00943	0.00653	0.00483
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.47063	0.48653	0.49127	0.49426	0.01955	0.00949	0.00657	0.00486
	-	2	2	0.46083	0.48156	0.48794	0.49175	0.01941	0.00943	0.00653	0.00483
	-	2	1	0.48087	0.49159	0.49463	0.49677	0.01853	0.00926	0.00647	0.00482
	-	3	1	0.47107	0.48662	0.49130	0.49427	0.01802	0.00910	0.00639	0.00476

Table (4-5): Estimated value ($\hat{\theta}_{WSE}$) and $MWSE(\hat{\theta})$ of PFD-I, under the WSELF, under different conjugate prior of the parameter for $\theta = 1$.

Method	parameter			Estimated value ($\hat{\theta}_{WSE}$)				$MWSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.03580	1.01840	1.01550	1.01410	0.04866	0.02182	0.01446	0.01129
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.92282	0.96004	0.97580	0.98406	0.02900	0.01658	0.01177	0.00949
	3	2	6	0.86440	0.92597	0.95172	0.96544	0.03603	0.01842	0.01244	0.00975
	1	2	6	0.80266	0.89103	0.92731	0.94669	0.05416	0.02386	0.01488	0.01107

	2	1	6	0.86075	0.92491	0.95125	0.96520	0.03809	0.01903	0.01274	0.00993
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.88695	0.94115	0.96297	0.97435	0.03566	0.01842	0.01255	0.00990
	-	4	2	0.85690	0.92382	0.95078	0.96495	0.04034	0.01968	0.01304	0.01011
	-	4	4	0.80266	0.89103	0.92731	0.94669	0.05416	0.02386	0.01488	0.01107
	-	6	4	0.75501	0.86052	0.90499	0.92912	0.07189	0.02986	0.01773	0.01265
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.88390	0.94035	0.96265	0.97419	0.03798	0.01909	0.01286	0.01009
	-	4	2	0.84855	0.92154	0.94981	0.96445	0.04552	0.02107	0.01368	0.01050
	-	3	3	0.84401	0.92035	0.94932	0.96420	0.04851	0.02182	0.01402	0.01070
	-	2	3	0.87734	0.93868	0.96199	0.97388	0.04343	0.02053	0.01352	0.01049
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.91353	0.95778	0.97501	0.98375	0.04107	0.01999	0.01338	0.01048
	-	2	2	0.87734	0.93868	0.96199	0.97388	0.04343	0.02053	0.01352	0.01049
	-	2	1	0.91548	0.95824	0.97516	0.98381	0.03805	0.01922	0.01303	0.01027
	-	3	1	0.88070	0.93952	0.96232	0.97403	0.04056	0.01979	0.01318	0.01029

Table (4-6): Estimated value ($\hat{\theta}_{WSE}$) and $MWSE(\hat{\theta})$ of PFD-I, under the WSELF, under different conjugate prior of the parameter for $\theta = 1.5$.

Method	parameter			Estimated value ($\hat{\theta}_{WSE}$)				$MWSE(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.55380	1.52760	1.52320	1.52120	0.07300	0.03273	0.02169	0.01693
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	1.25690	1.36480	1.41020	1.43460	0.06273	0.03028	0.01982	0.01520
	3	2	6	1.15140	1.29700	1.36050	1.39540	0.09736	0.04218	0.02548	0.01834
	1	2	6	1.06920	1.24810	1.32560	1.36830	0.13784	0.05594	0.03215	0.02219
	2	1	6	1.15890	1.30490	1.36710	1.40100	0.09566	0.04105	0.02486	0.01799
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	1.22330	1.34850	1.39960	1.42680	0.07539	0.03394	0.02149	0.01615
	-	4	2	1.16710	1.31330	1.37400	1.40670	0.09398	0.03999	0.02429	0.01768
	-	4	4	1.06920	1.24810	1.32560	1.36830	0.13784	0.05594	0.03215	0.02219
	-	6	4	0.98666	1.18920	1.28050	1.3319	0.18588	0.07563	0.04245	0.02836
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	1.23590	1.35860	1.40740	1.43310	0.07405	0.03334	0.02123	0.01605
	-	4	2	1.18650	1.33140	1.38860	1.41880	0.09092	0.03817	0.02337	0.01721
	-	3	3	1.19800	1.34130	1.39640	1.42510	0.08970	0.03745	0.02303	0.01705
	-	2	3	1.26600	1.38050	1.42400	1.44630	0.07277	0.03274	0.02103	0.01605
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	1.34260	1.42220	1.45270	1.46820	0.06279	0.03024	0.02011	0.01569
	-	2	2	1.26600	1.38050	1.42400	1.44630	0.07277	0.03274	0.02103	0.01605
	-	2	1	1.32110	1.40930	1.44350	1.46110	0.06084	0.02968	0.01978	0.01543
	-	3	1	1.25000	1.36920	1.41550	1.43960	0.07310	0.03293	0.02107	0.01602

Table (4-7): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the MLINEX, under different conjugate prior of the parameter for $\theta = 0.5$ and $w=1$ and $c=1$.

