

**BAYESIAN INFERENCE ON GOMPERTZ-LINDLEY
DISTRIBUTION BASED ON
DIFFERENT LOSS FUNCTIONS**

R. A. Al-Jarallah^{1,§}, M. E. Ghitany¹, D. Kundu²

¹ Department of Statistics and Operations Research

Faculty of Science, Kuwait University, KUWAIT

² Department of Mathematics and Statistics

Indian Institute of Technology Kanpur, Pin 208016, INDIA

Abstract

In the present paper, we consider Bayesian estimation of the parameters and reliability function of the Gompertz Lindley distribution based on three different loss functions. Bayesian estimators are obtained by using different priors under the symmetric squared error and asymmetric linear exponential and general entropy loss functions. Approximate Bayes estimators are computed using Markov chain Monte Carlo (MCMC) methods. These estimators are compared by using bias and mean squared error through simulation study. Finally, a real life data application is reported as an illustration.

MSC 2020: 62F10, 62F15

Key Words and Phrases: Gompertz-Lindley distribution; loss function; Bayesian estimation; Markov chain Monte Carlo (MCMC)

1. Introduction

The Gompertz-Lindley (GL) distribution with parameters α and λ was first introduced by [8]. This distribution has a probability density function

(PDF)

$$f(x) = \frac{\alpha^2 \lambda}{\alpha + 1} \frac{e^{\lambda x} (e^{\lambda x} + \alpha + 1)}{(e^{\lambda x} + \alpha - 1)^3}, \quad x > 0, \quad \alpha, \lambda > 0, \quad (1)$$

reliability function (RF)

$$\begin{aligned} R(x) &= P(X > x) \\ &= \frac{\alpha^2}{\alpha + 1} \frac{e^{\lambda x} + \alpha}{(e^{\lambda x} + \alpha - 1)^2}, \quad x > 0, \quad \alpha, \lambda > 0, \end{aligned} \quad (2)$$

and hazard rate function (HRF)

$$\begin{aligned} h(x) &= \frac{f(x)}{R(x)} \\ &= \frac{\lambda e^{\lambda x} (e^{\lambda x} + \alpha + 1)}{(e^{\lambda x} + \alpha - 1)(e^{\lambda x} + \alpha)}, \quad x > 0, \quad \alpha, \lambda > 0. \end{aligned} \quad (3)$$

Figure 1 shows that the PDF of the GL distribution can be decreasing or unimodal. This figure shows also that the RF of the GL distribution increases as α increases.

Figure 2 shows that the HRF of the GL distribution can be decreasing, increasing or unimodal. That is, the GL distribution has a flexible HRF than well known two-parameter distributions in the literature.

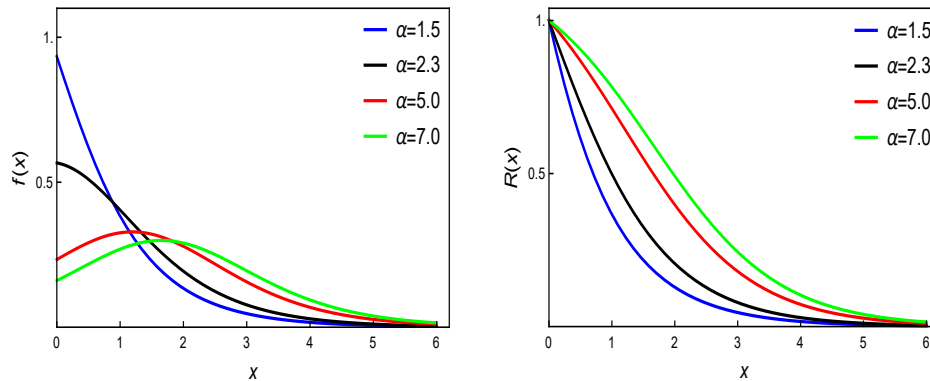


FIGURE 1. PDF and RF of the GL distribution for selected values of α and $\lambda = 1$.

The authors in [8] demonstrated that the GL distribution can be effectively used in reliability applications than the gamma and Weibull distributions.

Recently, [1] investigated several frequentist estimation methods of the GL distribution. In this paper, we consider Bayesian estimation of the GL distribution based on different loss functions.

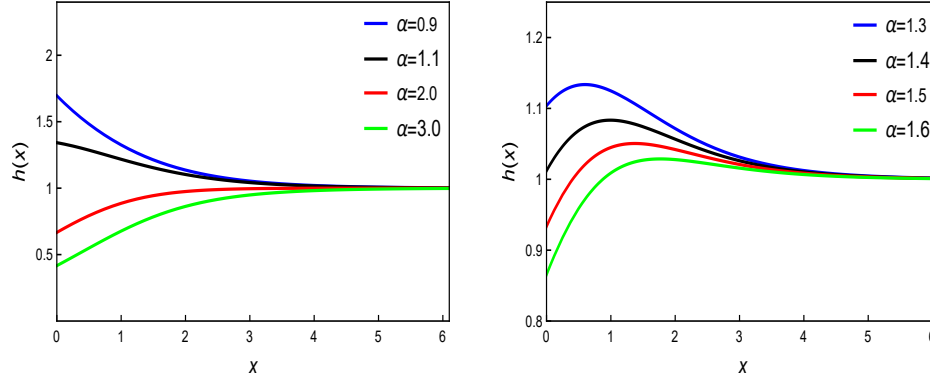


FIGURE 2. HRF of the GL distribution for selected values of α and $\lambda = 1$.

The rest of the paper is organized as follows. In Section 2, Bayesian estimation for model parameters, reliability and hazard rate functions are derived under squared error, linear exponential and general entropy loss functions. Markov chain Monte Carlo (MCMC) methods such as Metropolis-Hastings algorithm has been introduced in Section 3. The Bayesian estimators are compared using Monte Carlo simulation and results are presented in Section 4. Application to real data set is given in Section 5. Finally, we conclude the paper in Section 6.

