

A Comparative Analysis of Shrinkage Bayesian Estimation Methods for Parameters of the Lognormal Distribution

Nadiha Abed Habeeb

University of Karbala, College of Engineering, Department of Biomedical Engineering, Karbala, Iraq
nadia.habeeb@alzahraa.edu.iq, <https://orcid.org/0009-0008-3552-744X>

How to cite this article: Nadiha Abed Habeeb (2024). A Comparative Analysis of Shrinkage Bayesian Estimation Methods for Parameters of the Lognormal Distribution. *Library Progress International*, 44(3), 11240-11249.

ABSTRACT

The Lognormal distribution serves as a fundamental model across various disciplines, offering a versatile framework for analyzing positively skewed data. Lognormal distributions are frequently employed to represent stock prices, as these prices cannot be negative and typically adhere to a multiplicative process. Additionally, this distribution is utilized to model the time until failure for various products. It is also applicable in the analysis of certain environmental variables, such as the concentrations of pollutants. In Bayesian estimation, the choice of prior distribution significantly influences parameter estimation, particularly in scenarios with limited sample sizes. Shrinkage Bayesian estimation methods have emerged as powerful tools to address the challenges of small sample sizes by incorporating prior information effectively.

This research conducts a comparative analysis of shrinkage Bayesian estimation methods. Through comprehensive simulations, we investigate the impact of these priors on parameter estimation accuracy by limited sample sizes and initial parameter values.

simulation experiments showed that Shrinkage Bayesian estimator method give the best estimators with minimum mean square error comparing with Maximum likelihood estimator.

Other simulations can be performed for estimation methods (shrinkage, moments, modified moments). shrinkage Bayesian estimation methods can also be performed on (Weibull, Gamma distribution) distributions and the results compared.

Keywords

Shrinkage Bayesian estimators, Lognormal Distribution, Maximum likelihood estimation method, simulation experiments.

1-Research Problem

In numerous practical contexts, the estimation of parameters for the lognormal distribution is crucial, particularly when addressing skewed datasets. Conventional techniques, such as Maximum Likelihood Estimation (MLE), are frequently employed for this purpose; however, they may exhibit inefficiencies or biases, especially when sample sizes are limited or when prior information is accessible. Bayesian estimation offers a methodology that allows for the integration of prior knowledge, and the application of shrinkage techniques within Bayesian estimation enhances the robustness and precision of parameter estimates by "shrinking" them towards a central value or by amalgamating information from various sources.

2-Importance of the Research

The precision of parameter estimation in statistical models is of paramount importance, particularly in sectors such as finance, engineering, and environmental science, where critical decisions are heavily dependent on the accuracy of these estimates. Shrinkage Bayesian estimation techniques provide a means to integrate prior knowledge, thereby enhancing estimation precision, especially in scenarios involving limited sample sizes or incomplete datasets. This study aims to compare these methods specifically in the context of the lognormal distribution, with the goal of determining which approaches produce the most dependable outcomes across various conditions. The lognormal distribution finds extensive application in multiple domains, including finance (for stock price modeling), environmental science (for pollutant concentration modeling), and engineering (for reliability assessments). The accuracy of parameter estimation significantly influences decision-making in these areas. This research will equip practitioners with a comprehensive understanding of the most appropriate shrinkage Bayesian methods tailored to their specific requirements, ultimately facilitating more informed and effective decision-making.

3-Research Objectives

Conduct a comprehensive evaluation of different shrinkage Bayesian estimation techniques aimed at estimating the parameters of the lognormal distribution. This assessment should focus on their effectiveness in terms of bias, mean squared error (MSE), and coverage probability.

Engage in thorough simulation studies to compare the performance of the shrinkage Bayesian methods across various scenarios, which will include differing sample sizes, shape parameters, and scale parameters of the lognormal distribution. These simulations will provide insights into the robustness and dependability of each estimation method.

4-Introduction

The Lognormal distribution stands as a cornerstone in statistical modeling, finding wide applications in fields ranging from finance to biology due to its ability to represent positively skewed data. Parameter estimation for the Lognormal distribution plays a pivotal role in various statistical analyses, influencing decisions and insights drawn from data. In Bayesian estimation, the incorporation of prior information is paramount, particularly when dealing with limited sample sizes or when prior knowledge about the parameters is available.

However, selecting appropriate prior distributions for Bayesian estimation poses a significant challenge, especially in scenarios where prior information is scarce or when the sample size is small. Shrinkage Bayesian estimation methods offer a promising avenue to address these challenges by effectively incorporating prior information and improving the robustness of parameter estimates.

In this research, we undertake a comprehensive comparative analysis of shrinkage Bayesian estimation methods for parameters of the Lognormal distribution through simulation experiments. The primary objective is to compare various shrinkage priors, precision, and robustness under different sample size and prior specification scenarios. By systematically examining a range of shrinkage priors, including hierarchical and empirical Bayes approaches, we aim to elucidate their relative efficacy in Lognormal parameter estimation. Through extensive simulations, we explore how these methods perform across varying levels of prior information availability and sample sizes, providing insights into their strengths and limitations.

