

Article

Bayesian Inference on Stress–Strength Reliability with Geometric Distributions

Mohammed K. Shakhathreh 

Department of Mathematics and Statistics, Faculty of Science, Jordan University of Science and Technology, P.O. Box 3030, Irbid 22110, Jordan; mkshakhathreh6@just.edu.jo

Abstract

This paper investigates the estimation of the stress–strength reliability parameter $\rho = P(X \leq Y)$, where stress (X) and strength (Y) are independently modeled by geometric distributions. Objective Bayesian approaches are employed by developing Jeffreys, reference, and probability-matching priors for ρ , and their effects on the resulting Bayes estimates are examined. Posterior inference is carried out using the random-walk Metropolis–Hastings algorithm. The performance of the proposed Bayesian estimators is assessed through extensive Monte Carlo simulations based on average estimates, root mean squared errors, and frequentist coverage probabilities of the highest posterior density credible intervals. Furthermore, the applicability of the methodology is demonstrated using two real data sets.

Keywords: Jeffreys prior; reference prior; probability matching prior; symmetric proposal; stress–strength reliability

1. Introduction

The estimation of stress–strength models has long been recognized as a fundamental topic in reliability theory. These models play an important role in various applied fields, particularly in engineering, quality assurance, and medical sciences. A significant portion of the statistical literature has focused on the estimation of the reliability parameter $\rho = P(X < Y)$, where X denotes the stress applied to a system, Y represents its strength, and ρ quantifies the probability that the system can withstand the applied stress. It is worth emphasizing that numerous authors have investigated the estimation of the stress–strength reliability parameter ρ using both frequentist and subjective Bayesian approaches across a range of continuous distributions and under various sampling schemes, and this topic continues to attract considerable attention in the statistical literature. For instance, Ref. [1] addressed this problem for the generalized exponential distribution; Ref. [2] studied it for the inverse Pareto distribution under progressively censored data; and Ref. [3] explored it for the exponential distribution based on generalized order statistics. In addition, Ref. [4] examined the Bayesian estimation of ρ for the Lomax distribution under type-II hybrid censoring using an asymmetric loss function. On the other hand, several authors have employed objective Bayesian methods, for instance, Ref. [5] for the Weibull distribution, Ref. [6] for the generalized exponential stress–strength model, and Ref. [7] for the Fréchet stress–strength model. Related investigations based on discrete probability distributions include the works of [8] for Poisson data and [9] for the Poisson-geometric distribution. The geometric distribution has also been studied by [10–12] for complete data and for lower



Academic Editors: Michel Planat and Zhibin Du

Received: 10 September 2025

Revised: 30 September 2025

Accepted: 10 October 2025

Published: 13 October 2025

Citation: Shakhathreh, M.K. Bayesian Inference on Stress–Strength Reliability with Geometric Distributions. *Symmetry* **2025**, *17*, 1723. <https://doi.org/10.3390/sym17101723>

Copyright: © 2025 by the author. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

record data by [13]. However, in these latter studies, the Bayesian methods applied were subjective and limited in scope.

In this paper, we address the problem of estimating the stress–strength reliability parameter ρ when both the stress X and the strength Y are modeled as independent geometric random variables with parameters θ_1 and θ_2 , respectively. Accordingly, the reliability of the stress–strength system is given by

$$\rho = \Pr(X \leq Y) = \sum_{k=1}^{\infty} \sum_{j=k}^{\infty} \theta_1 \theta_2 (1 - \theta_1)^{k-1} (1 - \theta_2)^{j-1} = \frac{\theta_1}{1 - (1 - \theta_1)(1 - \theta_2)}. \quad (1)$$

The importance of studying stress–strength reliability under geometric distributions lies in its wide range of real-world applications, where reliable estimates with meaningful interpretation are required. For instance, in manufacturing quality control, let X denote the number of shocks a component tolerates before breaking and Y the number a stronger reference part withstands. The stress–strength reliability $\rho = \Pr(X \leq Y)$ then quantifies the probability that the component fails first. Since such failures typically occur after repeated shocks with a constant chance of breaking, the geometric distribution provides a natural model. A similar scenario arises in digital communication systems, where X represents the number of packet retransmissions until failure for a weaker channel and Y for a stronger channel. In this case, ρ gives the probability that the weaker channel fails before the stronger one. Because packet transmissions follow repeated Bernoulli trials, the geometric distribution again offers a suitable framework. Finally, and more importantly, two additional real data examples highlighting the relevance of the stress–strength model are presented in the applications section.

Methods for estimating ρ from discrete distributions have primarily relied on classical approaches, with the most common being the maximum likelihood (ML) estimator and subjective Bayesian methods, as clearly demonstrated in the works of [10–12]. While ML estimators have desirable properties, such as invariance and asymptotic efficiency, they are well known to exhibit bias, particularly with small sample sizes. Similar limitations arise when using subjective Bayesian methods, as these often involve multiple hyperparameters that must be elicited. However, the authors in the aforementioned papers employed the subjective Bayesian approach with hyperparameters that were subjectively chosen, rather than explicitly elicited, which can lead to potentially biased results. Specifically, the authors applied independent Beta priors to the parameters of the geometric distribution, i.e., $\theta_1 \sim \text{Beta}(a_1, b_1)$ and $\theta_2 \sim \text{Beta}(a_2, b_2)$, and arbitrarily selected the hyperparameters. One might ask why the subjective Bayesian approach did not involve a more rigorous elicitation of these parameters? The answer to this question lies in the criteria used for elicitation. For this case, four parameters need to be elicited, and there are numerous ways to do so, without a clear indication of which method will yield the most robust and reliable results. For example, one might set the mean and variance of the prior distributions, i.e., the considered beta distributions equal to the mean and variance of the geometric distribution. However, this approach can lead to infeasible solutions, such as negative values. Even if an alternative criterion is chosen and produces reasonable values, it still does not guarantee that the method will be applicable in all cases or lead to reliable results. Certainly, the most appealing approach would be to leverage prior information (past knowledge) about the parameters θ_1 and θ_2 . Unfortunately, such prior information is often unavailable in many situations. Furthermore, the authors did not train their methods on real data, which limits the generalizability and robustness of their results.

Alternatively, we adopt a Bayesian approach for estimating ρ using non-informative priors. Our motivation for using such priors, in addition to the importance of studying ρ under geometric distributions, lies in its wide range of real-world applications. The reasons

for choosing this approach are as follows: (i) this is the first study to consider objective Bayesian inference for ρ ; (ii) the priors considered in this study are informative yet contain minimal information; (iii) the priors eliminate the need for hyperparameter elicitation; (iv) the priors are obtained using formal rules; and (v) the priors possess the flavor of invariance, or certain modifications that offer improved performance, particularly in high-dimensional parameter settings, or alternative priors capable of producing credible intervals with coverage properties comparable to those of frequentist confidence intervals.

Before proceeding, we would like to point out that ρ is a function of the model parameters. In such cases, constructing non-informative priors requires particular care, except for Jeffreys' prior, which is invariant by design. Other types of priors, such as matching and reference priors, are generally not invariant under reparameterization. For instance, suppose matching priors are developed under the assumption that θ_1 is the parameter of interest and θ_2 is a nuisance parameter, or vice versa. If a one-to-one transformation is then applied to express the model in terms of ρ , the resulting prior is not guaranteed to be a matching prior for ρ . This is because the inverse transformation yields reference or matching priors that are valid for θ_1 and θ_2 , depending on which parameter was treated as influential, but not necessarily for a derived parameter like ρ . A similar issue arises with reference priors. While it might be argued, as shown by [14], that reference and matching priors are invariant, this invariance only holds when such priors are developed directly for ρ as a function of θ_1 and θ_2 , followed by an appropriate transformation. In that case, the resulting priors retain their desired properties with respect to ρ . A comparable situation is discussed in [15], who provided Bayesian inference for the entropy of the Lomax distribution, a quantity that also depends on the model parameters, and demonstrated that the reference and matching priors for this entropy cannot be derived from those of the model parameters. For further developments and illustrative examples, the reader is referred to [14], who specifically constructed matching priors for functions of model parameters.

The rest of the paper is organized as follows. Likelihood-based inference for ρ is presented in Section 2. In Section 3, we derive non-informative priors for ρ , including Jeffreys, reference, and probability-matching priors. Section 4 reports a simulation study conducted via MCMC to assess the performance of the Bayes estimators of ρ under various objective priors. In this section, we also provide the average estimates, standard errors, and coverage probabilities of the 95% credible intervals. The compatibility of the proposed priors is further examined in Section 5 using the posterior predictive distribution. In Section 6, the practical utility of the stress–strength reliability measure for the geometric distribution is demonstrated through two real data applications. Finally, concluding remarks are given in Section 7.