2. Bayesian Estimation

In general, when conducting a Bayesian analysis, an important aspect is considering a loss function. Loss function represents the loss associated with an error in estimation. In the statistical literature, a number of symmetric and asymmetric loss functions are available. The symmetric loss function is equally penalizes under estimation or overestimation which is the most real life situations. However, in other situations, positive loss maybe more severe than the negative loss and vice-versa requiring asymmetric loss functions.

Here we have considered one symmetric, that is, squared error loss function (SELF) and two asymmetric, that is, linear exponential (LINEX) and general entropy (GELF) loss functions. The mathematical form of these loss functions and their respective Bayes estimators are discussed below.

The most widely used loss function in Bayesian inference is the squared error loss function (SELF), that is defined as

$$L_S(\theta, \hat{\theta}) = (\theta - \hat{\theta})^2, \quad (4)$$

and the Bayes estimator of the parameter θ based on this type of loss function is known as the posterior mean, assuming it exists, and is given by

$$\hat{\theta}_S = E_\theta[\theta|data], \quad (5)$$

where $E_\theta[\cdot|data]$ is the expectation with respect to posterior distribution of θ .

The asymmetric loss function is the linear exponential (LINEX) loss function which was originally introduced by [15]. It penalizes underestimation and overestimation for negative and positive ν , respectively, and is defined as

$$L_L(\theta, \hat{\theta}) = \exp[\nu(\theta - \hat{\theta})] - \nu(\theta - \hat{\theta}) - 1, \quad \nu \neq 0, \quad (6)$$

where the sign of the shape parameter ν reflects the direction of symmetry, and its magnitude reflects the degree of asymmetry. The Bayes estimate of θ relative to LINEX loss function is given by

$$\hat{\theta}_L = -\frac{1}{\nu} \log E_\theta[\exp(-\nu\theta)|data]. \quad (7)$$

Another useful asymmetric loss function is the general entropy loss function (GELF) proposed by [3] as a simple generalization of the entropy loss, and is defined as

$$L_G(\theta, \hat{\theta}) = \left(\frac{\hat{\theta}}{\theta}\right)^w - w \log\left(\frac{\hat{\theta}}{\theta}\right) - 1, \quad w \neq 0, \quad (8)$$

where w is the loss parameter which reflects the departure from symmetry. The corresponding Bayes estimator θ under GELF loss function is given by

$$\hat{\theta}_G = (E_\theta[\theta^{-w}|data])^{-1/w}. \quad (9)$$

Note that, for $w = -1$, the above equation reduced to the Bayes estimator under SELF loss function.

Here, we study the performance of Bayes estimators that depends on the prior distribution and the loss functions that are considered. Many authors have been discussed the performance of the Bayes estimators under different assumptions of prior and loss functions, for example, [5], and [12] have obtained the Bayes estimators under different loss functions for generalized exponential distribution and inverse Gaussian distribution respectively. Moreover, [13] have obtained the Bayes estimator for the parameters of exponentiated gamma distribution under the general entropy loss function.

In the following, Bayesian estimation procedure is developed to estimate the parameters α , λ , the RF $R(t)$, and HRF $h(t)$, for given value $t > 0$, of the GL distribution under different loss functions considered in this paper.

Let $X_1, X_2, \dots, X_n$ be a random sample of size n from GL distribution with PDF given in (1). Let $\mathbf{x} = (x_1, x_2, \dots, x_n)$ be a vector of realization, the

likelihood function as

$$L(\alpha, \lambda; \mathbf{x}) = \frac{\alpha^{2n} \lambda^n}{(\alpha + 1)^n} \prod_{i=1}^n \frac{e^{\lambda x_i} (e^{\lambda x_i} + \alpha + 1)}{(e^{\lambda x_i} + \alpha - 1)^3}. \quad (10)$$

An important aspect to implement Bayesian analysis is the choice of prior distribution that reflects the prior information about the parameters of interest prior to collecting the data, however weakly informative prior could be considered if there is no such information.

Thus, we consider informative prior for each parameter, assuming independent gamma distributions with the following PDFs

$$\pi_1(\alpha; a_1, b_1) = \frac{b_1^{a_1}}{\Gamma(a_1)} \alpha^{a_1-1} e^{-b_1 \alpha}, \quad \alpha > 0, \quad a_1, b_1 > 0, \quad (11)$$

$$\pi_2(\lambda; a_2, b_2) = \frac{b_2^{a_2}}{\Gamma(a_2)} \lambda^{a_2-1} e^{-b_2 \lambda}, \quad \lambda > 0, \quad a_2, b_2 > 0, \quad (12)$$

where the hyper parameters a_1, b_1 (a_2, b_2) are chosen to reflect prior knowledge about α (λ). For the gamma prior to be weakly informative, the hyper parameters are assumed to equal small values such as 0.0001.

Combining the prior distributions given in (11)-(12) and the likelihood function in (10), then using Bayes theorem, the joint posterior PDF of α and λ given the data can be expressed as

$$\begin{aligned} \pi^*(\alpha, \lambda | \mathbf{x}) &= \frac{\alpha^{2n+a_1-1} e^{-b_1 \alpha} \lambda^{n+a_2-1} e^{-(b_2 - \sum_{i=1}^n x_i) \lambda}}{C(\alpha + 1)^n} \\ &\quad \times \prod_{i=1}^n \frac{(e^{\lambda x_i} + \alpha + 1)}{(e^{\lambda x_i} + \alpha - 1)^3}, \end{aligned} \quad (13)$$

where

$$\begin{aligned} C &= \int_0^\infty \int_0^\infty \frac{\alpha^{2n+a_1-1} e^{-b_1 \alpha} \lambda^{n+a_2-1} e^{-(b_2 - \sum_{i=1}^n x_i) \lambda}}{(\alpha + 1)^n} \\ &\quad \times \prod_{i=1}^n \frac{(e^{\lambda x_i} + \alpha + 1)}{(e^{\lambda x_i} + \alpha - 1)^3} d\alpha d\lambda. \end{aligned} \quad (14)$$

The posterior density in (13) involves complex integrals, and hence cannot be evaluated analytically.