Many researches have been conducted on it

The research of (Hamura, Y.) in (2023), In this paper, he considers the simultaneous estimation of Poisson parameters when auxiliary information from the aggregated data can be used. It uses standardized squared error and entropy loss functions. The Bayesian shrinkage estimator is derived based on the conjugate prior. It compares direct and Bayesian estimators for hazard functions with different priors constructed based on different subsets of observations. It creates the conditions for domination and demonstrates minimax and permissibility in a simple framework. [5]

The research of (Bargoshadi, S. A., & Bevrani, H.) in (2024), The main goal of this work is to apply linear shrinkage estimation techniques and pretest shrinkage estimation to estimate the parameters of two 2-parameter Burr-XII distributions. Furthermore, predictions of future observations are performed using classical and Bayesian methods under a common type II review scheme. The efficiency of shrinkage estimation is compared with maximum likelihood and Bayesian estimation obtained by expectation estimation. For Bayesian estimation, both informative and non-informative prior distributions are considered. Additionally, various loss functions are considered, including squared error, linear exponential, and generalized entropy. Calculate approximate confidence, credibility, and maximum likelihood density intervals. Evaluate the performance of estimation methods.[2]

The results of this research are expected to contribute to the further development of Bayesian estimation methods for lognormal distribution parameters and provide researchers and practitioners with information on selecting appropriate shrinkage priors to achieve robust parameter estimation in practical applications. guide. Furthermore, our results have the potential to improve the understanding of Bayesian inference techniques in the context of skewed data distributions and enable more accurate and reliable statistical analyzes in various research fields.

5-Log normal distribution

The lognormal distribution is a continuous probability distribution commonly used to model positive and partial variables. It is widely used in various fields such as finance, biology, and engineering, and data in these fields naturally follows an exponential growth pattern.

The Lognormal distribution is characterized by two parameters:

the location parameter, μ , which represents mean of the natural logarithm of the variable, and the scale parameter, σ , which measures the spread or variability of the logarithm of the variable. Unlike the Normal distribution, which describes variables on the linear scale, the Lognormal distribution describes variables on the logarithmic scale. As a result, the Lognormal distribution allows for modeling variables that have strictly positive values and exhibit right-skewness.

The probability density function (pdf) for the distribution is given by: - [4,6,8]

$$f(y|\mu, \sigma) = \frac{1}{y\sigma\sqrt{2\pi}} e^{-\frac{(\ln(y)-\mu)^2}{2\sigma^2}} \dots (1)$$

Where ($y > 0$)

(μ) and (σ) are (location and scale) parameters respectively

The cumulative distribution function (CDF) is:-

$$F(y|\mu, \sigma) = \frac{1}{2} + \frac{1}{2} \operatorname{erf}\left(\frac{\ln(y) - \mu}{\sigma\sqrt{2}}\right) \quad \dots (2)$$

(erf) denotes error function.

The Lognormal distribution is widely used due to its ability to model variables with asymmetric distributions and its relevance to real-world phenomena. Understanding and accurately estimating its parameters are crucial for various statistical analyses and decision-making processes in fields where Lognormally distributed data are encountered.

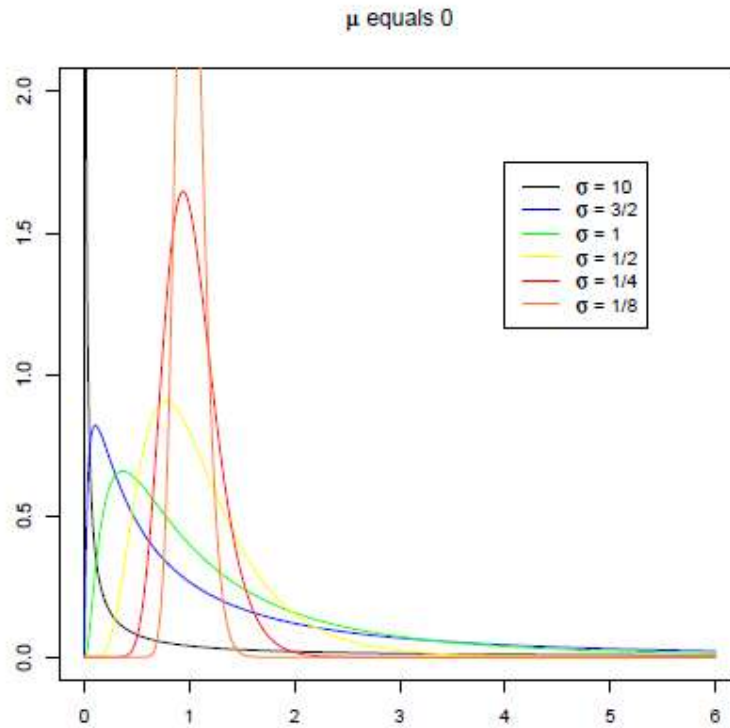


Figure 1: Lognormal Density Plots with ($\mu = 0$)[3]