2. Likelihood-Based Inference for ρ

Let $\mathbf{X} = (X_1, \dots, X_n)^t$ and $\mathbf{Y} = (Y_1, \dots, Y_m)^t$ be independent random samples from $\text{GE}(\theta_1)$ and $\text{GE}(\theta_2)$, respectively, where $\text{GE}(\cdot)$ denotes the geometric distribution. Then, the likelihood function of (θ_1, θ_2) based on the realizations $\mathbf{x} = (x_1, \dots, x_n)^t$ and $\mathbf{y} = (y_1, \dots, y_m)^t$, is

$$\begin{aligned} L(\theta_1, \theta_2 | \mathbf{x}, \mathbf{y}) &= \prod_{i=1}^n \left\{ \theta_1 (1 - \theta_1)^{x_i - 1} \right\} \prod_{j=1}^m \left\{ \theta_2 (1 - \theta_2)^{y_j - 1} \right\}, \\ &= \theta_1^n \theta_2^m (1 - \theta_1)^{T_1 - n} (1 - \theta_2)^{T_2 - m}, \quad (n, m) \in \mathbb{N} \times \mathbb{N}, \end{aligned} \quad (2)$$

where $T_1 = \sum_{i=1}^n X_i$ and $T_2 = \sum_{j=1}^m Y_j$. The ML estimates of (θ_1, θ_2) can be obtained by maximizing the log-likelihood function with a mathematical form given by

$$\ell(\theta_1, \theta_2) = n \ln(\theta_1) + m \ln(\theta_2) + (T_1 - n) \ln(1 - \theta_1) + (T_2 - m) \ln(1 - \theta_1).$$

The partial derivatives of $\ell(\theta_1, \theta_2)$ with θ_1 and θ_2 are given below:

$$\frac{\partial \ell(\theta_1, \theta_2)}{\partial \theta_1} = \frac{n}{\theta_1} - \frac{T_1 - n}{1 - \theta_1} \quad \text{and} \quad \frac{\partial \ell(\theta_1, \theta_2)}{\partial \theta_2} = \frac{m}{\theta_2} - \frac{T_2 - m}{1 - \theta_2}.$$

On solving the above equations, the ML estimates of (θ_1, θ_2) are given by $\hat{\theta}_1 = \frac{n}{T_1}$ and $\hat{\theta}_2 = \frac{m}{T_2}$. Since $E(X) = \frac{1}{\theta_1}$ and $E(Y) = \frac{1}{\theta_2}$, the elements of the Fisher information matrix are computed via

$$-\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_1^2}\right) = \left(\frac{n}{\theta_1^2} + \frac{\mathbb{E}(T_1 - n)}{(1 - \theta_1)^2} = \frac{n}{\theta_1^2(1 - \theta_1)}\right),$$

$$\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_1 \partial \theta_2}\right) = \mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_2 \partial \theta_1}\right) = 0,$$

and

$$-\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_2^2}\right) = \frac{m}{\theta_2^2} + \frac{\mathbb{E}(T_2 - m)}{(1 - \theta_2)^2} = \frac{m}{\theta_2^2(1 - \theta_2)}.$$

Therefore, we have that

$$\mathbf{F}(\theta_1, \theta_2) = \begin{pmatrix} -\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_1^2}\right) & -\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_1 \partial \theta_2}\right) \\ -\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_2 \partial \theta_1}\right) & -\mathbb{E}\left(\frac{\partial^2 \ell(\theta_1, \theta_2)}{\partial \theta_2^2}\right) \end{pmatrix} = \begin{pmatrix} \frac{n}{\theta_1^2(1 - \theta_1)} & 0 \\ 0 & \frac{m}{\theta_2^2(1 - \theta_2)} \end{pmatrix}, \quad (3)$$

The following results are required to establish the asymptotic normality of the ML estimate of ρ .

Lemma 1. The ML estimates $\hat{\theta}_1 = \frac{n}{T_1}$ and $\hat{\theta}_2 = \frac{m}{T_2}$ are consistent estimators of θ_1 and θ_2 , respectively.

Proof. We provide only the proof of the consistency of $\hat{\theta}_1 = \frac{n}{T_1}$ for θ_1 , and the proof for the other estimator can be handled similarly. Since $W_n := \frac{T_1}{n} = \frac{\sum_{i=1}^n X_i}{n}$, it follows from the strong law of large numbers that $\frac{T_1}{n}$ converges in probability to $\mathbb{E}(X_1) = \frac{1}{\theta_1}$. Now, since $g(x) = \frac{1}{x}$ for $x > 0$ is a continuous function, it follows that $g(W_n) = \frac{n}{T_1}$ converges to $g(\mathbb{E}(X_1)) = \theta_1$. $\square$

The following theorem establishes that the ML estimates of θ_1 and θ_2 are asymptotically bivariate normal. The conditions required for this are as follows: The likelihood function must be identifiable and differentiable, and the true parameter value must lie in the interior of the parameter space. The partial derivatives of the log-likelihood function must be continuous and have a unique solution at the true parameter. Additionally, the observed Fisher information must converge to its expected value as the sample size increases, and the Fisher information matrix must be finite and positive definite. Finally, the log-likelihood function should satisfy smoothness conditions. For more information on these regularity conditions and the establishment of the asymptotic normality of the MLE, see [16] (1993, p. 132).

Theorem 1. We have that, as $n \rightarrow \infty$, $m \rightarrow \infty$, and $\frac{n}{m} \rightarrow p$, $\left(\sqrt{n}(\hat{\theta}_1 - \theta_1), \sqrt{m}(\hat{\theta}_2 - \theta_2)\right)^t \Rightarrow N_2(0, \mathbf{H}^{-1}(\theta_1, \theta_2))$, where $\Rightarrow$ stand for convergence in distribution, and $N_2(\cdot, \cdot)$ stand for bivariate normal distribution, and

$$\mathbf{H}(\theta_1, \theta_2) = \begin{pmatrix} \frac{1}{\theta_1^2(1-\theta_1)} & 0 \\ 0 & \frac{p}{\theta_2^2(1-\theta_2)} \end{pmatrix}, \quad (4)$$

Proof. Thanks to the exponential family of distributions, which satisfies all the mentioned regularity conditions, asymptotic normality follows. Since the geometric distribution belongs to the exponential family, the asymptotic normality of its MLEs is guaranteed. The above matrix is obtained by virtue of the regularity condition given in the preceding paragraph, which states that the average observed Fisher information must converge to its expected value as the sample size increases. To see this, the average observed Fisher information matrix, after simple rearrangement, is

$$\frac{1}{n}\mathbf{J}(\theta_1, \theta_2) = \begin{pmatrix} \frac{1}{\theta_1^2} + \frac{\bar{X}-1}{(1-\theta_1)^2} & 0 \\ 0 & \frac{m}{n} \left(\frac{1}{\theta_2^2} + \frac{\bar{Y}-1}{(1-\theta_2)^2} \right) \end{pmatrix},$$

Since $\bar{X}$ and $\bar{Y}$ converges in probability to $1/\theta_1$ and $1/\theta_2$ in probability as $n, m \rightarrow \infty$ respectively and by mean of the assumption that $m/n \rightarrow p \in (0, 1]$, it follows that $\frac{1}{n}\mathbf{J}(\theta_1, \theta_2)$ converges in probability to $\mathbf{H}(\theta_1, \theta_2)$. Noticed that when $p = 1$ implies $m = n$ and hence $\mathbf{H}(\theta_1, \theta_2)$ is the Fisher information matrix of independent geometric distributions per unit of observation, and this completes the proof. $\square$

Theorem 2 ([16] (1993, p. 132)). Suppose that $\sqrt{n}(W_n - \boldsymbol{\theta}) \Rightarrow N_k(\mathbf{0}, \boldsymbol{\Sigma})$, where W_n is k -dimensional random vectors, $\boldsymbol{\theta} = (\theta_1, \dots, \theta_k)^t$, and $\boldsymbol{\Sigma}$ is the variance-covariance matrix. Let $g: \mathbb{R}^k \rightarrow \mathbb{R}$ be a real-valued function with $\nabla := (\partial g(\boldsymbol{\theta})/\partial \theta_1, \dots, \partial g(\boldsymbol{\theta})/\partial \theta_k) \neq \mathbf{0}$ and continuous in a neighborhood of $\boldsymbol{\theta}$. Then, $\sqrt{n}(g(W_n) - g(\boldsymbol{\theta})) \Rightarrow N_k(\mathbf{0}, \nabla^t \boldsymbol{\Sigma} \nabla)$.

The above theorem is usually referred to as the Delta method. Here, we establish the first results relating to the estimate of ρ .

Theorem 3. As $n \rightarrow \infty$, $m \rightarrow \infty$, and $\frac{n}{m} \rightarrow p$, we have $\sqrt{n}(\hat{\rho} - \rho) \Rightarrow N(0, \sigma_\rho^2)$, where

$$\sigma_\rho^2 = \frac{\theta_1^2 \theta_2^2 (1 - \theta_1)}{\delta^2} \{1 + p(1 - \theta_1)(1 - \theta_2)\},$$

and $\delta = \{1 - (1 - \theta_1)(1 - \theta_2)\}^2$.

Proof. We have from, Theorem 1 that as $n \rightarrow \infty$, $m \rightarrow \infty$, $\left(\sqrt{n}(\hat{\theta}_1 - \theta_1), \sqrt{m}(\hat{\theta}_2 - \theta_2)\right)^t \Rightarrow N_2(0, \mathbf{H}^{-1}(\theta_1, \theta_2))$, where $\mathbf{H}(\theta_1, \theta_2)$ is given in (4). Clearly, $\rho: (0, 1) \times (0, 1) \rightarrow (0, 1)$, where

$$\rho =: \rho(\theta_1, \theta_2) = \frac{\theta_1}{1 - (1 - \theta_1)(1 - \theta_2)}.$$

Notice that the ρ plays the role as the $g(\cdot)$ function in Theorem 2, and hence we have that as $n \rightarrow \infty$, $m \rightarrow \infty$, and $\frac{n}{m} \rightarrow p$, $\sqrt{n}(\hat{\rho} - \rho) \Rightarrow N(0, \sigma_\rho^2)$, where

$$\sigma_\rho^2 = \nabla^t \mathbf{H}^{-1}(\theta_1, \theta_2) \nabla, \quad (5)$$

and

$$\nabla = \left(\frac{\partial \rho}{\partial \theta_1}, \frac{\partial \rho}{\partial \theta_2} \right)^t = \left(\frac{\theta_2}{\{1 - (1 - \theta_1)(1 - \theta_2)\}^2}, -\frac{\theta_1(1 - \theta_1)}{\{1 - (1 - \theta_1)(1 - \theta_2)\}^2} \right)^t.$$

In view of (5), we have that

$$\sigma_\rho^2 = \frac{\theta_1^2 \theta_2^2 (1 - \theta_1)}{\delta^2} \{1 + p(1 - \theta_1)(1 - \theta_2)\},$$

as desired. $\square$

Since $\hat{\theta}_1$ and $\hat{\theta}_2$ are consistence estimators by the mean of Lemma 1 an approximate $(1 - \alpha)100\%$ confidence interval of ρ is

$$\hat{\rho} \pm z_{\frac{\alpha}{2}} \hat{\sigma}_\rho, \text{ where } \hat{\sigma}_\rho^2 = \frac{\hat{\theta}_1^2 \hat{\theta}_2^2 (1 - \hat{\theta}_1)}{\hat{\delta}^2} \{1 + p(1 - \hat{\theta}_1)(1 - \hat{\theta}_2)\}.$$