In the following, considering the loss function SELF, LINEX and GELF, respectively, we derive the Bayes estimates of $\alpha, \lambda, R(t)$ and $h(t)$.

(i) Bayes estimators under SELF:

$$\begin{aligned} \hat{\alpha}_S &= E[\alpha | \mathbf{x}] \\ &= \int_0^\infty \int_0^\infty \alpha \pi^*(\alpha, \lambda | \mathbf{x}) d\alpha d\lambda, \end{aligned} \quad (15)$$

$$\begin{aligned}\hat{\lambda}_S &= E[\lambda|\mathbf{x}] \\ &= \int_0^\infty \int_0^\infty \lambda \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda,\end{aligned}\quad (16)$$

$$\begin{aligned}\hat{R}_S(t) &= E[R(t; \alpha, \lambda)|\mathbf{x}] \\ &= \int_0^\infty \int_0^\infty R(t; \alpha, \lambda) \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda,\end{aligned}\quad (17)$$

$$\begin{aligned}\hat{h}_S(t) &= E[h(t; \alpha, \lambda)|\mathbf{x}] \\ &= \int_0^\infty \int_0^\infty h(t; \alpha, \lambda) \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda.\end{aligned}\quad (18)$$

(ii) Bayes estimators under LINEX:

$$\begin{aligned}\hat{\alpha}_L &= -\frac{1}{\nu} \log(E[\exp(-\nu\alpha)|\mathbf{x}]) \\ &= -\frac{1}{\nu} \log \left\{ \int_0^\infty \int_0^\infty e^{-\nu\alpha} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\},\end{aligned}\quad (19)$$

$$\begin{aligned}\hat{\lambda}_L &= -\frac{1}{\nu} \log(E[\exp(-\nu\lambda)|\mathbf{x}]) \\ &= -\frac{1}{\nu} \log \left\{ \int_0^\infty \int_0^\infty e^{-\nu\lambda} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\},\end{aligned}\quad (20)$$

$$\begin{aligned}\hat{R}_L(t) &= -\frac{1}{\nu} \log(E[\exp(-\nu R(t; \alpha, \lambda)|\mathbf{x})]) \\ &= -\frac{1}{\nu} \log \left\{ \int_0^\infty \int_0^\infty e^{-\nu R(t; \alpha, \lambda)} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\},\end{aligned}\quad (21)$$

$$\begin{aligned}\hat{h}_L(t) &= -\frac{1}{\nu} \log(E[\exp(-\nu h(t; \alpha, \lambda)|\mathbf{x})]) \\ &= -\frac{1}{\nu} \log \left\{ \int_0^\infty \int_0^\infty e^{-\nu h(t; \alpha, \lambda)} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\}.\end{aligned}\quad (22)$$

(iii) Bayes estimators under GELF:

$$\begin{aligned}\hat{\alpha}_G &= E[\alpha^{-w}|\mathbf{x}]^{-1/w} \\ &= \left\{ \int_0^\infty \int_0^\infty \alpha^{-w} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\}^{-1/w},\end{aligned}\quad (23)$$

$$\begin{aligned}\hat{\lambda}_G &= E[\lambda^{-w}|\mathbf{x}]^{-1/w} \\ &= \left\{ \int_0^\infty \int_0^\infty \lambda^{-w} \pi^*(\alpha, \lambda|\mathbf{x}) d\alpha d\lambda \right\}^{-1/w},\end{aligned}\quad (24)$$

$$\begin{aligned}\hat{R}_G(t) &= E [R(t; \alpha, \lambda)^{-w} | \mathbf{x}]^{-1/w} \\ &= \left\{ \int_0^\infty \int_0^\infty R(t; \alpha, \lambda)^{-w} \pi^*(\alpha, \lambda | \mathbf{x}) d\alpha d\lambda \right\}^{-1/w},\end{aligned}\quad (25)$$

$$\begin{aligned}\hat{h}_G(t) &= E [h(t; \alpha, \lambda)^{-w} | \mathbf{x}]^{-1/w} \\ &= \left\{ \int_0^\infty \int_0^\infty h(t; \alpha, \lambda)^{-w} \pi^*(\alpha, \lambda | \mathbf{x}) d\alpha d\lambda \right\}^{-1/w}.\end{aligned}\quad (26)$$

It is clear that all the above Bayes estimators are involved the double integrals for which simple closed forms cannot be obtained. Therefore, an approximation methods are needed such as Markov chain monte carlo (MCMC) methods, namely Gibbs sampler ([7], [14]) and Metropolis-Hastings (M-H) algorithm ([9], [2]) can be used to simulate from the posterior density and the sample based inference can be performed.

3. Markov Chain Monte Carlo

In this section, we apply the Markov chain monte carlo (MCMC) method to obtain the approximate Bayes estimators of $\alpha, \lambda, R(t)$, and $h(t)$ under the three loss functions considered in this paper.

The MCMC method is a general simulation method that is used to sample posterior distribution and compute posterior quantities of interest. This method is usually applied to solve integration and optimization problems in large dimensional spaces, however, in most cases, the integration does not have a closed structure. For more details of the MCMC methods, see [7]. By applying MCMC we can generate the random sample of unknown quantities by using the posterior densities. The generated sample is then used to obtain the Bayes estimators with respect to the proposed loss functions.