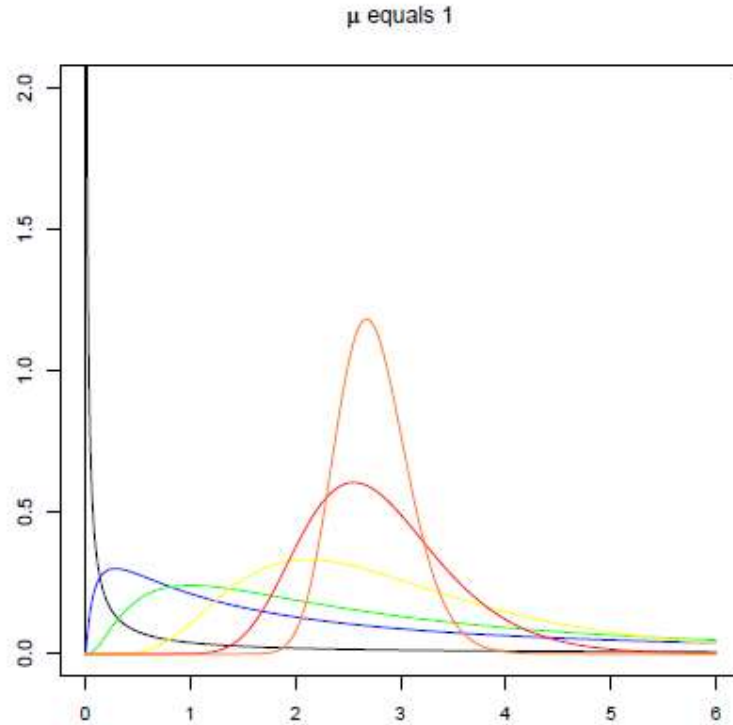


Figure 2: Lognormal Density Plots with ($\mu = 1$)[3]

6-Estimation method

Estimating parameters of a distribution is a crucial task in statistical analysis, especially when dealing with real-world data that exhibit positive skewness. Several methods can be employed for Lognormal parameter estimation, including

1.6-Maximum Likelihood Estimation (MLE)

MLE is a widely used method for estimating the parameters of the lognormal distribution. Given a data set, MLE attempts to find parameters that maximize the likelihood function, which measures the probability of observing the given data given a hypothetical distribution. In the lognormal distribution, the likelihood function is based on the product of the probability density function (pdf) values for each observation. MLE estimates of the lognormal distribution parameters (μ and σ) can be obtained by maximizing this likelihood function.

likelihood function of the lognormal distribution for a series of random variables with size (n) is: - [9]

$$L(\mu, \sigma^2 / y) = \prod_{i=1}^n [f(y_i / \mu, \sigma^2)] \quad \dots (3)$$

$$= \prod_{i=1}^n \left((2\pi\sigma^2)^{-\frac{1}{2}} y_i^{-1} e^{-\frac{(\ln(y_i) - \mu)^2}{2\sigma^2}} \right) \quad \dots (4)$$

$$= ((2\pi\sigma^2)^{-\frac{n}{2}} \prod_{i=1}^n e^{-\frac{(\ln(y_i) - \mu)^2}{2\sigma^2}} y_i^{-1}) \quad \dots (5)$$

the natural log of the likelihood function:

$$\mathcal{L}\left(\mu, \frac{\sigma^2}{y}\right) = -\frac{n}{2} \ln(2\pi\sigma^2) - \sum_{i=1}^n \ln(y_i) - \frac{\sum_{i=1}^n [\ln(y_i)]^2}{2\sigma^2} + \frac{\sum_{i=1}^n \mu \ln(y_i)}{\sigma^2} - \frac{n\mu^2}{2\sigma^2} \quad \dots (6)$$

the gradient of $\mathcal{L}$ with respect to (μ) and (σ^2) and set it equal to 0

Thus, the maximum likelihood estimators are:-

$$\hat{\mu} = \frac{\sum_{i=1}^n \ln(y_i)}{n} \quad \dots (7)$$

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n [\ln(y_i) - \frac{\sum_{i=1}^n \ln(y_i)}{n}]^2}{n} \quad \dots (8)$$

2.6-Shrinkage Bayesian estimators

Shrinkage estimators play a crucial role in Bayesian statistics as they enhance parameter estimation by leveraging prior information. These estimators are particularly valuable in scenarios involving high-dimensional data or a large number of parameters to estimate relative to the sample size.

The concept behind shrinkage estimators involves pulling the estimates towards a central value, typically the mean of the prior distribution. This process effectively reduces the variance of the estimator, resulting in more reliable and robust estimates, especially when the sample size is limited.

There exist various types of shrinkage estimators, each employing different loss functions. Bayesian shrinkage estimators, in particular, combine information from the data (likelihood function) with prior information (prior distribution) to derive a posterior distribution of the parameters. This posterior distribution represents an updated belief about the parameters after observing the data.