3. Bayes Inferences: Objective Priors

In this section, we construct several important objective (non-informative) priors for the parameter ρ , including the Jeffreys prior, reference priors, and probability matching priors. These priors are developed to ensure minimal informativeness while retaining desirable inferential properties. Following the construction of each prior, we examine the propriety and finiteness of the resulting posterior distributions. We then proceed to derive the Bayes estimates of ρ under the squared error loss (SEL). To facilitate the construction of certain priors, particularly the reference and matching priors, we reparameterize the original model in terms of the parameter ρ . While the Jeffreys prior is invariant under one-to-one transformations and thus can be obtained directly without reparameterization, such transformations are crucial for deriving the reference and matching priors, as they depend explicitly on the parameter of interest. Therefore, we consider a one-to-one transformation from

$$\Omega_{(\theta_1, \theta_2)} = \{(\theta_1, \theta_2) : 0 < \theta_1 < 1, 0 < \theta_2 < 1\} \text{ onto } \Omega_{(\rho, \omega)} = \{(\rho, \omega) : 0 < \rho < 1, 0 < \omega < 1\},$$

where

$$\rho = \frac{\theta_1}{1 - (1 - \theta_1)(1 - \theta_2)} \text{ and } \omega = \theta_2, \quad (6)$$

along with the inverse of the transformation,

$$\theta_1 = \frac{\rho\omega}{1 - (1 - \omega)\rho} \text{ and } \theta_2 = \omega.$$

Therefore, the likelihood function of (ρ, ω) based on $(\mathbf{x}, \mathbf{y})^t$ is

$$\mathbb{L}(\rho, \omega | \mathbf{x}, \mathbf{y}) \propto \rho^n \omega^{n+m} (1 - \rho)^{T_1 - n} (1 - \omega)^{T_2 - m} (1 - \rho(1 - \omega))^{-T_1}. \quad (7)$$

We now aim to obtain the Fisher information matrix corresponding to the parameters ρ and ω . To this end, we first evaluate the Jacobian matrix of the transformation $\mathbf{M} = \frac{\partial(\theta_1, \theta_2)}{\partial(\rho, \omega)}$ which is given by

$$\mathbf{M} = \begin{pmatrix} \frac{\omega}{(1 - (1 - \omega)\rho)^2} & \frac{\rho(1 - \rho)}{(1 - (1 - \omega)\rho)^2} \\ 0 & 1 \end{pmatrix}. \quad (8)$$

Using Equations (3) and (8), the Fisher information matrix for the parameters ρ and ω , denoted by $\Sigma(\rho, \omega)$, is obtained through the relation $\Sigma(\rho, \omega) = \mathbf{M}^T \mathbf{F} \mathbf{M}$. This result is summarized in the following lemma.

Theorem 4. *The Fisher information matrix of (ρ, ω) has the following form*

$$\Sigma(\rho, \omega) = \begin{pmatrix} \sigma_{\rho\rho} & \sigma_{\rho\omega} \\ \sigma_{\omega\rho} & \sigma_{\omega\omega} \end{pmatrix} = \begin{pmatrix} \frac{n}{\rho^2(1-\rho)(1-(1-\omega)\rho)} & \frac{n}{\rho\omega(1-(1-\omega)\rho)} \\ \frac{n}{\rho\omega(1-(1-\omega)\rho)} & \frac{m\omega+(n+m)(1-\omega)(1-\rho)}{\omega^2(1-\omega)(1-(1-\omega)\rho)} \end{pmatrix}. \quad (9)$$

3.1. Jeffreys Prior

The first non-informative prior we consider is the Jeffreys prior, which is commonly used in Bayesian inference due to its simple derivation from the Fisher information matrix and its invariance property. In general, this type of prior is improper, as are many other non-informative priors. Additionally, it may be inappropriate when the dimensionality of the model increases. Therefore, from Equation (9), the Jeffreys prior of (ρ, ω) is

$$\pi_J(\rho, \omega) \propto \frac{1}{\rho\omega\{(1-\rho)(1-\omega)(1-(1-\omega)\rho)\}^{1/2}}, \quad 0 < \rho < 1, \quad 0 < \omega < 1. \quad (10)$$

The joint posterior of (ρ, ω) under $\pi_J(\rho, \omega)$ in (10) is

$$\pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) \propto \rho^{n-1} \omega^{n+m-1} (1-\rho)^{T_1-n-1/2} (1-\omega)^{T_2-m-1/2} (1-\rho(1-\omega))^{-T_1-1/2}. \quad (11)$$

Theorem 5. *The posterior $\pi_J(\rho, \omega | \mathbf{x}, \mathbf{y})$ is proper if and only if $m > 2$, $T_1 > n - 1/2$ and $T_2 > m - 1/2$.*

Proof. We have that for sufficiently small values $(\rho, \omega) \rightarrow (0, 0)$, we have that $\pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) \sim \rho^{n-1} \omega^{n+m-1}$, and hence $\int_0^1 \int_0^1 \pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) d\omega d\rho = 1/(n(n+m)) < +\infty$. Similarly, for the values $(\rho, \omega) \rightarrow (1, 1)$, we have that $\pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) \sim (1-\rho)^{T_1-n-1/2} (1-\omega)^{T_2-m-1/2}$. So, we have that

$$\begin{aligned} \int_0^1 \int_0^1 \pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) d\omega d\rho &= \int_0^1 \int_0^1 (1-\rho)^{T_1-n-1/2} (1-\omega)^{T_2-m-1/2} d\omega d\rho \\ &= \frac{1}{(T_1-n+1/2)(T_2-m+1/2)}. \end{aligned}$$

Thus, the expression is finite and positive if and only if $T_1 > n - 1/2$ and $T_2 > m - 1/2$. Next, when the values of (ρ, ω) are close to the singularity, i.e., $(\rho, \omega) \rightarrow (1, 0)$, we consider the first-order Taylor expansion of $g(\rho, \omega) := 1 - \rho(1 - \omega)$ around $(1, 0)$. We then obtain $1 - \rho(1 - \omega) \approx 1 - \rho + \omega$, and it follows that

$$\begin{aligned} \int_0^1 \int_0^1 \pi_J(\rho, \omega | \mathbf{x}, \mathbf{y}) d\omega d\rho &\propto \int_0^1 \int_0^1 \omega^{n+m-1} (1-\rho)^{T_1-n-\frac{1}{2}} (1-\rho+\omega)^{-T_1-\frac{1}{2}} d\omega d\rho \\ &= \int_0^1 \int_0^1 u^{n+m-1} v^{T_1-n-\frac{1}{2}} (v+u)^{-T_1-\frac{1}{2}} du dv, \end{aligned} \quad (12)$$

where in the last equation we used the transformation $u = \omega$ and $v = 1 - \rho$. Now, to check the finiteness of the above double integral, we consider the transformation $u = ts$ and $v = t(1-s)$ with Jacobian of transformation $|J| = t^{-1}$,

$$\int_0^1 \int_0^1 \frac{u^{n+m-1} v^{T_1-n-1/2}}{(v+u)^{T_1+1/2}} du dv = \int_0^1 \int_0^1 t^{m-3} s^{n+m-1} (1-s)^{T_1-n-1/2} ds dt,$$

$$= \frac{\Gamma(n+m)\Gamma(T_1-n+1/2)}{(m-2)\Gamma(T_1+m+1/2)}.$$

Hence, the right-hand side of the above equation is finite and positive, provided that $m > 2$. $\square$

3.2. Approximate Reference Priors

The reference prior, originally introduced by [17] and further refined in the same work, serves as an enhancement of Jeffreys' prior, particularly in multi-parameter settings where Jeffreys' prior may not perform well; see [18]. A key feature of this approach is the partitioning of model parameters into groups based on their inferential priority. The derivation of the reference prior follows a general algorithm outlined in [17]. Notably, one significant result arising from this algorithm, which proves useful in various applications, is that it avoids the need for nested compact subsets in the nuisance parameter space. This is formally presented in the following proposition.

Proposition 1 ([19] (2000, p. 328)). *Let ψ be the parameter of interest and suppose that λ is the nuisance parameter. Let*

$$D_\psi(\psi, \lambda) := \frac{\det(\mathbf{I}(\psi, \lambda))}{I_{\lambda\lambda}(\psi, \lambda)},$$

where $\mathbf{I}(\psi, \lambda)$ is the Fisher information matrix and $I_{\lambda\lambda}(\psi, \lambda)$ is the bottom-right element of $\mathbf{I}(\psi, \lambda)$. If the parameter of interest ψ is independent of the nuisance parameter space $\Lambda(\psi)$, and if

$$D_\psi^{\frac{1}{2}}(\psi, \lambda) = h_1(\psi)g_1(\lambda) \text{ and } I_{\lambda\lambda}^{\frac{1}{2}}(\psi, \lambda) = h_2(\psi)g_2(\lambda),$$

then the reference prior with respect to the ordered group (ψ, λ) is

$$\pi(\psi, \lambda) = h_1(\psi)g_2(\lambda).$$

It is worth emphasizing that neither the procedure outlined in [17] nor the result in Proposition 1 yields an explicit form of the reference prior for ρ . As an alternative, we construct an approximate reference prior by leveraging the result in Proposition 1.