Integrating joint posterior $\pi^*(\alpha, \lambda | \mathbf{x})$ in (13) with respect to λ and α respectively, we obtain the marginal posterior densities of α and λ as

$$\pi_1^*(\alpha | \lambda, \mathbf{x}) \propto \frac{\alpha^{2n+a_1-1} e^{-b_1 \alpha} \lambda^n}{(\alpha + 1)^n} \prod_{i=1}^n \frac{e^{\lambda x_i} (e^{\lambda x_i} + \alpha + 1)}{(e^{\lambda x_i} + \alpha - 1)^3}, \quad (27)$$

$$\pi_2^*(\lambda | \alpha, \mathbf{x}) \propto \lambda^n \prod_{i=1}^n \frac{e^{\lambda x_i} (e^{\lambda x_i} + \alpha + 1)}{(e^{\lambda x_i} + \alpha - 1)^3}. \quad (28)$$

It is clear from (15) and (16) that the marginal posterior densities of α and λ are not in a closed form. Therefore, we consider the Metropolis Hastings (M-H) algorithm. This algorithm is one of the MCMC methods that first created by [10] and later extended by [9] and had many applications.

The M-H algorithm allows for sampling from an analytic target complicated distributions by filtering samples from a proposal distribution. Here, we adopt the M-H algorithm with normal distribution as a proposal distribution, to generate samples of α and λ from (15) and (16) (see [6]).

The Gibbs sampler algorithm proposed is summarized as follows:

- Step 1: Start with selecting an initial values of the parameters $\alpha^{(0)}, \lambda^{(0)}$ and set $j = 1$.
- Step 2 : Generate $\alpha^{(j)}$ from $\pi_1^*(\alpha|\lambda^{(j-1)}, \mathbf{x})$ by M-H considering the normal proposal distribution $q(\alpha) = I(\alpha > 0)N(\alpha^{(j-1)}, 1)$. The proposed value is accepted with probability

$$\rho(\alpha^{(j-1)}, \alpha^{(j)}) = \min \left\{ 1, \frac{\pi(\alpha^{(j)})q(\alpha^{(j-1)})}{\pi^*(\alpha^{(j-1)})q(\alpha^{(j)})} \right\},$$

or accept $\alpha^{(j-1)}$ with probability $1 - \rho(\alpha^{(j-1)}, \alpha^{(j)})$.

- Step 3: Generate $\lambda^{(j)}$ from $\pi_2^*(\lambda|\alpha^{(j)}, \mathbf{x})$ by M-H with the normal proposal distribution $q(\lambda) = I(\lambda > 0)N(\lambda^{(j-1)}, 1)$. The proposal value is accepted with probability

$$\rho(\lambda^{(j-1)}, \lambda^{(j)}) = \min \left\{ 1, \frac{\pi(\lambda^{(j)})q(\lambda^{(j-1)})}{\pi^*(\lambda^{(j-1)})q(\lambda^{(j)})} \right\},$$

or accept $\lambda^{(j-1)}$ with probability $1 - \rho(\lambda^{(j-1)}, \lambda^{(j)})$.

- Step 4: Compute, for given value t , $R(t; \alpha^{(j)}, \lambda^{(j)})$ using equation (2) and $h(t; \alpha^{(j)}, \lambda^{(j)})$ using equation (3) and set $j = j + 1$.
- Step 5: Repeat steps 2-4 for large number of iterations, say T , until convergence is assured and then obtain $\alpha^{(j)}, \lambda^{(j)}, R(t; \alpha^{(j)}, \lambda^{(j)})$, and $h(t; \alpha^{(j)}, \lambda^{(j)})$ for $j = 1, \dots, T$.

Samples of $\alpha, \lambda, R(t)$, and $h(t)$, for given t , are obtained from the resulted posterior samples. In the following we derive the Bayes estimates of $\alpha, \lambda, R(t)$ and $h(t)$ under the considered loss functions.

(i) MCMC estimators under SELF:

$$\hat{\alpha}_S = \frac{1}{T} \sum_{j=1}^T \alpha^{(j)} \quad (29)$$

$$\hat{\lambda}_S = \frac{1}{T} \sum_{j=1}^T \lambda^{(j)}, \quad (30)$$

$$\hat{R}_S(t) = \frac{1}{T} \sum_{j=1}^T R(t; \alpha^{(j)}, \lambda^{(j)}), \quad (31)$$

$$\hat{h}_S(t) = \frac{1}{T} \sum_{j=1}^T h(t; \alpha^{(j)}, \lambda^{(j)}). \quad (32)$$

(ii) MCMC estimators under LINEX:

$$\hat{\alpha}_L = -\frac{1}{\nu} \log \left(\frac{1}{T} \sum_{j=1}^T e^{-\nu \alpha^{(j)}} \right), \quad (33)$$

$$\hat{\lambda}_L = -\frac{1}{\nu} \log \left(\frac{1}{T} \sum_{j=1}^T e^{-\nu \lambda^{(j)}} \right), \quad (34)$$

$$\hat{R}_L(t) = -\frac{1}{\nu} \log \left(\frac{1}{T} \sum_{j=1}^T e^{-\nu R(t; \alpha^{(j)}, \lambda^{(j)})} \right), \quad (35)$$

$$\hat{h}_L(t) = -\frac{1}{\nu} \log \left(\frac{1}{T} \sum_{j=1}^T e^{-\nu h(t; \alpha^{(j)}, \lambda^{(j)})} \right). \quad (36)$$

(iii) MCMC estimators under GELF:

$$\hat{\alpha}_G = \left(\frac{1}{T} \sum_{j=1}^T \left(\frac{1}{\alpha^{(j)}} \right)^w \right)^{-1/w}, \quad (37)$$

$$\hat{\lambda}_G = \left(\frac{1}{T} \sum_{j=1}^T \left(\frac{1}{\lambda^{(j)}} \right)^w \right)^{-1/w}, \quad (38)$$

$$\hat{R}_G(t) = \left(\frac{1}{T} \sum_{j=1}^T \left(\frac{1}{R(t; \alpha^{(j)}, \lambda^{(j)})} \right)^w \right)^{-1/w}, \quad (39)$$

$$\hat{h}_G(t) = \left(\frac{1}{T} \sum_{j=1}^T \left(\frac{1}{h(t; \alpha^{(j)}, \lambda^{(j)})} \right)^w \right)^{-1/w}. \quad (40)$$