Method	parameter			$\hat{\theta}_{MLF}$ Estimated value (θ_{MLF})				$\hat{\theta}$ MLINEX(θ)			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	0.52269	0.51209	0.5082	0.50692	0.02126	0.00985	0.00672	0.00497
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.51803	0.51085	0.50762	0.50658	0.01652	0.00869	0.00619	0.00467
	3	2	6	0.49898	0.50102	0.50102	0.50160	0.01475	0.00813	0.00592	0.00450
	1	2	6	0.46334	0.48212	0.48817	0.49186	0.01758	0.00879	0.00620	0.00463
	2	1	6	0.49017	0.49634	0.49783	0.49918	0.01550	0.00832	0.00600	0.00454
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.49040	0.49644	0.49788	0.49922	0.01669	0.00865	0.00616	0.00463
	-	4	2	0.48102	0.49157	0.49460	0.49674	0.01663	0.00860	0.00613	0.00461
	-	4	4	0.46334	0.48212	0.48817	0.49186	0.01758	0.00879	0.00620	0.00463
	-	6	4	0.44695	0.47303	0.48191	0.48707	0.01980	0.00933	0.00643	0.00475
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.48092	0.49157	0.49461	0.49675	0.01791	0.00894	0.00630	0.00470
	-	4	2	0.46168	0.48174	0.48802	0.49179	0.02026	0.00948	0.00653	0.00482
	-	3	3	0.45145	0.47668	0.48466	0.48927	0.02291	0.01010	0.00680	0.00496
	-	2	3	0.46083	0.48156	0.48794	0.49175	0.02185	0.00986	0.00671	0.00491
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.47063	0.48653	0.49127	0.49426	0.02121	0.00973	0.00665	0.00489
	-	2	2	0.46083	0.48156	0.48794	0.49175	0.02185	0.00986	0.00671	0.00491
	-	2	1	0.48087	0.49159	0.49463	0.49677	0.01936	0.00931	0.00647	0.00480
	-	3	1	0.47107	0.48662	0.49130	0.49427	0.01961	0.00935	0.00648	0.00479

Table (4-8): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the MLINEX, under different conjugate prior of the parameter for $\theta=1$ and $w=1$ and $c=1$.

Method	parameter			$\hat{\theta}_{MLF}$				$MLINEX(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.03580	1.01840	1.01550	1.01410	0.02155	0.01028	0.00694	0.00541
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.92282	0.96004	0.97580	0.98406	0.01636	0.00882	0.00613	0.00484
	3	2	6	0.86440	0.92597	0.95172	0.96544	0.02173	0.01032	0.00675	0.00514
	1	2	6	0.80266	0.89103	0.92731	0.94669	0.03409	0.01385	0.00832	0.00601
	2	1	6	0.86075	0.92491	0.95125	0.96520	0.02307	0.01068	0.00691	0.00524
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.88695	0.94115	0.96297	0.97435	0.02108	0.01010	0.00668	0.00514
	-	4	2	0.85690	0.92382	0.95078	0.96495	0.02455	0.01106	0.00708	0.00534
	-	4	4	0.80266	0.89103	0.92731	0.94669	0.03409	0.01385	0.00832	0.00601
	-	6	4	0.75501	0.86052	0.90499	0.92912	0.04631	0.01768	0.01011	0.00700
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.88390	0.94035	0.96265	0.97419	0.02252	0.01047	0.00685	0.00524
	-	4	2	0.84855	0.92154	0.94981	0.96445	0.02799	0.01187	0.00744	0.00554
	-	3	3	0.84401	0.92035	0.94932	0.96420	0.02999	0.01231	0.00763	0.00565
	-	2	3	0.87734	0.93868	0.96199	0.97388	0.02591	0.01128	0.00721	0.00544
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.91353	0.95778	0.97501	0.98375	0.02318	0.01062	0.00696	0.00533
	-	2	2	0.87734	0.93868	0.96199	0.97388	0.02591	0.01128	0.00721	0.00544
	-	2	1	0.91548	0.95824	0.97516	0.98381	0.02150	0.01022	0.00678	0.00523
	-	3	1	0.88070	0.93952	0.96232	0.97403	0.02412	0.01087	0.00703	0.00534

Table (4-9): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the MLINEX, under different conjugate prior of the parameter for $\theta=1.5$ and $w=1$ and $c=1$.