Proposition 2. *The approximate reference prior for the ordered group parameters (ρ, ω) is*
(i) *for sufficiently small ρ ,*

$$\pi_{R_1}(\rho, \omega) \propto \frac{\sqrt{m\omega + (n+m)(1-\omega)}}{\rho \omega (1-\omega)^{1/2}}, \quad 0 < \rho < 1, 0 < \omega < 1. \quad (13)$$

(ii) *for $\rho \rightarrow 1$, we have that*

$$\pi_{R_2}(\rho, \omega) \propto \frac{1}{\sqrt{\omega(1-\rho)(1-\omega)}}, \quad 0 < \rho < 1, 0 < \omega < 1. \quad (14)$$

Proof. We provide the proof only for part (i). We have that

$$\begin{aligned} D_{\rho}^{\frac{1}{2}}(\rho, \omega) &= \frac{1}{\rho(1-\rho)^{1/2}\{m\omega + (n+m)(1-\omega)(1-\rho)\}^{1/2}}, \\ &\sim \frac{1}{\rho\{m\omega + (n+m)(1-\omega)\}^{1/2}}, \text{ as } \rho \rightarrow 0 \\ &:= h_1(\rho)g_1(\omega) \end{aligned}$$

where $h_1(\rho) = \rho^{-1}$ and $g_1(\omega) = [m\omega + (n+m)(1-\omega)]^{-1/2}$. Now

$$\begin{aligned} \sigma_{\omega\omega}^{\frac{1}{2}}(\rho, \omega) &\sim \frac{m\omega + (n+m)(1-\omega)}{\omega^2(1-\omega)}, \text{ as } \rho \rightarrow 0 \\ &:= h_2(\rho)g_2(\omega), \end{aligned}$$

where $h_2(\rho) = 1$ and $g_2(\omega) = \sigma_{\omega\omega}^{\frac{1}{2}}(\rho, \omega)$. In view of Proposition 1, it follows that the approximate reference prior for the ordered group (ρ, ω) is the one given in (13) $\square$

Theorem 6. Under the assumption of Theorem 5, the posterior distributions under $\pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y})$ and $\pi_{R_2}(\rho, \omega|\mathbf{x}, \mathbf{y})$ are proper distributions.

Proof. We provide the proof for $\pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y})$. The proof of $\pi_{R_2}(\rho, \omega|\mathbf{x}, \mathbf{y})$ is quite similar to the proof in Theorem. Since $(m\omega + (n+m)(1-\omega))^{1/2}$ is a decreasing function in ω , we have

$$\pi_{R_1}(\rho, \omega) \leq \frac{\sqrt{n+m}}{\rho\omega(1-\omega)^{1/2}}.$$

Therefore, we have that

$$\pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y}) \leq K\rho^{n-1}\omega^{n+m-1}(1-\rho)^{T_1-n}(1-\omega)^{T_2-m-1/2}(1-\rho(1-\omega))^{-T_1},$$

where $K = \sqrt{n+m}$. Similar to the proof in Theorem 5, we have for sufficiently small values of ρ and ω that $\pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y}) \leq K\rho^{n-1}\omega^{n+m-1}$ it then follows that the posterior distribution is finite, i.e., $\int_0^1 \int_0^1 \pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y})d\omega d\rho \leq K \int_0^1 \int_0^1 \pi_{R_1}\rho^{n-1}\omega^{n+m-1} \leq d\omega d\rho \leq K/(n(n+m)) < +\infty$. When ρ and ω are close to 1, it then follows for $T_1 > n-1$ and $T_2 > m-1/2$ that $\int_0^1 \int_0^1 \pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y})d\omega d\rho \leq K \int_0^1 \int_0^1 (1-\rho)^{T_1-n}(1-\omega)^{T_2-m-1/2}d\omega d\rho < +\infty$. Finally, when the values of (ρ, ω) close to $(\rho, \omega) \rightarrow (1, 0)$, the first-order Taylor expansion for $g(\rho, \omega) := (1-\rho(1-\omega))$ around $(1, 0)$, implies that $1-\rho(1-\omega) \approx 1-\rho+\omega$, it then follows that

$$\pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y}) \leq K \omega^{n+m-1}(1-\rho)^{T_1-n-1/2}(1-\rho+\omega)^{-T_1-1/2}.$$

Therefore, we have that

$$\int_0^1 \int_0^1 \pi_{R_1}(\rho, \omega|\mathbf{x}, \mathbf{y})d\omega d\rho \leq K \int_0^1 \int_0^1 \omega^{n+m-1}(1-\rho)^{T_1-n-1/2}(1-\rho+\omega)^{-T_1}d\omega d\rho.$$

Similarly to the analysis of (12), the above integral is finite. $\square$

3.3. Probability Matching Prior

Probability matching priors for the parameter of interest were introduced by [20] to ensure that the posterior probabilities of specified intervals exactly or approximately match their corresponding frequentist coverage probabilities. Mathematically speaking, let be $\pi(\rho, \omega)$ be a given prior of (ρ, ω) with parameter of interest ρ and ω is the nuisance param-

eter and $\rho^{(1-\alpha)}(\pi, \mathbf{X}, \mathbf{Y})$ is the $(1 - \alpha)$ th percentile of the marginal posterior distribution of ρ , where $\mathbf{X} = (X_1, X_2, \dots, X_n)$ and $\mathbf{Y} = (Y_1, Y_2, \dots, Y_m)$ are the sample data. Then, $\pi(\cdot)$ is called a second-order probability matching prior if

$$\Pr(\rho \leq \rho^{(1-\alpha)}(\pi, \mathbf{X}, \mathbf{Y})) = 1 - \alpha + o(n^{-1}),$$

holds for all $\alpha \in (0, 1)$. It is worth emphasizing that Jeffreys prior achieves probability matching with an asymptotic error of order n^{-1} when there is only a single parameter, as shown in [20]. However, this desirable property does not necessarily extend to models involving multiple parameters, which are common in many statistical applications. According to the results in [20], a prior π is a matching prior if and only if it satisfies the following partial differential equation

$$\frac{\partial}{\partial \rho} \left\{ \frac{\sigma^{\rho\rho} \pi(\rho, \omega)}{\sqrt{\sigma^{\rho\rho}}} \right\} + \frac{\partial}{\partial \omega} \left\{ \frac{\gamma^{\rho\omega} \pi(\rho, \omega)}{\sqrt{\sigma^{\rho\rho}}} \right\} = 0, \quad (15)$$

where σ^{ij} is the (i, j) th-element of the inverse of Fisher information matrix in (9) and $i, j \in \{\rho, \omega\}$. Clearly, this approach can be used to derive a matching prior for ρ using the Fisher information matrix $\Sigma(\rho, \omega)$ given in (9). However, when the parameter of interest is a function of the model parameters (as in our case, where ρ is such a function), it is first necessary to obtain the transformed Fisher information in terms of ρ . This is then followed by solving a partial differential equation, as shown in (15). A natural question arises: can a similar representation to (15) be developed using the Fisher information matrix of the original model parameters to obtain a matching prior? Interestingly, [21] successfully demonstrates that matching priors can be constructed not only for individual parameters but also for functions of these parameters. Specifically, let $\psi = \psi(\theta_1, \theta_2)$ be a real-valued function of (θ_1, θ_2) . Then, the gradient of ψ is given as $\nabla_\psi(\theta_1, \theta_2) := \{(\partial/\partial\theta_1)\psi(\theta_1, \theta_2), (\partial/\partial\theta_2)\psi(\theta_1, \theta_2)\}^T$ and defined as

$$\zeta_\psi(\theta_1, \theta_2) = \frac{\mathbf{I}^{-1}(\theta_1, \theta_2) \nabla_\psi(\theta_1, \theta_2)}{\sqrt{\{\nabla_\psi^T(\theta_1, \theta_2) \mathbf{I}^{-1}(\theta_1, \theta_2) \nabla_\psi(\theta_1, \theta_2)\}}} = (\zeta_1(\theta_1, \theta_2), \zeta_2(\theta_1, \theta_2))^T.$$

Then, $\pi_\psi(\theta_1, \theta_2)$ is a matching prior for $\psi(\theta_1, \theta_2)$ if and only if it satisfies the partial differential equation

$$\frac{\partial}{\partial \theta_1} \{\zeta_1(\theta_1, \theta_2) \pi(\theta_1, \theta_2)\} + \frac{\partial}{\partial \theta_2} \{\zeta_2(\theta_1, \theta_2) \pi(\theta_1, \theta_2)\} = 0. \quad (16)$$

Since ρ is a function of θ_1 and θ_2 the matching prior for ρ can thus be derived directly without requiring an explicit one-to-one transformation into the (ρ, ω) space, using the Fisher information matrix in terms of the original parameters as given in (3).

Theorem 7. The matching prior for ρ given in (1) is

$$\pi_\rho(\theta_1, \theta_2) = \frac{\sqrt{m + n(1 - \theta_1)(1 - \theta_2)}}{\theta_1 \theta_2 (1 - \theta_1)^{1/2} (1 - \theta_2)}. \quad (17)$$

Proof. Put $\psi = \rho$, it then follows that

$$\nabla_\rho(\theta_1, \theta_2) = \left(\frac{\theta_2}{\delta}, \frac{-\theta_1(1 - \theta_1)}{\delta} \right),$$

where δ is defined in Theorem 3. Next, let $c = m/n$

$$\zeta(\theta_1, \theta_2) = \left(\frac{c^{1/2} \theta_1 (1 - \theta_1)^{1/2}}{\sqrt{m + n(1 - \theta_1)(1 - \theta_2)}}, \frac{c^{-1/2} \theta_2 (1 - \theta_1)^{1/2} (1 - \theta_2)}{\sqrt{m + n(1 - \theta_1)(1 - \theta_2)}} \right).$$

Therefore, $\pi_\rho(\theta_1, \theta_2)$ is a matching prior for ρ if and only if

$$\frac{\partial}{\partial \theta_1} \left\{ \frac{c^{1/2} \theta_1 (1 - \theta_1)^{1/2} \pi_\rho(\theta_1, \theta_2)}{\sqrt{m + n(1 - \theta_1)(1 - \theta_2)}} \right\} + \frac{\partial}{\partial \theta_2} \left\{ \frac{c^{-1/2} \theta_2 (1 - \theta_1)^{1/2} (1 - \theta_2) \pi_\rho(\theta_1, \theta_2)}{\sqrt{m + n(1 - \theta_1)(1 - \theta_2)}} \right\} = 0. \quad (18)$$

A solution of the above partial differential equation is given in (17). $\square$

Remark 1. The reference prior for irrespective of which the parameter of interest θ_1 or θ_2 is

$$\pi_{RM}(\theta_1, \theta_2) = \frac{1}{\theta_1 \theta_2 \sqrt{(1 - \theta_1)(1 - \theta_2)}}. \quad (19)$$

The reference prior in (19) is straightforward to derive since the Fisher information matrix is diagonal. Interestingly, the corresponding prior is also a matching prior. However, the prior in (19) is not equal to the prior in (17), which implies that the prior in (19) cannot serve as a matching prior for ρ . Specifically, if a one-to-one transformation is performed to express the parameters in terms of ρ and $\omega = \theta_2$ using the prior in (19), the resulting prior, though now a function of ρ and $\omega = \theta_2$, cannot be regarded as a matching prior for ρ . In contrast, applying the same transformation to the prior in (17) does yield a valid matching prior for ρ , as also confirmed by the invariance results of [14].