4. Simulation Study

This section presents the simulation study which is carried out to assess the performance of the various derived Bayes estimators under the three proposed loss functions. The MCMC method have been applied for Bayesian analysis so that sample based inference is carried out. All the computational algorithms are performed using R Software [11]. We mainly compare the performance of these estimates in terms of the average bias and mean square error (MSE). The bias and MSE of an estimate $\hat{\theta}$ of θ is given respectively as

$$\text{Bias}(\hat{\theta}) = \frac{1}{T} \sum_{i=1}^T (\hat{\theta}_i - \theta_i),$$

$$\text{MSE}(\hat{\theta}) = \frac{1}{T} \sum_{i=1}^T (\hat{\theta}_i - \theta_i)^2,$$

where $T = 50,000$ is the number of iterations in the simulation process. In this simulation, we considered $(\alpha, \lambda) = (2, 1)$ and $t = 1$. Also, we considered different sample sizes, $n = 20, 40, 60, 80, 100$, to represent small, moderate and large sample sizes. We also considered two different choices of $\nu, w = -0.5, 1.5$ for both LINEX and GELF loss functions, respectively.

For conducting Bayesian analysis two configurations of priors are considered. The first prior is a non-informative while the second is informative having the following sets of hyper-parameters:

$$\text{Prior I : } a_1 = b_1 = a_2 = b_2 = 0.0001,$$

$$\text{Prior II : } a_1 = 4, b_1 = 2, a_2 = 2, b_2 = 2.$$

In case of informative prior, the hyper parameters are chosen in such away that the prior mean equals to the true value of the parameter with varying prior variance. The prior variance varies from large, moderate to small reflecting the confidence of our prior guess.

In each case, a random sample of size n is generated from GL distribution and the Bayes estimators of the unknown parameters are obtained by applying the M-H algorithm provided in Section 3. We generated 50,000 realization of the parameters $\alpha, \lambda, R(1)$ and $h(1)$ from posterior densities in (15) and (16) using M-H with 5000 burn in period to reduce the dependence of the starting values. The MCMC run shows fine mixing of the chain. For reducing the autocorrelation among the generated values of α and λ , we only recored every 5-th generated values of each parameter. The convergence of MCMC sample is checked, and it was found that the Markov chain converges rapidly with any arbitrary initial starting values.

Tables 1 - 2 show the performance of the proposed Bayes estimators of the model parameters $\alpha, \lambda, R(1)$ and $h(1)$ for different values of n, ν and w , under Prior I and Prior II, respectively.

These tables show that the bias of the Bayes estimators under the proposed loss functions are small and can be negative or positive. Also, MSE of Bayes estimators under different loss functions decreases as the sample size increases. Finally, Bayes estimator of α under SELF has smaller MSE under LINEX and GELF for both priors.

5. Illustrative Example

In this section, we present the analysis of a real data set to illustrate the performance of the proposed Bayes estimators. The data set was reported by [4] and represents 107 failure times (in hours) for right rear brakes on D9G-66A caterpillar tractors.

56, 753, 1153, 1586, 2150, 2624, 3826, 83, 763, 1154, 1599, 2156, 2675, 3995, 104, 806, 1193, 1608, 2160, 2701, 4007, 116, 834, 1201, 1723, 2190, 2755, 4159, 244, 838, 1253, 1769, 2210, 2877, 4300, 305, 862, 1313, 1795, 2220, 2879, 4487, 429, 897, 1329, 1927, 2248, 2922, 5074, 452, 904, 1347, 1957, 2285, 2986, 5579, 453, 981, 1454, 2005, 2325, 3092, 5623, 503, 1007, 1464, 2010, 2337, 3160, 6869, 552, 1008, 1490, 2016, 2351, 3185, 7739, 614, 1049, 1491, 2022, 2437, 3191, 661, 1069, 1532, 2037, 2454, 3439, 673, 1107, 1549, 2065, 2546, 3617, 683, 1125, 1568, 2096, 2565, 3685, 685, 1141, 1574, 2139, 2584, 3756

Summary of descriptive statistics of this data set is shown in Table 3.

The MLE estimates, and the Bayes estimators of $\alpha, \lambda, R(1)$ and $h(1)$ under SELF, LINEX and GELF loss functions are presented in Table 4. For this data set, Bayesian analysis is carried out in case of non-informative prior, that is Prior I : $a_1 = b_1 = a_2 = b_2 = 0.0001$, since we do not have any prior information.

By applying Metropolis-Hastings algorithm with normal proposal distribution, we generated 50,000 random samples each of size 107 and discarded the initial 5000 samples as a burn-in. The Bayes estimators are computed based on the remaining observations.

TABLE 1. MSE and Bias (in parenthesis) of Bayesian estimates based on non-informative Prior I when $(\alpha, \lambda) = (2, 1)$.