Method	parameter			$\hat{\theta}_{MLF}$				$MLINEX(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.55380	1.52760	1.52320	1.52120	0.02155	0.01028	0.00694	0.00541
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	1.25690	1.36480	1.41020	1.43460	0.02566	0.01151	0.00728	0.00543
	3	2	6	1.15140	1.29700	1.36050	1.39540	0.04125	0.01656	0.00965	0.00676
	1	2	6	1.06920	1.24810	1.32560	1.36830	0.06053	0.02240	0.01237	0.00830
	2	1	6	1.15890	1.30490	1.36710	1.40100	0.04058	0.01609	0.00939	0.00660
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	1.22330	1.34850	1.39960	1.42680	0.03142	0.01305	0.00796	0.00581
	-	4	2	1.16710	1.31330	1.37400	1.40670	0.03992	0.01564	0.00915	0.00646
	-	4	4	1.06920	1.24810	1.32560	1.36830	0.06053	0.02240	0.01237	0.00830
	-	6	4	0.98666	1.18920	1.28050	1.33190	0.08449	0.03088	0.01659	0.01078
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							

	-	4	1	1.23590	1.35860	1.40740	1.43310	0.03082	0.01274	0.00782	0.00573
	-	4	2	1.18650	1.33140	1.38860	1.41880	0.03868	0.01483	0.00873	0.00623
	-	3	3	1.19800	1.34130	1.39640	1.42510	0.03816	0.01449	0.00856	0.00614
	-	2	3	1.26600	1.38050	1.42400	1.44630	0.02999	0.01231	0.00763	0.00565
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	1.34260	1.42220	1.45270	1.46820	0.02436	0.01091	0.00706	0.00538
	-	2	2	1.26600	1.38050	1.42400	1.44630	0.02999	0.01231	0.00763	0.00565
	-	2	1	1.32110	1.40930	1.44350	1.46110	0.02412	0.01087	0.00703	0.00534
	-	3	1	1.25000	1.36920	1.41550	1.43960	0.03032	0.01250	0.00770	0.00568

Table (4-10): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the $MLINEX$, under different conjugate prior of the parameter for $\theta = 0.5$ and $w=1$ and $c=2$.

Method	parameter			$\hat{\theta}_{MLF}$				$MLINEX(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	0.52269	0.51209	0.50820	0.50692	0.09065	0.04095	0.02766	0.02035
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.50869	0.50600	0.50435	0.50411	0.06606	0.03498	0.02493	0.01880
	3	2	6	0.48999	0.49627	0.49780	0.49916	0.05922	0.03280	0.02385	0.01812
	1	2	6	0.45435	0.47737	0.48495	0.48942	0.07158	0.03562	0.02503	0.01866
	2	1	6	0.48100	0.49155	0.49459	0.49673	0.06248	0.03360	0.02421	0.01829
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.48088	0.49155	0.49460	0.49674	0.06732	0.03495	0.02486	0.01866
	-	4	2	0.47168	0.48673	0.49134	0.49427	0.06728	0.03478	0.02474	0.01856
	-	4	4	0.45435	0.47737	0.48495	0.48942	0.07158	0.03562	0.02503	0.01866
	-	6	4	0.43827	0.46837	0.47873	0.48465	0.08081	0.03789	0.02598	0.01915
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.47120	0.48663	0.49131	0.49426	0.07246	0.03618	0.02541	0.01894
	-	4	2	0.45196	0.47680	0.48471	0.48930	0.08229	0.03843	0.02636	0.01941
	-	3	3	0.44152	0.47169	0.48133	0.48677	0.09299	0.04094	0.02746	0.02000
	-	2	3	0.45070	0.47651	0.48459	0.48924	0.08863	0.03996	0.02707	0.01980
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.46029	0.48144	0.48789	0.49173	0.08591	0.03939	0.02685	0.01971
	-	2	2	0.45070	0.47651	0.48459	0.48924	0.08863	0.03996	0.02707	0.01980
	-	2	1	0.47074	0.48655	0.49127	0.49425	0.07833	0.03767	0.02610	0.01933
	-	3	1	0.46116	0.48162	0.48797	0.49176	0.07948	0.03784	0.02614	0.01932

Table (4-11): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the $MLINEX$, under different conjugate prior of the parameter for $\theta=1$ and $w=1$ and $c=2$.

Method	parameter			$\hat{\theta}_{MLF}$				$MLINEX(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.03580	1.01840	1.01550	1.01410	0.09177	0.04239	0.02836	0.02211
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	0.90619	0.95094	0.96953	0.97927	0.06686	0.03560	0.02459	0.01944
	3	2	6	0.84882	0.91720	0.94560	0.96074	0.08855	0.04170	0.02712	0.02067
	1	2	6	0.78707	0.88225	0.92119	0.94199	0.13621	0.05556	0.03335	0.02408
	2	1	6	0.84466	0.91597	0.94506	0.96045	0.09386	0.04311	0.02776	0.02105
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	0.86973	0.93187	0.95662	0.96951	0.08591	0.04078	0.02682	0.02064
	-	4	2	0.84026	0.91471	0.94451	0.96016	0.09968	0.04459	0.02842	0.02144
	-	4	4	0.78707	0.88225	0.92119	0.94199	0.13621	0.05556	0.03335	0.02408
	-	6	4	0.74034	0.85205	0.89902	0.92451	0.18137	0.07030	0.04031	0.02799
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							
	-	4	1	0.86605	0.93089	0.95621	0.96931	0.09162	0.04225	0.02749	0.02104
	-	4	2	0.83068	0.91209	0.94337	0.95957	0.11312	0.04780	0.02983	0.02226
	-	3	3	0.82545	0.91071	0.94279	0.95926	0.12093	0.04954	0.03057	0.02269
	-	2	3	0.85806	0.92886	0.95537	0.96889	0.10502	0.04545	0.02890	0.02187
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	0.89345	0.94775	0.96830	0.97872	0.09414	0.04281	0.02790	0.02144
	-	2	2	0.85806	0.92886	0.95537	0.96889	0.10502	0.04545	0.02890	0.02187
	-	2	1	0.89621	0.94841	0.96855	0.97883	0.08740	0.04119	0.02718	0.02101
	-	3	1	0.86216	0.92989	0.9558	0.96910	0.09796	0.04381	0.02818	0.02145