Shrinkage Bayesian estimator depend on MLE estimators

the posterior distribution of the parameter (ϑ) is the gamma distribution ($G(n, t)$) which has the following probability density function: - [1]

$$f(\vartheta/n) = \frac{t^n \vartheta^{n-1} e^{-\beta t}}{\Gamma(n)} \quad \dots (9)$$

The shrinkage Bayesian estimators depend on the loss function such that: -

1-Shrinkage Bayesian estimator with square loss function (SBE1)

$$L_1(\vartheta, \hat{\vartheta}_s) = [\vartheta - \hat{\vartheta}_s]^2 \quad \dots (10)$$

the posterior risk function for (ϑ) will be

$$\rho_1(\vartheta, \hat{\vartheta}_{s1}) = E_{\vartheta}[\vartheta - \hat{\vartheta}_{s1}]^2 = \hat{\vartheta}_{s1}^2 - 2\hat{\vartheta}_{s1} \frac{n}{t} + \frac{n(n+1)}{[t]^2} \quad \dots (11)$$

$$\frac{\partial \rho_1(\vartheta, \hat{\vartheta}_{s1})}{\partial \hat{\vartheta}_{s1}} = 2\hat{\vartheta}_{s1}^2 - 2 \frac{n}{t}$$

$$\frac{\partial \rho_1(\vartheta, \hat{\vartheta}_{s1})}{\partial \hat{\vartheta}_{s1}} = 0$$

the generalized Bayesian estimator of the parameter(ϑ) under the squared error loss function will be

$$\hat{\vartheta}_{s1} = \frac{n}{t} \quad \dots (12)$$

For ($\vartheta = \alpha$)

$$\hat{\alpha}_{s1} = \frac{n}{\sum_{i=1}^n \ln(y_i)} \quad \dots (13)$$

$$\hat{\beta}_{s1} = \frac{n}{\sum_{i=1}^n [\ln(y_i) - \frac{\sum_{i=1}^n \ln(y_i)}{n}]^2} \quad \dots (14)$$

the shrinkage estimator is

$$\hat{\vartheta}_{sh1} = \tau(\hat{\vartheta}_{s1} - \vartheta_0) + \vartheta_0 \quad \dots (15)$$

$$\rho_1(\vartheta, \hat{\vartheta}_{sh1}) = E[\vartheta - \hat{\vartheta}_{sh1}]^2$$

$$\rho_1(\vartheta, \hat{\vartheta}_{sh1}) = E_{\vartheta}[\tau(\hat{\vartheta}_{s1} - \vartheta_0) + \vartheta_0 - \vartheta]^2$$

$$\rho_1(\vartheta, \hat{\vartheta}_{sh1}) = [\tau\hat{\vartheta}_{s1}]^2 + (2\tau + 2\tau^2)\vartheta_0\hat{\vartheta}_{s1} + [1 - \tau]^2\vartheta_0^2 - 2\tau\hat{\vartheta}_{s1}\frac{n}{t} + 2(\tau - 1)\frac{n}{t}\vartheta_0 + \frac{n(n+1)}{[t]^2} \quad \dots (16)$$

$$\frac{\partial \rho_1(\vartheta, \hat{\vartheta}_{sh1})}{\partial \tau} = 2\tau\vartheta_{sh1}^2 + 2(1 - 2\tau)\vartheta_0\hat{\vartheta}_{s1} - 2(1 - \tau)\vartheta_0^2 - \frac{2n(\hat{\vartheta}_{s1} + \vartheta_0)}{t}$$

$$\frac{\partial \rho_1(\vartheta, \hat{\vartheta}_{sh1})}{\partial \tau} = 0$$

$$\tau = \frac{t(\vartheta_0^2 - \vartheta_0(\frac{n}{t}) + n(\frac{n}{t} + \vartheta_0))}{t(\frac{n}{t} - \vartheta_0)^2} \quad \dots (17)$$

generalized Bayesian shrinkage estimator for loss function will be: -

$$\hat{\vartheta}_{sh1} = \left(\frac{t(\vartheta_0^2 - \vartheta_0(\frac{n}{t}) + n(\frac{n}{t} + \vartheta_0))}{t(\frac{n}{t} - \vartheta_0)^2} \right) \left(\frac{n}{t} - \vartheta_0 \right) + \vartheta_0 \quad \dots (18)$$

2-Bayesian Shrinkage Estimator with LINEX Loss Function (SBE2)

With [5]

$$L_2(\Delta) = e^{a\Delta} - (a\Delta + 1) \quad \dots (19)$$

Such that $(\Delta = \frac{\vartheta_{sh2}}{\vartheta})$

$$\rho_2(\vartheta, \hat{\vartheta}_{s2}) = E_b \left[e^{a(\frac{\vartheta_{sh2}}{\vartheta} - 1)} - a \left(\frac{\vartheta_{sh2}}{\vartheta} - 1 \right) - 1 \right] \dots (20)$$

$$\rho_2(\vartheta, \hat{\vartheta}_{s2}) = e^a E_b \left[e^{a(\frac{\vartheta_{sh2}}{\vartheta})} \right] - a \hat{\vartheta}_{sh} E_b \left(\frac{1}{\vartheta} \right) + a - 1$$