n		SELF	LINEX		GELF	
			$v = -0.5$	$v = 1.5$	$w = -0.5$	$w = 1.5$
20	α	0.001659 (0.018563)	0.029721 (-0.166501)	0.257839 (0.503411)	0.015901 (0.126334)	0.393111 (0.527721)
	λ	0.017113 (-0.132388)	0.038887 (-0.191177)	0.002977 (0.051776)	0.004133 (-0.069881)	0.059877 (0.2229001)
	$R(1)$	0.001877 (0.044771)	0.001526 (0.032771)	0.009444 (0.078661)	0.005433 (0.077119)	0.055255 (0.234450)
	$h(1)$	0.015934 (-0.126231)	0.031133 (-0.176447)	0.000442 (0.021039)	0.003141 (-0.056045)	0.024083 (-0.155188)
40	α	0.000979 (0.014982)	0.028977 (-0.167334)	0.251821 (0.501496)	0.015122 (0.1200314)	0.342420 (0.589822)
	λ	0.0175422 (-0.133281)	0.0381343 (-0.189117)	0.002811 (0.051221)	0.004565 (-0.067822)	0.051744 (0.223441)
	$R(1)$	0.001077 (0.035519)	0.001255 (0.025517)	0.006198 (0.087895)	0.005512 (0.067899)	0.052296 (0.251118)
	$h(1)$	0.015925 (-0.126355)	0.031016 (-0.176115)	0.000382 (0.019556)	0.003091 (-0.055597)	0.020703 (0.020703)
60	α	0.000773 (0.013441)	0.029654 (-0.166771)	0.247854 (0.489979)	0.014544 (0.117440)	0.337003 (0.571121)
	λ	0.018112 (-0.127660)	0.037665 (-0.193441)	0.003219 (0.051455)	0.014433 (-0.067231)	0.051987 (0.230113)
	$R(1)$	0.001977 (0.042551)	0.000877 (0.031226)	0.005126 (0.078301)	0.005541 (0.074331)	0.051776 (0.245572)
	$h(1)$	0.015921 (-0.126180)	0.030462 (-0.174535)	0.000356 (0.018883)	0.011569 (-0.107560)	0.021037 (-0.145044)
80	α	0.000452 (0.006991)	0.026544 (-0.175402)	0.241338 (0.487550)	0.014781 (0.118909)	0.343031 (0.589110)
	λ	0.018119 (-0.005881)	0.037119 (-0.087221)	0.002187 (0.228777)	0.003661 (0.105779)	0.051228 (0.899760)
	$R(1)$	0.001233 (-0.412287)	0.000897 (-0.458601)	0.004977 (-0.241781)	0.005692 (-0.346327)	0.0517748 (-0.085498)
	$h(1)$	0.015873 (-0.125991)	0.030213 (-0.173820)	0.000336 (0.018346)	0.003011 (-0.054874)	0.0209241 (-0.144651)
100	α	0.000766 (0.007211)	0.022883 (-0.175543)	0.242190 (0.491887)	0.013790 (0.116101)	0.327001 (0.571322)
	λ	0.017088 (-0.302117)	0.038101 (-0.193881)	0.002577 (0.052188)	0.005421 (-0.0621188)	0.050198 (0.226558)
	$R(1)$	0.001981 (0.040113)	0.000799 (0.027332)	0.005302 (0.077916)	0.005788 (0.075447)	0.0524331 (0.224087)
	$h(1)$	0.015568 (-0.124773)	0.030219 (-0.174411)	0.000361 (0.019000)	0.002960 (-0.054405)	0.020921 (-0.144641)

TABLE 2. MSE and Bias (in parenthesis) of Bayesian estimates based on informative Prior II when $(\alpha, \lambda) = (2, 1)$.

n		SELF	LINEX		GELF	
			$\nu = -0.5$	$\nu = 1.5$	$w = -0.5$	$w = 1.5$
20	α	0.000811	0.037822	0.146658	0.009113	0.177299
		(0.000211)	(-0.187712)	(0.279009)	(0.087299)	(0.371331)
	λ	0.000766	0.004792	0.026504	0.006711	0.103801
		(0.019559)	(-0.045547)	(0.164474)	(0.077622)	(0.333281)
	$R(1)$	0.000265	0.000733	0.000344	0.000311	0.032119
		(-0.016678)	(-0.024224)	(0.014767)	(0.010331)	(0.179960)
	$h(1)$	0.000709	0.002385	0.014346	0.005182	0.075405
		(-0.026643)	(-0.048843)	(-0.119775)	(-0.071988)	(-0.274600)
40	α	0.000650	0.031776	0.143288	0.006991	0.174110
		(0.001771)	(-0.181771)	(0.375111)	(0.081655)	(0.412899)
	λ	0.000433	0.004661	0.024114	0.004855	0.099166
		(0.007221)	(-0.059144)	(0.153677)	(0.066821)	(0.315117)
	$R(1)$	0.000211	0.000455	0.000410	0.000344	0.032773
		(-0.012201)	(-0.022788)	(0.019144)	(0.014188)	(0.171335)
	$h(1)$	0.000391	0.002472	0.012325	0.007340	0.072125
		(-0.019793)	(-0.049724)	(-0.111020)	(-0.061107)	(-0.268561)
60	α	0.000601	0.033554	0.143177	0.007311	0.173120
		(0.002766)	(-0.179332)	(0.377885)	(0.083455)	(0.415101)
	λ	0.000344	0.004866	0.023166	0.004122	0.098765
		(0.002011)	(-0.065122)	(0.151665)	(0.062111)	(0.31233)
	$R(1)$	0.000161	0.000478	0.000411	0.000344	0.031877
		(-0.010055)	(-0.020721)	(0.021358)	(0.016570)	(0.175222)
	$h(1)$	0.000274	0.002420	0.011407	0.003080	0.071537
		(-0.016566)	(-0.049194)	(-0.106807)	(-0.055505)	(-0.267465)
80	α	0.000455	0.032655	0.141280	0.007187	0.171122
		(0.000299)	(-0.183422)	(0.375237)	(0.080933)	(0.412899)
	λ	0.000346	0.004677	0.022899	0.004213	0.099874
		(0.000877)	(-0.067221)	(0.15012)	(0.061890)	(0.313542)
	$R(1)$	0.000102	0.000398	0.000410	0.000298	0.031158
		(-0.010329)	(-0.021782)	(0.021223)	(0.0168264)	(0.177892)
	$h(1)$	0.000266	0.002408	0.011324	0.002994	0.071661
		(-0.016330)	(-0.049073)	(-0.106418)	(-0.054723)	(-0.267696)
100	α	0.000422	0.032144	0.139521	0.006913	0.173258
		(-0.001860)	(-0.185344)	(0.372538)	(0.080131)	(0.411879)
	λ	0.000320	0.004401	0.023488	0.003564	0.098769
		(-0.001239)	(-0.066761)	(0.149583)	(0.057690)	(0.311767)
	$R(1)$	0.000167	0.000301	0.000367	0.000315	0.0312776
		(-0.008767)	(-0.021341)	(0.021075)	(0.0167331)	(0.181766)
	$h(1)$	0.000248	0.002251	0.011065	0.002775	0.070800
		(-0.015778)	(-0.047451)	(-0.105194)	(-0.211113)	(-0.266083)