Table (4-12): Estimated value ($\hat{\theta}_{MLF}$) and $MLINEX(\hat{\theta})$ of PFD-I, under the $MLINEX$, under different conjugate prior of the parameter for $\theta=1.5$ and $w=1$ and $c=2$.

Method	parameter			$\hat{\theta}_{MLF}$				$MLINEX(\hat{\theta})$			
				Sample Size(n)				Sample Size(n)			
				25	50	75	100	25	50	75	100
MLE	-	-	-	1.55380	1.52760	1.52320	1.52120	0.09177	0.04239	0.02836	0.02211
Bayes	a	b	λ	for $\pi_1(\theta \setminus t)$ when $P_1(\theta) \propto \text{gamma}(a, b) \times \text{erlang}(\lambda)$							
	3	2	4	1.23430	1.35180	1.40110	1.42760	0.10388	0.04641	0.02923	0.02179
	3	2	6	1.13070	1.28470	1.35170	1.38860	0.16234	0.06601	0.03855	0.02703
	1	2	6	1.04840	1.23580	1.31680	1.36150	0.23227	0.08814	0.04905	0.03305
	2	1	6	1.13720	1.29230	1.35820	1.39410	0.16012	0.06422	0.03754	0.02643
Bayes	-	λ	λ_1	for $\pi_2(\theta \setminus t)$ when $P_2(\theta) \propto \text{erlang}(\lambda) \times \text{exponential}(\lambda_1)$							
	-	4	1	1.19960	1.33520	1.39040	1.41970	0.12613	0.05244	0.03191	0.02330
	-	4	2	1.14450	1.30030	1.36500	1.39980	0.15792	0.06251	0.03659	0.02588
	-	4	4	1.04840	1.23580	1.31680	1.36150	0.23227	0.08814	0.04905	0.03305
	-	6	4	0.96750	1.17740	1.27200	1.32530	0.31523	0.11953	0.06505	0.04255
Bayes	-	λ	c_1	for $\pi_3(\theta \setminus t)$ when $P_3(\theta) \propto \text{erlang}(\lambda) \times \text{non-informative}(c_1)$							

	-	4	1	1.21090	1.34490	1.39800	1.42590	0.12394	0.05125	0.03134	0.02301
	-	4	2	1.16150	1.31770	1.37920	1.41160	0.15378	0.05940	0.03493	0.02497
	-	3	3	1.17160	1.32720	1.38680	1.41780	0.15201	0.05806	0.03425	0.02461
	-	2	3	1.23820	1.36610	1.41420	1.43890	0.12093	0.04954	0.03057	0.02269
Bayes	-	λ_1	c_1	for $\pi_4(\theta \setminus t)$ when $P_4(\theta) \propto \text{exponential}(\lambda_1) \times \text{non-informative}(c_1)$							
	-	1	2	1.31310	1.40730	1.44270	1.46070	0.09887	0.04394	0.02831	0.02160
	-	2	2	1.23820	1.36610	1.41420	1.43890	0.12093	0.04954	0.03057	0.02269
	-	2	1	1.29320	1.39480	1.43370	1.45370	0.09796	0.04381	0.02818	0.02145
	-	3	1	1.22370	1.35520	1.40590	1.43230	0.12214	0.05027	0.03089	0.02281

6. Discussion

From empirical results in tables (4 -1) to (4-3), corresponding to the smallest values of MSE, we listed the best estimators using Bayes estimation, under the squared error loss function (SELF), under different conjugate prior of the parameter for all the true value of $\theta = 0.5, 1, 1.5$. We see Bayes estimators under different double prior selection are too close the true values $\theta = 0.5, 1, 1.5$, and the values of MSE for all sample size. Bayes estimation gave the best estimation according to the to the smallest values of MSE comparative with the values of MSE of MLE. As shown below in Table-A.

Table-A: The estimation of MLE and the best estimators of Bayes estimation for θ under the squared error loss function (SELF).