With

$\vartheta \sim \text{Gamma dist.}(n, t)$

$\frac{1}{\vartheta} \sim \text{inv. Gamma dist.}(n, t)$ with the following probability density function: -

$$f\left(\frac{1}{\vartheta}\right) = \frac{[t]^n}{\Gamma(n)} \vartheta^{-(n+1)} e^{-\frac{t}{\vartheta}} \vartheta \dots (21)$$

$$E\left(\frac{1}{\vartheta}\right) = \frac{t}{n-1}$$

$$\rho_2(\vartheta, \hat{\vartheta}_{s2}) = e^a \left(\frac{t}{t - a\hat{\vartheta}_{s2}} \right)^n - a\hat{\vartheta}_{s2} \left[\frac{t}{n-1} \right] + a - 1$$

$$\frac{\partial \rho_2(\vartheta, \hat{\vartheta}_{s2})}{\partial \hat{\vartheta}_{s2}} = an[t]^n e^{-a} [t - a\hat{\vartheta}_{s2}]^{-(n+1)} - \frac{a[t]}{n-1}$$

$$\frac{\partial \rho_2(\vartheta, \hat{\vartheta}_{s2})}{\partial \hat{\vartheta}_{s2}} = 0$$

Then generalized Bayesian estimator under LINEX loss function will be: -

$$\hat{\beta}_{s2} = \frac{1}{a} \left(\sum_{i=1}^n \text{Log}(1 - y_i) - (n \sum_{i=1}^n \text{Log}(1 - y_i)^{n-1} e^{-a}) \right) \frac{1}{n+1} \dots (22)$$

The Bayesian Shrinkage Estimator under LINEX Loss Function will be: -

$$\rho_2(\vartheta, \hat{\vartheta}_{sh2}) = E_b \left(e^{a(\frac{\vartheta_{sh2}}{\vartheta} - 1)} \right) - a\hat{\vartheta}_{sh2} E_b \left(\frac{1}{\vartheta} \right) + a - 1$$

$$\rho_2(\vartheta, \hat{\vartheta}_{sh2}) = e^{-a} \left(\frac{t}{t - a\hat{\vartheta}_{sh2}} \right)^n - a\hat{\vartheta}_{sh2} \frac{t}{n-1} + a - 1$$

the shrinkage estimator is defined as

$$\hat{\vartheta}_{sh2} = \tau(\hat{\vartheta}_{s2} - \vartheta_0) + \vartheta_0 \dots (23)$$

$$\rho_2(\vartheta, \hat{\vartheta}_{sh2}) = \frac{e^{-a}(t)^n}{[t - a\tau(\hat{\vartheta}_{s2} - \vartheta) - a\vartheta_0]^n} - a \frac{t}{n-1} \tau(\hat{\vartheta}_{s2} - \vartheta) - a\vartheta_0 \frac{t}{n-1}$$

$$\frac{\partial \rho_2(\vartheta, \hat{\vartheta}_{sh2})}{\partial \tau} = \frac{an(\hat{\vartheta}_{s2} - \vartheta)e^{-a}(t)^n}{[t - a\tau(\hat{\vartheta}_{s2} - \vartheta) - a\vartheta_0]^{(n+1)}} - a \frac{t(\hat{\vartheta}_{s2} - \vartheta)}{n-1}$$

$$\frac{\partial \rho_2(\vartheta, \hat{\vartheta}_{sh2})}{\partial \tau} = 0$$

We get

$$\tau = \frac{[t - a\vartheta_0] - [n(n-1)t^{n-1}e^{-a}]^{\frac{1}{n+1}}}{a[\frac{1}{a}[t - [nt^{n-1}e^{-a}]^{\frac{1}{n+1}} - \vartheta_0]]} \dots (24)$$

$$\hat{\vartheta}_{sh2} = \left(\frac{[t - a\vartheta_0] - [n(n-1)t^{n-1}e^{-a}]^{\frac{1}{n+1}}}{a[\frac{1}{a}[t - [nt^{n-1}e^{-a}]^{\frac{1}{n+1}} - \vartheta_0]]} \right) \left(\frac{1}{a} [t - [nt^{n-1}e^{-a}]^{\frac{1}{n+1}}] - \vartheta_0 \right) + \vartheta_0 \dots (25)$$

7-Simulation experiments

The simulation experiments depend on generating samples that follow a Lognormal distribution according to the following simulation parameters

$(n_1 = 50, n_2 = 75, n_3 = 100)$, $(\alpha_1 = 0.5, \alpha_2 = 1, \alpha_3 = 0.5)$, $(\beta_1 = 0.3, \beta_2 = 0.6, \beta_3 = 0.9)$ and the estimation methods (MLE, SH1, SH2)