Bayesian graphical diagnostics tools involving trace and the autocorrelation function (ACF) plots are used to check the convergence of Metropolis-Hastings algorithm. Figures 3 to 6 show the MCMC trace plots and autocorrelation function (ACF) plots of $\alpha, \lambda, R(1)$ and $h(1)$, respectively. It is clear from the trace plots, for the simulated values of $\alpha, \lambda, R(1)$ and $h(1)$, that we have random scatters about some mean value represented by a solid line with a fine mixing of the chains.

The MLEs and the Bayes estimates of $\alpha, \lambda, R(1)$ and $h(1)$, respectively, are given in Table 4. This table shows that the MLE of α and $\hat{\alpha}_S$ are quite close. The value of $\hat{\alpha}_S$ is more than that of $\hat{\alpha}_L$, and $\hat{\alpha}_G$. However, the value of $\hat{\alpha}_L$ for $\nu = 1.5$ and $\hat{\alpha}_G$ for $w = 1.5$ have smaller values than the corresponding estimates for $\nu = -0.5$ and for $w = -0.5$. A similar trend is noted for the estimates of $R(1)$.

All trace plots in Figures 3 to 6 indicate that the MCMC samples are well mixed and stationary. All ACF plots in Figures 3 to 6 show that the chains have low autocorrelations indicating rapid convergence of the Metropolis-Hastings algorithm.

TABLE 3. Descriptive Statistics for Data set

Min	Q1	Median	Q3	Max	Mean	Standard deviation
56	1018.25	1795	2614	7739	2042.262	1397.770

TABLE 4. MLE and Bayesian Estimates for Data set parameters

<i>Estimate</i>		α	λ	$R(1)$	$h(1)$
<i>MLE</i>		6.109700	0.000100	0.999981	0.000019
<i>SELF</i>		5.992000	0.039930	0.992200	0.007868
<i>LINEX</i>	$\nu = -0.5$	5.985142	0.040023	0.992179	0.007897
	$\nu = 1.5$	5.760748	0.039486	0.992157	0.008122
<i>GELF</i>	$w = -0.5$	5.978675	0.035419	0.992168	0.006971
	$w = 1.5$	5.922888	0.000727	0.992146	0.000141

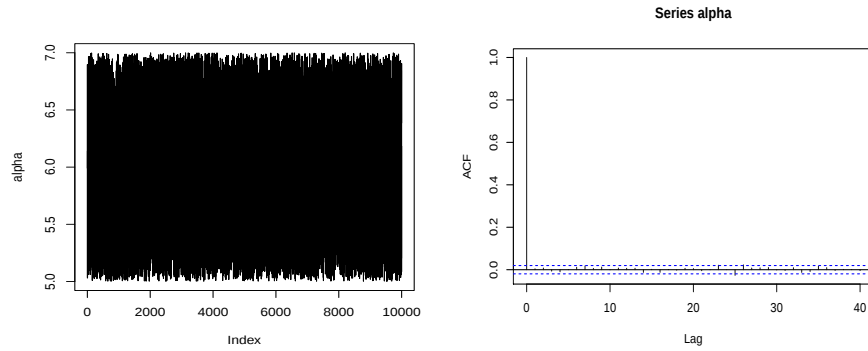


FIGURE 3. MCMC trace plot (left) and ACF plot (right) of simulated posterior samples by Metropolis-Hastings algorithm for α .

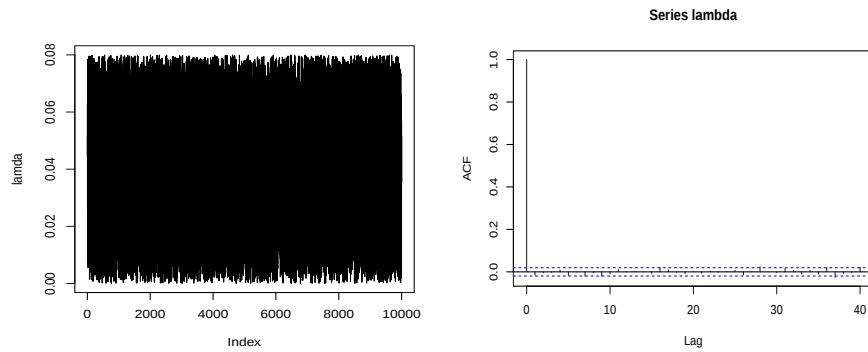


FIGURE 4. MCMC trace plot (left) and ACF plot (right) of simulated posterior samples by Metropolis-Hastings algorithm for λ .

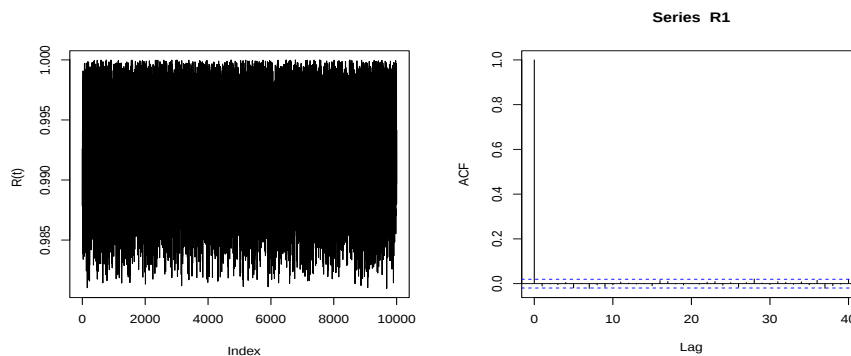


FIGURE 5. MCMC trace plot (left) and ACF plot (right) of simulated posterior samples by Metropolis-Hastings algorithm for the $R(1)$ function.