Method of estimation	Estimated value ($\hat{\theta}_{SE}$)				MSE($\hat{\theta}$)			
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 0.5$	0.52269	0.51209	0.5082	0.50692	0.01202	0.00531	0.00355	0.00259
Bayes estimation when the true value of θ is $\theta = 0.5$ when								
$\pi_1(\theta \setminus t)$ with $P_1(\theta) \propto \text{gamma}(a=1, b=2) \times \text{erlang}(\lambda=6)$	0.48116	0.49157	0.49459	0.49673	0.00732	0.00409	0.00298	0.00226
$\pi_2(\theta \setminus t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=6) \times \text{exponential}(\lambda_1=4)$	0.46414	0.48230	0.48825	0.49189	0.00730	0.00403	0.00294	0.00223
$\pi_3(\theta \setminus t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=6) \times \text{exponential}(\lambda_1=4)$	0.48092	0.49157	0.49461	0.49675	0.00853	0.00444	0.00314	0.00236
$\pi_4(\theta \setminus t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=3) \times \text{non-informative}(c_1=1)$	0.4907	0.49655	0.49794	0.49926	0.00895	0.00456	0.00321	0.00239
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1$	0.04866	0.02182	0.01446	0.01129	0.04866	0.02182	0.01446	0.01129
Bayes estimation when the true value of θ is $\theta = 1$ when								
$\pi_1(\theta \setminus t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	0.95577	0.97816	0.98831	-	0.02667	0.01603	0.01161	-
$\pi_2(\theta \setminus t)$ with $P_2(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=6)$	-	-	-	0.97481	-	-	-	0.00935
$\pi_3(\theta \setminus t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	0.92107	0.95960	0.97564	0.98400	0.03090	0.01718	0.01206	0.00968
$\pi_4(\theta \setminus t)$ with $P_4(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c_1=1)$	0.91926	0.95915	0.97548	0.98393	0.03302	0.01782	0.01237	0.00987
$\pi_4(\theta \setminus t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c_1=1)$	-	-	0.98834	-	-	-	0.01288	-
$\pi_4(\theta \setminus t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=3) \times \text{non-informative}(c_1=1)$	0.91740	0.95870	-	0.98387	0.03539	0.01850	-	0.01007
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1.5$	1.55380	1.52760	1.52320	1.52120	0.10950	0.04910	0.03254	0.02540
Bayes estimation when the true value of θ is $\theta = 1.5$ when								
$\pi_1(\theta \setminus t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	1.30180	1.39050	1.42830	1.44850	0.07683	0.04015	0.02737	0.02154
$\pi_2(\theta \setminus t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	1.27040	1.37490	1.41800	1.44090	0.09213	0.04471	0.02946	0.02273
$\pi_3(\theta \setminus t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c_1=1)$	1.28540	1.38570	1.42620	1.44740	0.09078	0.04428	0.02934	0.02276
$\pi_4(\theta \setminus t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c_1=1)$	1.37610	1.43800	1.46300	1.47580	0.07963	0.04163	0.02857	0.02266

And for empirical results in tables (4-4) to (4-7) , corresponding to the smallest values of MWSE ,we listed the best estimators using Bayes estimation , under the weighted error loss function (WSELF), under different conjugate prior of the parameter for all the true value of $\theta = 0.5, 1, 1.5$. We see Bayes estimators under different double prior selection are too close the true values $\theta = 0.5, 1, 1.5$, and the values of MWSE for all sample size. Bayes estimation gave the best estimation according to the to the smallest values of MWSE comparative with the values of MWSE of MLE.As shown below in Table-B.

Table-B: The estimation of MLE and the best estimators of Bayes estimation for under the weighted error loss function (WSELF).

Method of estimation	Estimated value ($\hat{\theta}_{WSE}$)				MWSE($\hat{\theta}$)			
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 0.5$	0.52269	0.51209	0.5082	0.50692	0.02405	0.01062	0.00710	0.00519
Bayes estimation when the true value of θ is $\theta = 0.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=6)$	0.49898	0.50102	-	-	0.01499	0.00836	-	-
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=1, b=2) \times \text{erlang}(\lambda=6)$	-	-	0.48817	0.49186	-	-	0.00603	0.00455
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda, =4)$	0.46334	0.48212	0.48817	0.49186	0.01561	0.00838	0.00603	0.00455
$\pi_2(\theta t)$ when $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c, =1)$	0.48092	0.49157	0.49461	0.49675	0.01706	0.00888	0.00629	0.00472
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda, =3) \times \text{non-informative}(c, =1)$	0.47107	0.48662	0.49130	0.49427	0.01802	0.00910	0.00639	0.00476
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1$	1.03580	1.01840	1.01550	1.01410	0.04866	0.02182	0.01446	0.01129
Bayes estimation when the true value of θ is $\theta = 1$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	0.92282	0.96004	0.97580	0.98406	0.02900	0.01658	0.01177	0.00949
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda, =1)$	0.88695	0.94115	0.96297	0.97435	0.03566	0.01842	0.01255	0.00990
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c, =1)$	0.88390	0.94035	0.96265	0.97419	0.03798	0.01909	0.01286	0.01009
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda, =2) \times \text{non-informative}(c, =1)$	0.91548	0.95824	0.97516	0.98381	0.03805	0.01922	0.01303	0.01027
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1.5$	1.55380	1.52760	1.52320	1.52120	0.07300	0.03273	0.02169	0.01693
Bayes estimation when the true value of θ is $\theta = 1.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	1.25690	1.36480	1.41020	1.43460	0.06273	0.03028	0.01982	0.01520
$\pi_1(\theta t)$ when $P_1(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda, =1)$	1.22330	1.34850	1.39960	1.42680	0.07539	0.03394	0.02149	0.01615
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=2) \times \text{non-informative}(c, =3)$	1.26600	1.38050	1.42400	1.44630	0.07277	0.03274	0.02103	0.01605
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda, =1) \times \text{non-informative}(c, =2)$	1.34260	1.42220	1.45270	1.46820	0.06279	0.03024	0.02011	0.01569