Comparing simulation results with the following criteria's:-[7]

$$\delta = \text{Min}(|\vartheta - \hat{\vartheta}|) \dots (26)$$

Such that

(δ) Represent minimum absolute deference

(ϑ) Represent true value for the parameter

($\hat{\vartheta}$) Represent the estimated value for the parameter

$$MSE = \frac{\sum_{h=1}^r [\vartheta - \hat{\vartheta}]^2}{r} \dots (27)$$

(*MSE*) Represent mean square error

(*r*) Represent number of iteration

8-Experimental Results

After completing the simulation experiments, the following results were obtained

Table (1) first parameter estimators and the best method to estimate it

α	β	n	$\hat{\alpha}_{MLE}$	$\hat{\alpha}_{SH1}$	$\hat{\alpha}_{SH2}$	δ_{α}	<i>best</i>
0.5	0.3	50	0.502155	0.498209	0.501056	1.06E-03	3
0.5	0.3	75	0.500136	0.499941	0.499857	5.85E-05	2
0.5	0.3	100	0.500003	0.500003	0.499984	2.72E-06	2
0.5	0.6	50	0.597319	0.498996	0.500357	3.57E-04	3
0.5	0.6	75	0.563788	0.499883	0.500183	1.17E-04	2
0.5	0.6	100	0.533381	0.500003	0.500002	1.66E-06	3
0.5	0.9	50	0.538938	0.499519	0.502016	4.81E-04	2
0.5	0.9	75	0.516913	0.499739	0.500027	2.67E-05	3
0.5	0.9	100	0.535674	0.499974	0.500013	1.33E-05	3
1	0.3	50	0.999801	1.000481	0.998359	1.99E-04	1
1	0.3	75	0.999934	1.00004	0.999986	1.43E-05	3
1	0.3	100	0.999982	0.999998	1.000028	1.93E-06	2
1	0.6	50	1.419507	1.000165	0.999109	1.65E-04	2
1	0.6	75	1.042923	0.99999	1.000223	9.61E-06	2
1	0.6	100	1.103148	1.00002	0.999997	2.31E-07	2
1	0.9	50	1.211894	1.001656	1.001579	0.001579	3
1	0.9	75	1.041806	1.000133	0.999829	1.33E-04	2
1	0.9	100	1.072302	0.999993	1.000004	1.07E-07	3
1.5	0.3	50	1.499797	1.497603	1.498686	2.03E-04	1
1.5	0.3	75	1.500207	1.500195	1.499912	8.83E-05	3
1.5	0.3	100	1.499998	1.50001	1.499998	1.80E-06	1
1.5	0.6	50	1.772489	1.501783	1.500591	5.91E-04	3
1.5	0.6	75	1.57779	1.500043	1.49985	4.29E-05	2
1.5	0.6	100	1.600765	1.499999	1.499988	1.16E-06	2
1.5	0.9	50	1.73851	1.497833	1.498886	1.11E-03	3
1.5	0.9	75	1.698424	1.49974	1.499884	1.16E-04	3
1.5	0.9	100	1.690641	1.499992	1.499996	3.66E-06	3

From the previous table and according to the estimates of the first parameter, the best estimation method is (SH2) with (48%)

Table (2) second parameter estimators and the best method to estimate it

α	β	n	$\hat{\beta}_{MLE}$	$\hat{\beta}_{SH1}$	$\hat{\beta}_{SH2}$	δ_{β}	<i>best</i>
0.5	0.3	50	0.311022	0.300751	0.298689	7.51E-04	2
0.5	0.3	75	0.297503	0.300181	0.299976	2.45E-05	3
0.5	0.3	100	0.291682	0.300006	0.300015	6.05E-06	2
0.5	0.6	50	0.590871	0.599306	0.598154	0.290871	1
0.5	0.6	75	0.582855	0.599948	0.599965	0.282855	1
0.5	0.6	100	0.599396	0.599998	0.600007	0.299396	1
0.5	0.9	50	0.903305	0.900197	0.901802	0.600197	2
0.5	0.9	75	0.923548	0.90008	0.900017	0.600017	3
0.5	0.9	100	0.901504	0.899992	0.899994	0.599992	2
1	0.3	50	0.308574	0.29749	0.299302	0.291426	1
1	0.3	75	0.285951	0.299786	0.300158	0.299842	3
1	0.3	100	0.293517	0.299981	0.299979	0.300019	2
1	0.6	50	0.637851	0.60111	0.600548	5.48E-04	3
1	0.6	75	0.598435	0.600105	0.600044	4.43E-05	3
1	0.6	100	0.587017	0.600003	0.600024	2.82E-06	2
1	0.9	50	0.886351	0.900066	0.901602	0.286351	1
1	0.9	75	0.923318	0.899935	0.900201	0.299935	2
1	0.9	100	0.924103	0.90002	0.899999	0.299999	3
1.5	0.3	50	0.3076	0.299334	0.300841	0.5924	1
1.5	0.3	75	0.302792	0.300042	0.299727	0.597208	1

1.5	0.3	100	0.310106	0.299989	0.29998	0.589894	1
1.5	0.6	50	0.580539	0.600028	0.597834	0.299972	2
1.5	0.6	75	0.593113	0.599843	0.599754	0.300157	2
1.5	0.6	100	0.590893	0.599994	0.599988	0.300006	2
1.5	0.9	50	0.98943	0.899383	0.900427	4.27E-04	3
1.5	0.9	75	0.882366	0.899859	0.900389	1.41E-04	2
1.5	0.9	100	0.895338	0.900016	0.900009	9.05E-06	3

From the previous table and according to the estimates of the second parameter, the best estimation method is (SH1) with (41%)