Now the default matching prior for ρ is given in (17). Since any 1-1 transformation preserves the matching prior propriety, we use the transformation (6) and consequently the matching prior is then given by

$$\pi_M(\rho, \omega) \propto \frac{\sqrt{m\omega + (n + m)(1 - \rho)(1 - \omega)}}{\rho \omega (1 - \rho)^{1/2} (1 - \omega) (1 - \rho(1 - \omega))}. \quad (20)$$

Theorem 8. Under the assumption of Theorem 5, the posterior $\pi_M(\rho, \omega | \mathbf{x}, \mathbf{y})$ is proper.

Proof. We have that

$$\begin{aligned} \int_0^1 \int_0^1 \pi_M(\rho, \omega | \mathbf{x}, \mathbf{y}) d\rho d\omega &= \int_0^1 \int_0^1 \rho^n \omega^{n+m} (1 - \rho)^{T_1 - n} (1 - \omega)^{T_2 - m} (1 - \rho(1 - \omega))^{-T_1} \\ &\quad \times \frac{\sqrt{m\omega + (n + m)(1 - \rho)(1 - \omega)}}{\rho \omega (1 - \rho)^{1/2} (1 - \omega) (1 - \rho(1 - \omega))} d\rho d\omega \\ &\leq C \int_0^1 \int_0^1 \rho^{n-1} \omega^{n+m-1} (1 - \rho)^{T_1 - n - \frac{1}{2}} (1 - \omega)^{T_2 - m - 1} \\ &\quad \times (1 - \rho(1 - \omega))^{-T_1 - 1} d\rho d\omega, \end{aligned}$$

where $C = \sqrt{2m + n}$. It can be verified, in conjunction with the analysis of (12), that the above integral is finite. $\square$

4. Numerical Computations

In this section, we investigate the performance of Bayes estimators for ρ under the squared error loss function (SELF) using Jeffreys, reference, and matching priors, and also consider the maximum likelihood estimate. The comparison is based on bias and root mean squared error, along with the frequentist coverage probabilities of the highest posterior density (HPD) credible intervals derived from these priors.

4.1. Random Walk Metropolis–Hastings Algorithm

The Bayes estimates of ρ under the various non-informative priors examined in this study do not admit closed-form solutions. Likewise, the marginal posterior distribution of ρ cannot be expressed analytically. To handle this, we apply the random walk Metropolis–Hastings (MH) algorithm to generate Markov Chain Monte Carlo (MCMC) samples, denoted by ρ_i . This approach is selected for its versatility in producing samples from a wide variety of proposal distributions, particularly when the conditional posterior distributions of the parameters are not available in explicit form. In our implementation, we adopt a symmetric random-walk MH algorithm with a bivariate normal proposal to draw samples from the joint posterior distribution of (ρ, ω) . These samples of ρ are subsequently used to construct HPD credible intervals for the reliability parameter ρ . The procedure for generating samples from the joint posterior distribution via the random walk MH algorithm is outlined in Algorithm 1.

Algorithm 1 Metropolis–Hastings Algorithm for Estimating ρ and ω

- 1: **Initialize:** Choose initial values ρ_0 and ω_0 .
- 2: **for** $\ell = 1$ to N **do**
- 3: Given the current state $(v^{(\ell)}, w^{(\ell)})$, where $v^{(\ell)} = \log(\rho^{(\ell)} / 1 - \rho^{(\ell)})$ and $w^{(\ell)} = \log(\omega^{(\ell)} / 1 - \omega^{(\ell)})$, propose a new state

$$(v^{(\ell+1)}, w^{(\ell+1)}) \sim N_2\left((v^{(\ell)}, w^{(\ell)}), c\Sigma\right),$$

Here, Σ denotes the variance-covariance matrix, and c is a scaling factor used to adjust the acceptance rate, which is typically maintained between 20% and 40%, as suggested by Neal and Roberts [22].

- 4: Compute the acceptance probability

$$\gamma = \min\left\{1, \frac{\pi_{v,w}(v^{(\ell+1)}, w^{(\ell+1)} | \mathbf{x}, \mathbf{y})}{\pi_{v,w}(v^{(\ell)}, w^{(\ell)} | \mathbf{x}, \mathbf{y})}\right\},$$

where $\pi_{v,w}(v, w | \mathbf{x}, \mathbf{y}) = \pi_{\rho,\omega}(h_1(v), g_1(w) | \mathbf{x}, \mathbf{y}) h_1(v) g_1(w) (1 + \exp(v))^{-1} (1 + \exp(w))^{-1}$, $h_1(v) := \exp(v) / (1 + \exp(v))$, and $g_1(w) := \exp(w) / (1 + \exp(w))$

- 5: Draw $U \sim \text{Uniform}[0, 1]$.
- 6: **if** $U \leq \gamma$ **then**
- 7: Accept the proposed values: $(v^{(\ell+1)}, w^{(\ell+1)})$.
- 8: **else**
- 9: Reject and set $(v^{(\ell+1)}, w^{(\ell+1)}) = (v^{(\ell)}, w^{(\ell)})$.
- 10: **end if**
- 11: **end for**
- 12: Discard the first B samples as burn-in, and retain $(v^{(B+1)}, w^{(B+1)}), \dots, (v^{(N)}, w^{(N)})$
- 13: Compute the Bayes estimate of ρ under SELF using the relation $\rho = \exp(v) / (1 + \exp(v))$

$$\hat{\rho} = \frac{1}{N-B} \sum_{\ell=B+1}^N \rho_{\ell}.$$

- 14: Compute the $(1 - \alpha)\%$ HPD credible interval for ρ :

- Sort $\rho_1, \dots, \rho_M$ as $\rho_{(1)} < \rho_{(2)} < \dots < \rho_{(M)}$.
- For $\ell = B, B+1, \dots, (N-B) - \lfloor (N-B)(1-\alpha) \rfloor$, define

$$I_{\ell} = \left[\rho_{(\ell)}, \rho_{(\ell + \lfloor (N-B)(1-\alpha) \rfloor)} \right].$$

- Select the interval I_{ℓ} with the smallest width as the HPD credible interval.
-

4.2. Simulation Experiment

We perform a simulation Study to evaluate the performance of Bayes estimators of ρ under the non-informative priors considered in this paper. The investigation emphasizes small sample sizes and different parameter configurations. Specifically, we examine sample size combinations $(n, m) = (5, 5), \dots, (20, 20)$ together with parameter pairs (θ_1, θ_2) provided in Table 1. For each scenario, independent samples $\mathbf{X}$ and $\mathbf{Y}$ are generated from geometric distributions with parameters θ_1 and θ_2 , respectively. Bayes estimates are computed under the squared error loss function (SELF). For a given sample, the algorithm described in Algorithm 1 is applied to produce 5500 MCMC draws, discarding the first 500 as burn-in. The remaining samples are then used to obtain Bayes estimates and construct credible intervals. The estimators are assessed over 1000 replications of size (n, m) in terms of average estimate (ES), standard errors (SD), frequentist coverage probability (CP) of the 95% HPD credible intervals. A detailed summary of the results is presented in Table 1. From these tables, several finding can be concluded:

- While the sample sizes considered in this article are relatively small, it is observed that the performance of all estimators of ρ improves as the sample size increases, with the average estimates approaching the true values and the standard deviations decreasing. The reduction in standard deviations demonstrates the consistency of these estimators.
- It is interesting to note that the Bayes estimator $\hat{\rho}$ based on π_M outperforms the other estimators, as its 95% HPD credible intervals achieve frequentist coverage probabilities that remain close to 0.95 across all considered scenarios.
- For $0.1493 \leq \rho \leq 0.1818$, the Bayes estimator $\hat{\rho}$ based on π_M and the ML estimator outperform the other estimators by exhibiting smaller standard deviations across all considered scenarios. Moreover, both estimators provide good frequentist coverage probabilities that approach 0.95, with the Bayes estimator based on π_M showing better convergence than the ML estimator.
- For $0.6667 \leq \rho \leq 0.9091$, the Bayes estimators using π_{R_1} , π_{R_2} , and π_J perform much better than the other two estimators in terms of exhibiting smaller standard deviations. However, the Bayes estimator under π_M is superior in terms of frequentist coverage probabilities, which remain close to 0.95, followed closely by the ML estimator.
- For $\rho = 0.989$, the Bayes estimator based on π_M performs significantly better than the others in terms of exhibiting smaller standard deviations, followed closely by the Bayes estimator under π_J . Moreover, the Bayes estimator based on π_M outperforms the one using π_J in terms of achieving frequentist coverage probabilities that remain close to 0.95.

Table 1. Simulation results showing estimates (Est), standard errors (SE), and coverage probabilities (CP) for different methods under various parameter settings.