Also for empirical results in tables (4-8) to (4-11), corresponding to the smallest

values of $\hat{MLINEX}(\theta)$,we listed the best estimators using Bayes estimation , under modified linear exponential (MLINEX) loss function (MLF)with $w=1$ and $c=1$, under different conjugate prior of the parameter for all the true value of $\theta = 0.5, 1, 1.5$. We see Bayes estimators under different double prior selection are

too close the true values $\theta = 0.5, 1, 1.5$, and the values of $\hat{MLINEX}(\theta)$ for all sample size .Bayes estimation gave the best estimation according to the to the smallest values of $\hat{MLINEX}(\theta)$ comparative with the values of $\hat{MLINEX}(\theta)$ of MLE.As shown below in Table-C.

Table-C: The estimation of MLE and the best estimators of Bayes estimation for under modified linear exponential (MLINEX) loss function (MLF) with $w=1$ and $c=1$.

Method of estimation	$\hat{\theta}_{MLF}$				$\hat{MLINEX}(\theta)$			
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 0.5$	0.52269	0.51209	0.5082	0.50692	0.02126	0.00985	0.00672	0.00497
Bayes estimation when the true value of θ is $\theta = 0.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=6)$	0.49898	0.50102	0.50102	0.50160	0.01475	0.00813	0.00592	0.00450
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=2)$	0.48102	0.49157	0.49460	0.49674	0.01663	0.00860	0.00613	0.00461
$\pi_3(\theta t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c_1=1)$	0.48092	0.49157	0.49461	0.49675	0.01791	0.00894	0.00630	0.00470
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c_1=1)$	0.48087	0.49159	0.49463	0.49677	0.01936	0.00931	0.00647	0.00480
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1$	1.03580	1.01840	1.01550	1.01410	0.02155	0.01028	0.00694	0.00541
Bayes estimation when the true value of θ is $\theta = 1$ when								
$\pi_1(\theta t)$ with $f P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	0.92282	0.96004	0.97580	0.98406	0.01636	0.00882	0.00613	0.00484
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	0.88695	0.94115	0.96297	0.97435	0.02108	0.01010	0.00668	0.00514
$\pi_3(\theta t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c_1=1)$	0.88390	0.94035	0.96265	0.97419	0.02252	0.01047	0.00685	0.00524
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=1) \times \text{non-informative}(c_1=2)$	0.91353	0.95778	0.97501	0.98375	0.02318	0.01062	0.00696	0.00533
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1.5$	1.55380	1.52760	1.52320	1.52120	0.02155	0.01028	0.00694	0.00541
Bayes estimation when the true value of θ is $\theta = 1.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	1.25690	1.36480	1.41020	1.43460	0.02566	0.01151	0.00728	0.00543
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	1.22330	1.34850	1.39960	1.42680	0.03142	0.01305	0.00796	0.00581
$\pi_3(\theta t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=2) \times \text{non-informative}(c_1=3)$	1.26600	1.38050	1.42400	1.44630	0.02999	0.01231	0.00763	0.00565
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c_1=1)$	1.32110	1.40930	1.44350	1.46110	0.02412	0.01087	0.00703	0.00534

Finally, for empirical results in tables (7 -1) to (7-3) , corresponding to the smallest values of $\hat{MLINEX}(\theta)$,we listed the best estimators using Bayes estimation , under modified linear exponential (MLINEX) loss function (MLF) with $w=1$ and $c=2$, under different conjugate prior of the parameter for all the true value of $\theta = 0.5, 1, 1.5$. We see Bayes estimators under different double prior selection are too close the true values $\theta = 0.5, 1, 1.5$, and the values of $\hat{MLINEX}(\theta)$ for all sample size .Bayes estimation gave the best estimation according to the to the smallest values of $\hat{MLINEX}(\theta)$ comparative with the values of $\hat{MLINEX}(\theta)$ of MLE.As shown below in Table-D.

Table-D: The estimation of MLE and the best estimators of Bayes estimation for under modified linear exponential (MLINEX) loss function (MLF) with $w=1$ and $c=2$.