Table (3) mean square error for first parameter estimators and the best method

α	β	n	$MSE\alpha_{MLE}$	$MSE\alpha_{SH1}$	$MSE\alpha_{SH2}$	MIN_{α}	<i>best</i>
0.5	0.3	50	1.80E-05	1.74E-05	8.41E-06	8.41E-06	3
0.5	0.3	75	1.35E-07	1.64E-07	7.36E-08	7.36E-08	3
0.5	0.3	100	1.45E-09	1.14E-09	1.89E-09	1.14E-09	2
0.5	0.6	50	3.44E-02	1.88E-05	1.64E-05	1.64E-05	3
0.5	0.6	75	1.45E-02	1.28E-07	2.66E-07	1.28E-07	2
0.5	0.6	100	4.17E-03	2.26E-09	1.62E-09	1.62E-09	3
0.5	0.9	50	6.04E-03	1.42E-05	1.25E-05	1.25E-05	3
0.5	0.9	75	1.27E-03	1.86E-07	1.34E-07	1.34E-07	3
0.5	0.9	100	5.97E-03	2.76E-09	2.56E-09	2.56E-09	3
1	0.3	50	2.42E-05	1.45E-05	1.31E-05	1.31E-05	3
1	0.3	75	1.47E-07	6.02E-08	7.17E-08	6.02E-08	2
1	0.3	100	5.41E-09	1.23E-09	3.17E-09	1.23E-09	2
1	0.6	50	0.596561	2.02E-05	5.34E-06	5.34E-06	3
1	0.6	75	7.18E-03	8.45E-08	2.27E-07	8.45E-08	2
1	0.6	100	2.24E-02	1.13E-09	1.92E-09	1.13E-09	2
1	0.9	50	0.168349	2.43E-05	1.10E-05	1.10E-05	3
1	0.9	75	1.54E-02	2.09E-07	1.20E-07	1.20E-07	3
1	0.9	100	2.40E-02	1.16E-09	2.28E-09	1.16E-09	2
1.5	0.3	50	9.74E-06	2.36E-05	1.19E-05	9.74E-06	1
1.5	0.3	75	1.64E-07	1.15E-07	1.92E-07	1.15E-07	2
1.5	0.3	100	2.33E-09	1.68E-09	2.38E-09	1.68E-09	2
1.5	0.6	50	0.307735	1.59E-05	1.94E-05	1.59E-05	2
1.5	0.6	75	3.37E-02	2.49E-07	1.66E-07	1.66E-07	3
1.5	0.6	100	0.036037	1.68E-09	2.12E-09	1.68E-09	2
1.5	0.9	50	0.145603	1.81E-05	1.85E-05	1.81E-05	2
1.5	0.9	75	8.87E-02	2.59E-07	1.38E-07	1.38E-07	3
1.5	0.9	100	6.63E-02	1.60E-09	1.63E-09	1.60E-09	2

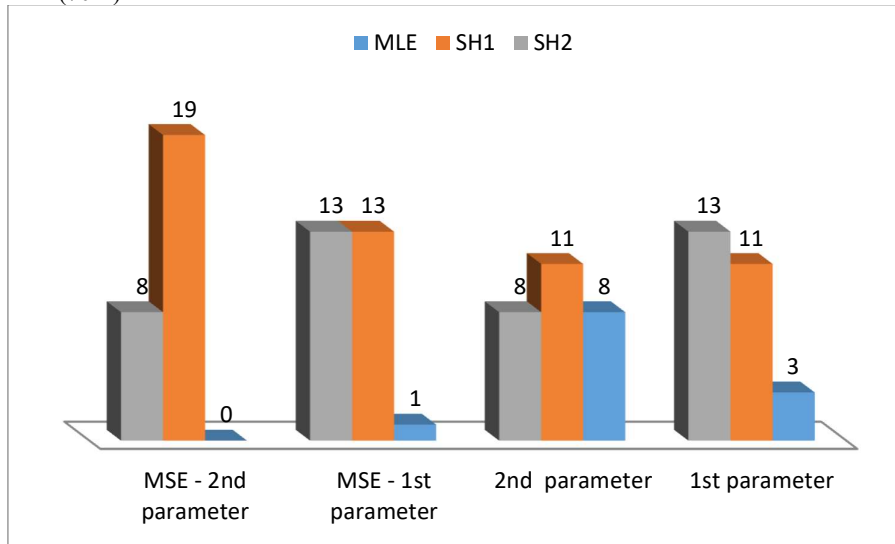
From the previous table and according to the estimates of the first parameter, the best estimation methods are (SH1 & SH2) with (48%) for each one of them