Scenario 1: $\theta_1 = 0.05, \theta_2 = 0.3, \rho = 0.1493$															
(n,m)	ML			π_J			π_{R_1}			π_{R_2}			π_M		
	ES	SD	PC	ES	SD	PC	ES	SD	PC	ES	SD	PC	ES	SD	PC
(5,5)	0.1851	0.0867	0.919	0.1881	0.0870	0.920	0.1970	0.0849	0.923	0.2245	0.0916	0.864	0.1738	0.0862	0.945
(10,5)	0.1733	0.0707	0.914	0.1754	0.0702	0.924	0.1861	0.0682	0.930	0.2016	0.0744	0.880	0.1580	0.0602	0.942
(5,10)	0.1761	0.0804	0.912	0.1779	0.0811	0.911	0.1812	0.0800	0.915	0.2057	0.0854	0.868	0.1732	0.0803	0.940
(10,10)	0.1667	0.0594	0.917	0.1683	0.0595	0.915	0.1732	0.0590	0.926	0.1862	0.0619	0.893	0.1607	0.0588	0.947
(15,10)	0.1639	0.0532	0.907	0.1652	0.0530	0.903	0.1706	0.0526	0.916	0.1799	0.0551	0.878	0.1569	0.0517	0.925
(10,15)	0.1652	0.0574	0.912	0.1661	0.0574	0.917	0.1688	0.0566	0.914	0.1813	0.0594	0.883	0.1617	0.0573	0.948
(15,15)	0.1600	0.0466	0.920	0.1609	0.0463	0.921	0.1642	0.0463	0.922	0.1730	0.0483	0.897	0.1559	0.0459	0.946
(20,15)	0.1578	0.0426	0.904	0.1587	0.0427	0.913	0.1621	0.0423	0.912	0.1687	0.0440	0.890	0.1531	0.0418	0.945
(15,20)	0.1601	0.0441	0.928	0.1607	0.0441	0.921	0.1630	0.0439	0.924	0.1713	0.0457	0.902	0.1573	0.0438	0.953
(20,20)	0.1574	0.0411	0.907	0.1582	0.0412	0.917	0.1606	0.0410	0.912	0.1671	0.0425	0.892	0.1543	0.0408	0.942
Scenario 2: $\theta_1 = 1, \theta_2 = 0.5, \rho = 0.1818$															
(5,5)	0.2099	0.0922	0.9180	0.2173	0.0936	0.9310	0.2374	0.0920	0.9230	0.2513	0.0959	0.8850	0.2033	0.0920	0.9180
(10,5)	0.2039	0.0756	0.8950	0.2090	0.0758	0.9080	0.2322	0.0746	0.9000	0.2334	0.0784	0.8780	0.1919	0.0732	0.9250
(5,10)	0.2112	0.0923	0.9190	0.2147	0.0925	0.9210	0.2235	0.0927	0.9240	0.2445	0.0959	0.8770	0.2114	0.0924	0.9350
(10,10)	0.1949	0.0637	0.9190	0.1982	0.0641	0.9250	0.2090	0.0645	0.9280	0.2161	0.0666	0.8950	0.1910	0.0636	0.9390
(15,10)	0.1952	0.0548	0.9230	0.1980	0.0552	0.9240	0.2091	0.0549	0.9200	0.2121	0.0568	0.8940	0.1896	0.0540	0.9080
(10,15)	0.1965	0.0623	0.9120	0.1989	0.0629	0.9120	0.2056	0.0632	0.9190	0.2150	0.0649	0.8900	0.1949	0.0633	0.9390
(15,15)	0.1960	0.0556	0.9230	0.1987	0.0554	0.9240	0.2103	0.0561	0.9180	0.2129	0.0576	0.9010	0.1905	0.0547	0.9150
(20,15)	0.1913	0.0474	0.9090	0.1932	0.0477	0.9080	0.2009	0.0477	0.9050	0.2031	0.0489	0.8970	0.1878	0.0471	0.9140
(15,20)	0.1922	0.0484	0.9140	0.1939	0.0487	0.9160	0.1993	0.0487	0.9180	0.2052	0.0499	0.8950	0.1906	0.0484	0.9560
(20,20)	0.1914	0.0420	0.9290	0.1928	0.0421	0.9300	0.1982	0.0422	0.9290	0.2020	0.0430	0.9100	0.1891	0.0419	0.9520
Scenario 3: $\theta_1 = 0.5, \theta_2 = 0.5, \rho = 0.6667$															
(5,5)	0.6488	0.1403	0.915	0.6614	0.1380	0.9120	0.6773	0.1276	0.9290	0.6820	0.1272	0.9110	0.6794	0.1523	0.946
(10,5)	0.6498	0.1096	0.9250	0.6622	0.1071	0.9410	0.6847	0.0967	0.9390	0.6722	0.1005	0.9390	0.6613	0.1150	0.9300
(5,10)	0.6537	0.1353	0.928	0.6604	0.1334	0.9230	0.6654	0.1303	0.9330	0.6823	0.1255	0.9200	0.6873	0.1457	0.997
(10,10)	0.6612	0.1027	0.9330	0.6670	0.1027	0.9240	0.6769	0.0983	0.9290	0.6782	0.0976	0.9270	0.6765	0.1077	0.9580
(15,10)	0.6594	0.0891	0.9280	0.6649	0.0888	0.9280	0.6770	0.0840	0.9300	0.6721	0.0852	0.9250	0.6682	0.0921	0.9450
(10,15)	0.6595	0.1005	0.9250	0.6641	0.0999	0.9260	0.6687	0.0973	0.9310	0.6748	0.0956	0.9230	0.6763	0.1044	0.9720
(15,15)	0.6637	0.0884	0.9200	0.6677	0.0879	0.9150	0.6741	0.0851	0.9130	0.6747	0.0847	0.9100	0.6738	0.0908	0.9570
(20,15)	0.6641	0.0750	0.9370	0.6680	0.0750	0.9190	0.6762	0.0724	0.9220	0.6740	0.0725	0.9200	0.6714	0.0770	0.9530
(15,20)	0.6631	0.0841	0.9310	0.6665	0.0838	0.9250	0.6708	0.0825	0.9260	0.6737	0.0816	0.9190	0.6742	0.0866	0.9700
(20,20)	0.6620	0.0772	0.9150	0.6654	0.0768	0.9210	0.6706	0.0750	0.9190	0.6712	0.0747	0.9180	0.6700	0.0786	0.9540

Table 1. Cont.

Scenario 4: $\theta_1 = 0.6, \theta_2 = 0.2, \rho = 0.8824$															
(5,5)	0.8736	0.0912	0.925	0.8491	0.0902	0.913	0.8434	0.0873	0.920	0.8487	0.0841	0.921	0.8668	0.0920	0.946
(10,5)	0.8554	0.0761	0.915	0.8598	0.0743	0.903	0.8588	0.0707	0.911	0.8576	0.0716	0.906	0.8697	0.0771	0.955
(5,10)	0.8844	0.0796	0.936	0.8565	0.0813	0.911	0.8493	0.0824	0.915	0.8573	0.0767	0.911	0.8776	0.0821	0.953
(10,10)	0.8786	0.0620	0.940	0.8659	0.0620	0.924	0.8627	0.0623	0.921	0.8643	0.0609	0.920	0.8830	0.0635	0.948
(15,10)	0.8654	0.0534	0.935	0.8679	0.0528	0.929	0.8661	0.0521	0.936	0.8659	0.0518	0.936	0.8756	0.0558	0.949
(10,15)	0.8646	0.0585	0.949	0.8661	0.0592	0.913	0.8631	0.0592	0.916	0.8653	0.0576	0.920	0.8800	0.0591	0.954
(15,15)	0.8696	0.0493	0.950	0.8714	0.0491	0.923	0.8690	0.0488	0.924	0.8701	0.0482	0.925	0.8799	0.0486	0.948
(20,15)	0.8703	0.0450	0.938	0.8722	0.0446	0.930	0.8704	0.0445	0.932	0.8704	0.0441	0.931	0.8780	0.0455	0.950
(15,20)	0.8706	0.0482	0.943	0.8720	0.0481	0.915	0.8696	0.0484	0.928	0.8711	0.0473	0.920	0.8810	0.0485	0.955
(20,20)	0.8735	0.0432	0.945	0.8747	0.0430	0.920	0.8730	0.0432	0.916	0.8734	0.0429	0.920	0.8841	0.0438	0.951
Scenario 5: $\theta_1 = \theta_2 = 0.9, \rho = 0.9091$															
(5,5)	0.8798	0.0780	0.943	0.8864	0.0739	0.967	0.8953	0.0670	0.981	0.8952	0.0668	0.973	0.9383	0.0829	0.953
(10,5)	0.9076	0.0617	0.955	0.9142	0.0568	0.974	0.9233	0.0500	0.983	0.9164	0.0543	0.985	0.9381	0.0621	0.951
(5,10)	0.8890	0.0722	0.943	0.8923	0.0705	0.979	0.8958	0.0672	0.982	0.9012	0.0643	0.983	0.9274	0.0771	0.955
(10,10)	0.9087	0.0563	0.956	0.9116	0.0548	0.981	0.9161	0.0518	0.983	0.9151	0.0524	0.984	0.9183	0.0573	0.953
(15,10)	0.9157	0.0520	0.949	0.9185	0.0506	0.934	0.9230	0.0477	0.930	0.9200	0.0494	0.920	0.9054	0.0528	0.950
(10,15)	0.9097	0.0568	0.954	0.9118	0.0556	0.980	0.9140	0.0540	0.983	0.9150	0.0532	0.985	0.8988	0.0580	0.965
(15,15)	0.9169	0.0502	0.950	0.9187	0.0493	0.942	0.9216	0.0473	0.927	0.9210	0.0476	0.932	0.9164	0.0507	0.949
(20,15)	0.9234	0.0406	0.951	0.9254	0.0397	0.898	0.9281	0.0382	0.893	0.9265	0.0389	0.898	0.9182	0.0408	0.954
(15,20)	0.9198	0.0470	0.952	0.9213	0.0462	0.944	0.9234	0.0448	0.944	0.9233	0.0449	0.944	0.9395	0.0474	0.951
(20,20)	0.9253	0.0408	0.948	0.9267	0.0398	0.894	0.9286	0.0390	0.886	0.9279	0.0391	0.891	0.9100	0.0407	0.950
Scenario 6: $\theta_1 = 0.9, \theta_2 = 0.1, \rho = 0.989$															
(5,5)	0.9710	0.0249	0.970	0.9727	0.0237	0.976	0.9694	0.0249	0.970	0.9702	0.0242	0.968	0.9879	0.0197	0.955
(10,5)	0.9782	0.0184	0.962	0.9795	0.0178	0.969	0.9780	0.0182	0.958	0.9777	0.0186	0.963	0.9865	0.0157	0.952
(5,10)	0.9742	0.0202	0.953	0.9750	0.0197	0.971	0.9719	0.0217	0.966	0.9728	0.0208	0.967	0.9893	0.0166	0.954
(10,10)	0.9806	0.0147	0.964	0.9812	0.0143	0.969	0.9799	0.0152	0.962	0.9800	0.0151	0.969	0.9878	0.0130	0.954
(15,10)	0.9830	0.0123	0.944	0.9836	0.0121	0.920	0.9828	0.0124	0.925	0.9826	0.0125	0.929	0.9878	0.0111	0.949
(10,15)	0.9811	0.0138	0.952	0.9815	0.0135	0.968	0.9804	0.0143	0.956	0.9806	0.0140	0.957	0.9882	0.0122	0.951
(15,15)	0.9840	0.0105	0.956	0.9844	0.0103	0.953	0.9836	0.0106	0.956	0.9837	0.0106	0.944	0.9886	0.0095	0.951
(20,15)	0.9854	0.0090	0.953	0.9858	0.0088	0.899	0.9852	0.0091	0.910	0.9851	0.0091	0.906	0.9888	0.0084	0.953
(15,20)	0.9840	0.0106	0.946	0.9843	0.0104	0.940	0.9837	0.0108	0.948	0.9837	0.0108	0.942	0.9885	0.0098	0.951
(20,20)	0.9853	0.0092	0.945	0.9856	0.0091	0.894	0.9850	0.0094	0.907	0.9850	0.0094	0.900	0.9887	0.0086	0.950