Method of estimation	$\hat{\theta}_{MLF}$				$\hat{\theta}_{MLINEX}$			
	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 0.5$	0.52269	0.51209	0.50820	0.50692	0.09065	0.04095	0.02766	0.02035
Bayes estimation when the true value of θ is $\theta = 0.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	0.48999	0.49627	0.49780	0.49916	0.05922	0.03280	0.02385	0.01812
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=2)$	0.47168	0.48673	0.49134	0.49427	0.06728	0.03478	0.02474	0.01856
$\pi_3(\theta t)$ when $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c=1)$	0.47120	0.48663	0.49131	0.49426	0.07246	0.03618	0.02541	0.01894
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c=1)$	0.47074	0.48655	0.49127	0.49425	0.07833	0.03767	0.02610	0.01933
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1$	1.03580	1.01840	1.01550	1.01410	0.09177	0.04239	0.02836	0.02211
Bayes estimation when the true value of θ is $\theta = 1$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	0.90619	0.95094	0.96953	0.97927	0.06686	0.03560	0.02459	0.01944
$\pi_2(\theta t)$ with $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	0.86973	0.93187	0.95662	0.96951	0.08591	0.04078	0.02682	0.02064
$\pi_3(\theta t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=4) \times \text{non-informative}(c=1)$	0.86605	0.93089	0.95621	0.96931	0.09162	0.04225	0.02749	0.02104
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=1) \times \text{non-informative}(c=2)$	0.89345	0.94775	0.96830	0.97872	0.09414	0.04281	0.02790	0.02144
Sample Size(n)	25	50	75	100	25	50	75	100
MLE when the true value of θ is $\theta = 1.5$	1.55380	1.52760	1.52320	1.52120	0.09177	0.04239	0.02836	0.02211
Bayes estimation when the true value of θ is $\theta = 1.5$ when								
$\pi_1(\theta t)$ with $P_1(\theta) \propto \text{gamma}(a=3, b=2) \times \text{erlang}(\lambda=4)$	1.23430	1.35180	1.40110	1.42760	0.10388	0.04641	0.02923	0.02179
$\pi_2(\theta t)$ when $P_2(\theta) \propto \text{erlang}(\lambda=4) \times \text{exponential}(\lambda_1=1)$	1.19960	1.33520	1.39040	1.41970	0.12613	0.05244	0.03191	0.02330
$\pi_3(\theta t)$ with $P_3(\theta) \propto \text{erlang}(\lambda=2) \times \text{non-informative}(c=3)$	1.23820	1.36610	1.41420	1.43890	0.12093	0.04954	0.03057	0.02269
$\pi_4(\theta t)$ with $P_4(\theta) \propto \text{exponential}(\lambda_1=2) \times \text{non-informative}(c=1)$	1.29320	1.39480	1.43370	1.45370	0.09796	0.04381	0.02818	0.02145

In general, the parameter estimates using Bayes estimation methods are close to the true values comparative to the other estimated values using maximum likelihood estimation, at certain values for the parameters of the conjugate prior of the parameter θ was considered as combination for the two different prior distribution.

5. Conclusions

When we compared the estimated values for the shape parameter (θ) of the Power Function Distribution (PFD-I) by using the methods in this study. We can conclude that Bayes' estimators for the unknown shape parameter θ was considered, under various double priors at certain values for the parameters of the double prior distribution and under the squared error loss function and the weighted error loss function and modified linear exponential (MLINEX) loss function with $w=1$ and $c=1,2$, for all the true value of θ and for all samples sizes (n) are better than other estimators by using maximum likelihood estimation. we find that best estimation under the squared error loss function (SELF) for each of the value of $\theta = 0.5, 1, 1.5$, according to the smallest values of MSE, by using the conjugate prior of the parameter θ was considered as

- (erlang($\lambda = 6$) $\times$ exponential($\lambda_1 = 4$)) distribution for all samples, when the true value $\theta = 0.5$.
- (gamma($a = 3, b = 2$) $\times$ erlang($\lambda = 4$)) distribution and for all samples except $n=100$, when the true value $\theta = 1$ and $\theta = 1.5$.

And we find that best estimation under the weighted error loss function (WSELF) for each of the value of $\theta = 0.5, 1, 1.5$, according to the smallest values of MWSE, by using the conjugate prior of the parameter θ was considered as

- (gamma($a = 3, b = 2$) $\times$ erlang($\lambda = 6$)) distribution for all samples, when the true value $\theta = 0.5$.
- (gamma($a = 3, b = 2$) $\times$ erlang($\lambda = 4$)) distribution and for all samples except $n=100$, when the true value $\theta = 1$.
- (exponential($\lambda_1 = 2$) $\times$ non - informative($c_1 = 1$)) distribution and for all samples except $n=100$, when the true value $\theta = 1.5$.

Also we find that best estimation under modified linear exponential (MLINEX) loss function (MLF) with $w=1$ and ($c=1, 2$), for each of the value of $\theta = 0.5, 1, 1.5$,

according to the smallest values of $\text{MLINEX}(\hat{\theta})$, by using the conjugate prior of the parameter θ was considered as

- (gamma($a = 3, b = 2$) $\times$ erlang($\lambda = 6$)) distribution for all samples, when the true value $\theta = 0.5$.
- (gamma($a = 3, b = 2$) $\times$ erlang($\lambda = 4$)) distribution and for all samples, when the true value $\theta = 1$.
- (exponential($\lambda_1 = 2$) $\times$ non - informative($c_1 = 1$)) distribution and for all samples, when the true value $\theta = 1.5$.