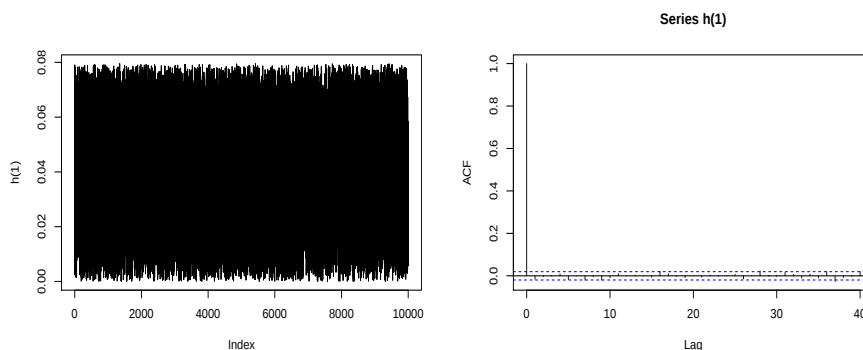


FIGURE 6. MCMC trace plot (left) and ACF plot (right) of simulated posterior samples by Metropolis-Hastings algorithm for the $h(1)$ function.

6. Conclusion

In this paper, we have considered the statistical inference of the Gompertz-Lindley distribution under Bayesian framework. We have estimated the unknown parameters, reliability and hazard rate functions under the squared error, linear, and general entropy loss functions. We obtained the Bayes estimators by applying Markov chain Monte Carlo and Metropolis-Hastings methods. In a simulation study, we compared the performance of the Bayes estimators on the basis of bias and mean-squared error by varying sample sizes and prior distributions. We hope that the Bayesian approach to the Gompertz-Lindley distribution presented here will be found useful for data analysts.

References

- [1] F. Almathkour, A. Alothman, M.E. Ghitany, R.C. Gupta, J. Mazucheli, A Comparative study of various methods of estimation for Gompertz-Lindley distribution. *International Journal of Applied Mathematics*, **35**, No 2 (2022), 347-364; DOI: 10.12732/ijam.v35i2.12.
- [2] S. Brooks, Markov chain Monte Carlo method and its application. *Journal of the Royal Statistical Society: Ser. D (The Statistician)*, **47**, No 1 (1998), 69-100; DOI: 10.1111/1467-9884.00117.
- [3] R. Calabria, G. Pulcini, Point estimation under asymmetric loss functions for left-truncated exponential samples. *Communications in Statistics - Theory and Methods*, **25** (1996), 585-600; DOI: 10.1080/03610929608831715.
- [4] M.N. Chang, P.V. Rao, Improved estimation of survival functions in the new-better-than-used class. *Technometrics*, **35**, No 2 (1993), 192-203.
- [5] S. Dey, Bayesian estimation of the shape parameter of the generalised exponential distribution under different loss functions. *Pakistan Journal of Statistics and Operation Research*, **6**, No 2 (2010), 163-174.
- [6] A. Gelman, J.B. Carlin, H.S. Stern, D.B. Rubin, Bayesian Data Analysis, Chapman and Hall/CRC (1995).
- [7] S. Geman, D. Geman, Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, **6**, No 6 (1984), 721-741; DOI: 10.1109/tpami.1984.4767596.
- [8] M. E. Ghitany, S. M. Aboukhamseen, A. A. Baqer, R. C. Gupta, Gompertz-Lindley distribution and associated inference. *Communications in Statistics - Simulation and Computation*, **51**, No 5 (2022), 2599-2618; DOI: 10.1080/03610918.2019.1699113.
- [9] W.K. Hastings, Monte Carlo sampling methods using Markov chains and their applications, *Biometrika*, **55**, No 1 (1970), 97-109; DOI: 10.1093/biomet/57.1.97.
- [10] N. Metropolis, A.W. Rosenbluth, M.N. Rosenbluth, A.H. Teller, E. Teller, Equation of state calculations by fast computing machines. *The Journal of Chemical Physics*, **21**, No 6 (1953), 1087-1092; DOI: 10.1063/1.1699114.
- [11] R Core Team. R: A language and environment for statistical computing. R Foundation for Statistical Computing. Vienna, Austria (2011); <http://www.r-project.org/index.html>.
- [12] S.K. Singh, U. Singh, D. Kumar, Bayesian estimation of the exponentiated gamma parameters and reliability function under asymmetric loss functions. *RAVSTAT - Statistical Journal*, **9** (2011), 247-260.

- [13] P.K. Singh, S.K. Singh, U. Singh, Bayes estimator of inverse Gaussian parameters under general entropy loss function using Lindley's approximation. *Communications in Statistics - Simulation and Computation*, **37**, No 9 (2008), 1750-1762; DOI: 10.1080/03610910701884054.
- [14] A.F.M. Smith, G.O. Roberts, Bayesian computation via the Gibbs sampler and related Markov chain Monte Carlo methods. *Journal of the Royal Statistical Society: Ser. B (Methodological)*, **55**, No 1 (1993), 3-23; DOI: 10.2307/2346063.
- [15] H. Varian, A Bayesian approach to real estate assessment. In: *Studies in Bayesian Econometrics and Statistics in Honor of Leonard J. Savage*, Stephen E. Fienberg and A. Zellner, Eds. (1975), 195–208, North-Holland Publishing Company, Amsterdam, The Netherlands.