5. Posterior Predictive Assessment of the Model

Once posterior propriety has been established under the proposed priors, the subsequent task is to assess how well the model represents the observed data. A widely used strategy in Bayesian analysis is to perform posterior predictive checks, which consist of drawing predictive samples from the model by conditioning on the posterior distribution of the parameters and then comparing these replications with the actual data through graphical diagnostics. To complement the visual assessment, we also consider the posterior predictive p -value, originally suggested by [23] and later formalized by [24], as a quantitative measure of fit. This method requires the specification of a discrepancy statistic, ideally one that does not involve unknown parameters. In this study, we use the discrepancy statistic, $T(\mathbf{x}, \mathbf{y}) = \min(\text{median}(\mathbf{x}), \text{median}(\mathbf{y}))$. Naturally, other discrepancy measures may also be considered, such as $T(\mathbf{x}, \mathbf{y}) = \max(\mathbf{x}, \mathbf{y})$ or $T(\mathbf{x}, \mathbf{y}) = s_x + s_y$, where s_x and s_y denote the sample variances. An algorithm for computing the posterior predictive p -value is provided below.

The posterior predictive p -value is then approximated by

$$p\text{-value} = \frac{\#\{T(\mathbf{x}^{\text{rep}}, \mathbf{y}^{\text{rep}}) \geq T(\mathbf{x}, \mathbf{y})\}}{N}.$$

6. Real Data Analysis

In this section, we consider two real data sets to illustrate the proposed methodology. The first data set relates to post-weld treatment methods designed to improve the fatigue life of welded joints, while the second consists of recorded lifetimes of steel specimens tested under varying stress levels. Both data sets are analyzed using the estimation methods developed in this article.

6.1. Data I

The first data set concerns post-weld treatment methods aimed at enhancing the fatigue life of welded joints. It was initially investigated by [25] and later revisited in [26]. The primary goal of these studies was to identify effective treatment techniques that could be applied in large-scale production of crane components. Two post-weld treatments are considered: burr grinding (BG) and TIG dressing, with their performance compared against the untreated, as-welded (AW) condition. The data set provides an opportunity to assess whether BG and TIG dressing improve fatigue strength relative to AW. In our analysis, since we are interested in a stress–strength setup, we treat the number of cycles to failure of specimens under BG, denoted by X , as the stress variable, while the cycles to failure under TIG, denoted by Y , represent the strength variable. The observed values are $X = 25, 25, 39, 44, 100, 72, 102, 74, 76, 144, 172$ and $Y = 84, 93, 156, 352, 666, 91, 112, 179, 136, 36, 94, 48, 44$. It is worth noting that this data set was also reported in [26], where the authors further established that the data follow geometric distributions.

We employ Algorithm 1 to generate samples from the posterior distribution based on the specified prior assumptions and the observed data. The posterior variance-covariance matrix is approximated using the Laplace method, and with a chosen scale parameter of $c = 1.5$, the resulting acceptance rate is approximately 41%, which falls within the recommended range of 10–40% reported by [22]. For initialization of the random walk Metropolis–Hastings algorithm, we use the maximum likelihood estimates of ρ and ω . The Markov chain is executed for 5500 iterations, discarding the first 500 as burn-in, with the remaining draws utilized to compute Bayes estimates under squared error loss (SELF) for the priors introduced in this work. As shown in Figure 1, the conditional posterior

distribution of ρ is well approximated by a Gaussian distribution. Furthermore, graphical diagnostics, including trace plots and auto-correlation function (ACF) plots of ρ under the different prior settings, indicate good mixing: the trace plots exhibit random fluctuation around the mean values (denoted by solid red line), while the ACFs decay reasonably toward zero, suggesting minimal auto-correlation. To further check the convergence of the MCMC, we run three chains and compute the Gelman–Rubin statistics ($\hat{R}$); a value close to 1 or below 1.05 indicates convergence [27]. We found that $\hat{R}$ is 1.0003 for π_J , 1.0012 for π_{R_1} , 1.0007 for π_{R_2} , and 1.0008 for π_M . Clearly, the Gelman–Rubin statistics under the four priors are well below the threshold of 1.05, confirming convergence. To further examine the suitability of the prior choices, we implement Algorithm 2, presenting histograms and kernel density estimates of the discrepancy statistics in Figure 2. The results demonstrate that replications from the Bayesian predictive density closely align with the observed data. In addition, the posterior predictive p -values, 0.27 for π_J , 0.28 for π_{R_1} , 0.30 for π_{R_2} , and 0.34 for π_M , confirm the compatibility of the proposed priors.

Algorithm 2 Posterior Predictive Assessment

- 1: Input observed data $\mathbf{x} = (x_1, \dots, x_n)$ and $\mathbf{y} = (y_1, \dots, y_m)$.
- 2: Compute the discrepancy measure

$$T(\mathbf{x}, \mathbf{y}) = \min(\text{median}(\mathbf{x}), \text{median}(\mathbf{y})).$$

- 3: **for** $\ell = 1, \dots, N$ **do**
- 4: Sample $\theta_\ell = (\theta_{1_\ell}, \theta_{2_\ell})$ from the posterior distribution (given prior assumptions and observed data $\mathbf{x}, \mathbf{y}$).
- 5: Generate replicated samples of the same size as the observed data:

$$\mathbf{x}^{\text{rep}} \sim \text{Geo}(\theta_{1_\ell}), \quad \mathbf{y}^{\text{rep}} \sim \text{Geo}(\theta_{2_\ell}).$$

- 6: Compute the replicated discrepancy measure

$$T(\mathbf{x}^{\text{rep}}, \mathbf{y}^{\text{rep}}).$$

- 7: **end for**

- 8: Repeat Steps 3–6 for a sufficiently large number N of posterior samples.
-

Before presenting the Bayes estimates of ρ obtained under different priors, we first provide the empirical estimate, defined as

$$\hat{\rho} = \frac{1}{nm} \sum_{i=1}^n \sum_{j=1}^m \mathbb{I}(X_i \leq Y_j).$$

Table 2 reports all estimates of ρ together with their standard deviations (SDs) and 95% HPD credible intervals. Among them, the Bayes estimate based on π_{R_2} outperforms the others by achieving the smallest SD, the narrowest 95% credible interval, and remaining closest to the empirical estimate. The Bayes estimate under π_M ranks second in terms of proximity to the empirical estimate. Furthermore, the Bayes estimate using π_{R_1} demonstrates better performance than both the Bayes estimate under π_J and the ML estimate, as it yields a smaller SD and a shorter 95% credible interval.