References

1. Bickel, P.J. & Doksum, K. A., (1977), *Mathematical Statistics: Basic Ideas and Selected Topics*. Holden- Day, Inc., San Francisco.
2. Hanif, S., Al-Ghamdi, S.D., Khan, K., Shahbaz, M.Q., (2015), " Bayesian estimation for parameters of power function distribution under various Priors", *Mathematical Theory and Modeling*, Vol.5, No.13, ISSN 2224-5804 (Paper).
3. Kifayat, T., Aslam, M and Ali, S. (2012), "Bayesian Inference for the Parameter of the Power Distribution", *Journal of Reliability and Statistical Studies*; 5, p.45-58.
4. M. Meniconi, D. M. Barry, (1996), " Microelectronics reliability", volume 36, issue 9, pages 1207-1212.
5. Rahman, H., Roy, M. K. and Baizid, A. R. (2012), " Bayes estimation under conjugate prior for the case of power function distribution", *American Journal of Mathematics and Statistics*, 2(3): 44-48.
6. Ronak, M., and Achyut, S., Patel, S., (2016), "The double prior selection for the parameter for the power function under type- II censoring", *International Journal of Current Research* Vol. 8, Issue, 05, pp.30454-30465.
7. Sultan, R., Sultan, H. and Ahmad, S. P., (2014), " Bayesian analysis of power function distribution under double Priors ", *J. Stat. Appl. Pro.* 3, No. 2, 239-249.
8. Encyclopedia (2016) "The Erlang distribution" ,This page was last modified on 3 June 2016, at 16:46. <http://en.wikipedia.org/wiki>.
9. Encyclopedia (2016) "The gamma distribution", This page was last modified on 17 August 2016, at 22:09. <http://en.wikipedia.org/wiki>
10. Zaka, A. and Akhtar, A. S. (2013), " Methods for estimating the parameters of the Power function distribution", *Pakistan Journal of Statistics and Operational Research*, 9(2), 213-224.

التقدير البيزي لمعلمة الشكل لتوزيع (PFD-I) دالة القوى لاستعمال دوال أولية فوقية

أ.م.د. جنان عباس ناصر العبيدي
الكلية التقنية الادارية – بغداد، العراق
Jinan.abbas@mtu.edu.iq

Received:6/12/2020

Accepted:24/1/2021

Published: March/ 2021

هذا العمل مرخص تحت اتفاقية المشاع الابداعي نسب المصنف - غير تجاري - الترخيص العمومي الدولي 4.0

[Attribution-NonCommercial 4.0 International \(CC BY-NC 4.0\)](https://creativecommons.org/licenses/by-nc-sa/4.0/)



مستخلص البحث

في هذا البحث، نختبر خصائص مقدرات بيز لمعلمة الشكل لتوزيع دالة القوى من النوع الأول باستعمال نوعين من التوزيعات أولية ، واستعمال دوال خسارة مختلفة ، ومقارنتها مع مقدرات الإمكان الأعظم . في العديد من التطبيقات العملية، ربما يكون لدينا معلومتين مختلفة حول التوزيع الاولي لمعلمة الشكل لتوزيع دالة القوى من النوع الأول، الذي يكون له تأثير على تقدير معلمة الشكل. لذا فقد استعملنا نوعين مختلفين من الدوال الأولية المرافقة لمعلمة الشكل (θ) لتوزيع دالة القوى من النوع الأول (PFD-I) لتقديرها . افترضت الدالة المرافقة الأولية لمعلمة الشكل (θ) كتوليفة من توزيعين أولية مختلفة، كتوزيع كاما مع توزيع ارلنك وتوزيع ارلنك مع الاسي وتوزيع ارلنك مع توزيع غير معلوماتي وتوزيع الاسي مع توزيع غير معلوماتي. فقد تم اشتقاق مقدرات بيز لمعلمة الشكل لتوزيع دالة القوى من النوع الأول باستعمال ثلاث أنواع لدالة الخسارة كدالة الخسارة التربيعية ودالة الخسارة التربيعية الموزونة ودالة الخسارة الاسية الخطية المحورة. بالإضافة الى طريقة التقدير الكلاسيكية المتمثلة بطريقة تقدير الامكان الاعظم . استعملنا أسلوب المحاكاة للحصول على نتائج هذا البحث لحالات مختلفة لمعلمة الشكل (θ) لتوزيع دالة القوى من النوع الأول (PFD-I) استعملت لتوليد البيانات و لأحجام مختلفة من العينات.

نوع البحث: ورقة بحثية.

المصطلحات الرئيسية للبحث: توزيع دالة القوى من النوع الأول، تقدير الإمكان الأعظم، تقدير بيز، دالة الخسارة التربيعية، ودالة الخسارة التربيعية الموزونة، ودالة الخسارة الاسية الخطية المحورة.