Table (4) mean square error for second parameter estimators and the best method

α	β	n	$MSE\beta_{MLE}$	$MSE\beta_{SH1}$	$MSE\beta_{SH2}$	MIN_{β}	<i>best</i>
0.5	0.3	50	7.48E-04	1.47E-05	1.22E-05	1.22E-05	3
0.5	0.3	75	1.25E-03	2.76E-07	2.34E-07	2.34E-07	3
0.5	0.3	100	7.27E-04	2.04E-09	1.42E-09	1.42E-09	3
0.5	0.6	50	4.62E-03	8.14E-06	2.16E-05	8.14E-06	2
0.5	0.6	75	7.92E-03	8.31E-08	2.49E-07	8.31E-08	2
0.5	0.6	100	5.59E-04	1.03E-09	2.49E-09	1.03E-09	2
0.5	0.9	50	3.17E-02	9.77E-06	2.08E-05	9.77E-06	2
0.5	0.9	75	6.80E-03	2.92E-07	2.10E-07	2.10E-07	3
0.5	0.9	100	1.17E-02	1.76E-09	2.65E-09	1.76E-09	2
1	0.3	50	2.09E-03	3.07E-05	4.15E-06	4.15E-06	3
1	0.3	75	6.40E-04	1.02E-07	2.15E-07	1.02E-07	2
1	0.3	100	6.27E-04	1.50E-09	1.95E-09	1.50E-09	2
1	0.6	50	1.96E-02	1.80E-05	2.11E-05	1.80E-05	2
1	0.6	75	3.22E-03	1.12E-07	2.50E-07	1.12E-07	2
1	0.6	100	5.07E-03	1.87E-09	1.80E-09	1.80E-09	3

1	0.9	50	7.20E-03	1.60E-05	1.53E-05	1.53E-05	3
1	0.9	75	4.27E-03	1.04E-07	1.79E-07	1.04E-07	2
1	0.9	100	4.02E-03	1.77E-09	1.91E-09	1.77E-09	2
1.5	0.3	50	9.42E-04	1.76E-05	2.00E-05	1.76E-05	2
1.5	0.3	75	5.51E-04	9.30E-08	1.87E-07	9.30E-08	2
1.5	0.3	100	1.01E-03	1.55E-09	1.88E-09	1.55E-09	2
1.5	0.6	50	6.29E-03	1.82E-05	2.39E-05	1.82E-05	2
1.5	0.6	75	4.60E-03	1.51E-07	1.98E-07	1.51E-07	2
1.5	0.6	100	4.77E-03	1.31E-09	2.08E-09	1.31E-09	2
1.5	0.9	50	1.91E-02	7.23E-06	2.47E-05	7.23E-06	2
1.5	0.9	75	5.50E-03	1.41E-07	3.00E-07	1.41E-07	2
1.5	0.9	100	4.02E-03	2.53E-09	1.06E-09	1.06E-09	3

From the previous table and according to the estimates of the first parameter, the best estimation method is (SH1) with (70%)



The Figure (3) represents the number of times each method is distinguished

For all experiments the best method was (SH1) by using Minimum absolute deference and Mean square error criteria.

9-Conclusions and suggestions

Simulation results showed that shrinkage Bayesian estimators gives best results comparing with maximum likelihood estimation method for different sample size and initial values of the parameter of log normal distribution

Simulation results indicated that the estimation technique was influenced by both the sample size and the initial values of the distribution parameters. Additionally, the superiority of the (SH1) method was demonstrated in the majority of simulation experiments. Alternative estimation methods for different distributions (Weibull, Gamma) could be utilized, and their results could be contrasted with (moment method, modified moment).

References

- 1 Data from Two Burr-XII Populations. *arXiv e-prints*, arXiv: 2401.03081.
- 2 Bargoshadi, S. A., & Bevrani, H. (2024). Shrinkage Estimation and Prediction for Joint Type-II Censored Data from Two Burr-XII Populations. *arXiv preprint arXiv:2401.03081*.
- 3 Dang, C., Valdebenito, M. A., Wei, P., Song, J., & Beer, M. (2024). Bayesian active learning line sampling with log-normal process for rare-event probability estimation. *Reliability Engineering & System Safety*, 110053.
- 4 De Jager, J. (2022). *Why Subsurface Uncertainties don't have to be Lognormal and other Practices to Avoid*. Paper presented at the Asia Petroleum Geoscience Conference and Exhibition ().
- 5 Hamura, Y. (2023). Bayesian shrinkage estimation for stratified count data. *Japanese Journal of Statistics and Data Science*, 1-23.
- 6 Hodson, T. O. (2022). Root mean square error (RMSE) or mean absolute error (MAE): When to use them or not. *Geoscientific Model Development Discussions*, 2022, 1-10.
- 7 McGrath, S., Zhao, X., Steele, R., Thombs, B. D., Benedetti, A., & Collaboration, D. S. D. (2020). Estimating the sample mean and standard deviation from commonly reported quantiles in meta-analysis. *Statistical methods in medical research*, 29(9), 2520-2537.

- 8 Shi, X., Wang, X.-S., & Wong, A. (2022). Explicit Gaussian Variational Approximation for the Poisson Lognormal Mixed Model. *Mathematics*, 10(23), 4542.
- 9 Yin, M., Iannelli, A., & Smith, R. S. (2021). Maximum likelihood estimation in data-driven modeling and control. *IEEE Transactions on Automatic Control*, 68(1), 317-328.
- 10 Abd Aliwie AN. A Pragmatic Analysis of Wish Strategies Used by Iraqi EFL Learners. *Salud, Ciencia y Tecnología - Serie de Conferencias* [Internet]. 2024 Aug. 12 [cited 2024 Sep. 6];3:..1151. Available from: <https://conferencias.ageditor.ar/index.php/sctconf/article/view/1151>