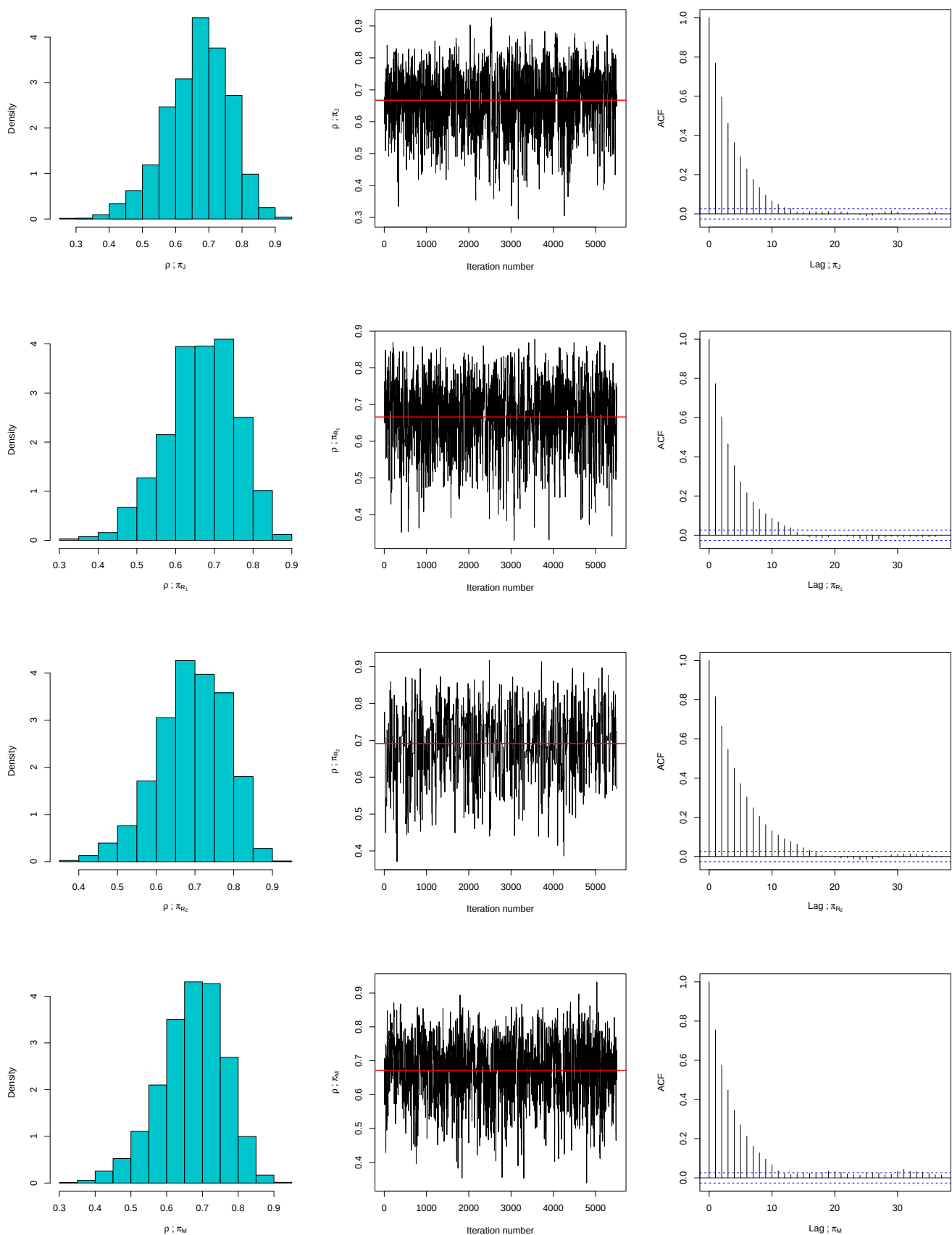


Figure 1. Diagnostic plots of the random walk MH algorithm for ρ under π_J , π_{R_1} , π_{R_2} , and π_M for the first data set.

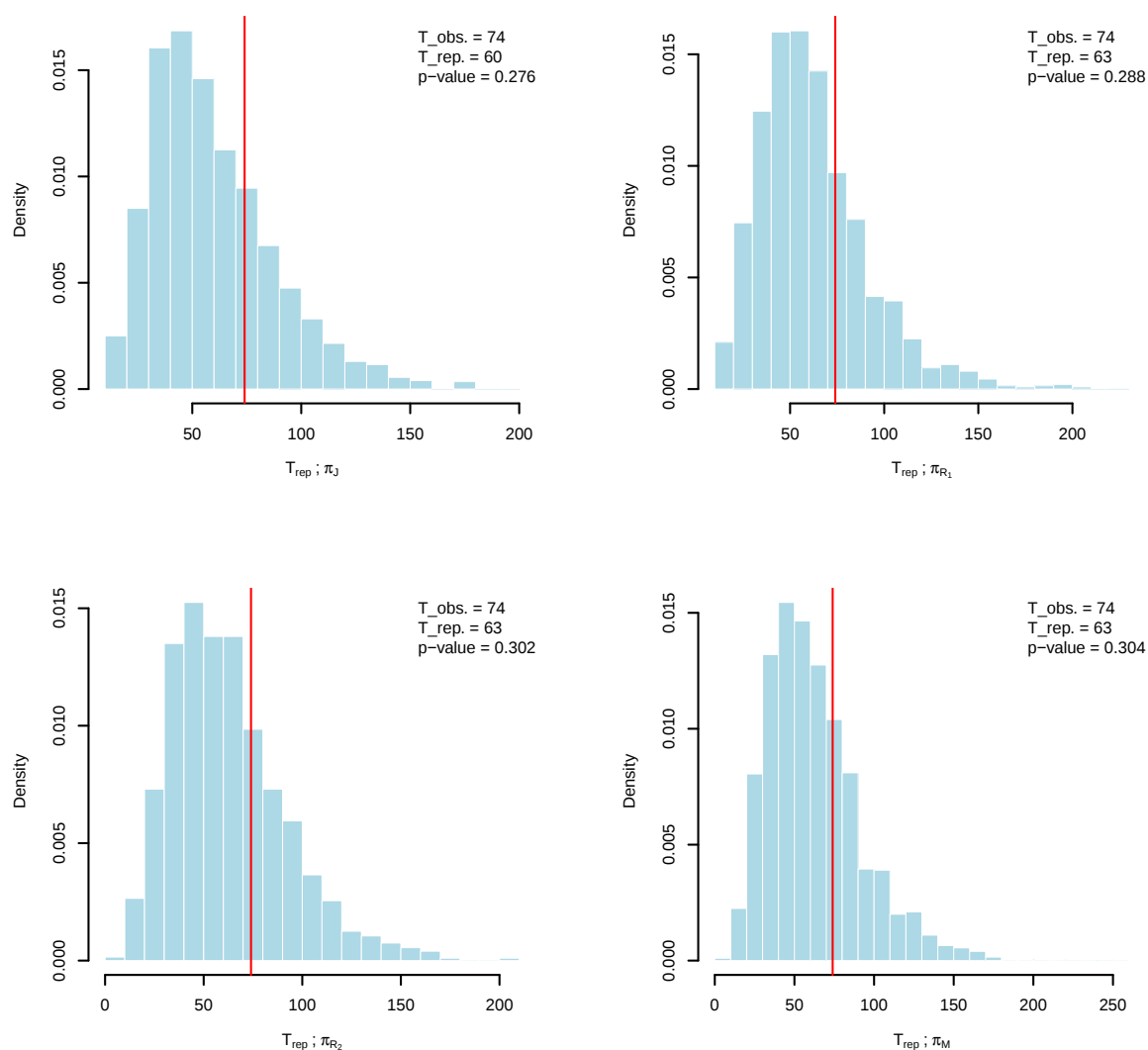


Figure 2. Histogram and kernel density of $T(x, y)$ along with its mean (red line) using π_J , π_{R_1} , π_{R_2} , and π_M (clockwise from the top left) for the first data set.

Table 2. The summary of Bayesian estimates of ρ for the first data.

Priors	Estimator	SD	95% CI
Empirical	0.6923	0.1120	(0.4726, 0.9120)
ML	0.6724	0.0916	(0.4929, 0.8519)
π_J	0.6672	0.0954	(0.4543, 0.8345)
π_{R_1}	0.6661	0.0914	(0.4721, 0.8280)
π_{R_2}	0.6912	0.0878	(0.4966, 0.8333)
π_M	0.6718	0.0906	(0.4705, 0.8297)

6.2. Data-II

The second data set comprises the recorded lifetimes of steel specimens tested under 14 different stress levels. These data were originally reported by [28] and subsequently analyzed in greater detail by [26,29]. To illustrate our theoretical results, we focus on the lifetimes corresponding to a stress level of 32, which we denote by the stress variable X , with the following observed values: 1144, 231, 523, 474, 4510, 3107, 815, 6297, 1580, 605, 1786, 206, 1943, 935, 283, 1336, 727, 370, 1056, 413, 619, 2214, 1826, and 597. Similarly, the lifetimes measured at a stress level of 32.5 are treated as the strength variable Y , given by the following values: 4257, 879, 799, 1388, 271, 308, 2073, 227, 347, 669, 1154, 393, 250, 196, 548, 475, 1705, 2211, 975, and 2925. It is worth noting that the data set can be reasonably modeled using the geometric distribution, as confirmed by [26].

In a similar manner, Algorithms 1 and 2 are employed to draw posterior samples under the assumed priors and observed data, and to evaluate the adequacy of these priors, respectively. As illustrated in Figure 3, the posterior distribution of ρ is well captured by a symmetric Gaussian proposal. Diagnostic plots further support this conclusion: the trace plots display stable fluctuations around the posterior means (marked by solid lines), and the corresponding autocorrelation functions (ACFs) decrease toward zero, indicating efficient mixing and low serial dependence. To assess prior suitability more thoroughly, Algorithm 2 is applied, with histograms and kernel density estimates of the discrepancy statistics shown in Figure 4. These results reveal that replicated samples from the Bayesian predictive distribution align closely with the observed data. Similarly, we run three MCMC chains and compute the Gelman–Rubin statistics. We found that $\hat{R}$ is 1.0016 for π_J , 1.0003 for π_{R_1} , 1.0013 for π_{R_2} , and 1.00041 for π_M . Since all values of the statistic are well below the threshold of 1.05, the convergence of the MCMC is confirmed. Additionally, the posterior predictive p -values, 0.42 for π_J , 0.406 for π_{R_1} , 0.444 for π_{R_2} , and 0.412 for π_M , provide further evidence in favor of the compatibility of the proposed priors.

Similarly, Table 3 presents the estimates of ρ along with their standard deviations (SDs) and 95% HPD credible intervals. The ML and the Bayes estimates are relatively close to empirical estimate. On the other hand, the Bayes estimates obtained under π_J and π_M perform best, yielding the smallest SDs, the narrowest credible intervals, and the closest agreement with the empirical estimate. Furthermore, the Bayes estimates based on π_{R_1} and π_{R_2} also outperform the ML estimate, as they are associated with smaller SDs and shorter intervals.

Table 3. The summary of Bayesian estimates of ρ for the second data.

Priors	Estimator	SD	95% CI
Empirical	0.5350	0.0935	(0.3518, 0.7182)
ML	0.5484	0.0891	(0.3739, 0.7228)
π_J	0.5428	0.0766	(0.3898, 0.6880)
π_{R_1}	0.5432	0.0805	(0.3819, 0.6887)
π_{R_2}	0.5552	0.0888	(0.4079, 0.7103)
π_M	0.5429	0.0798	(0.3824, 0.6954)

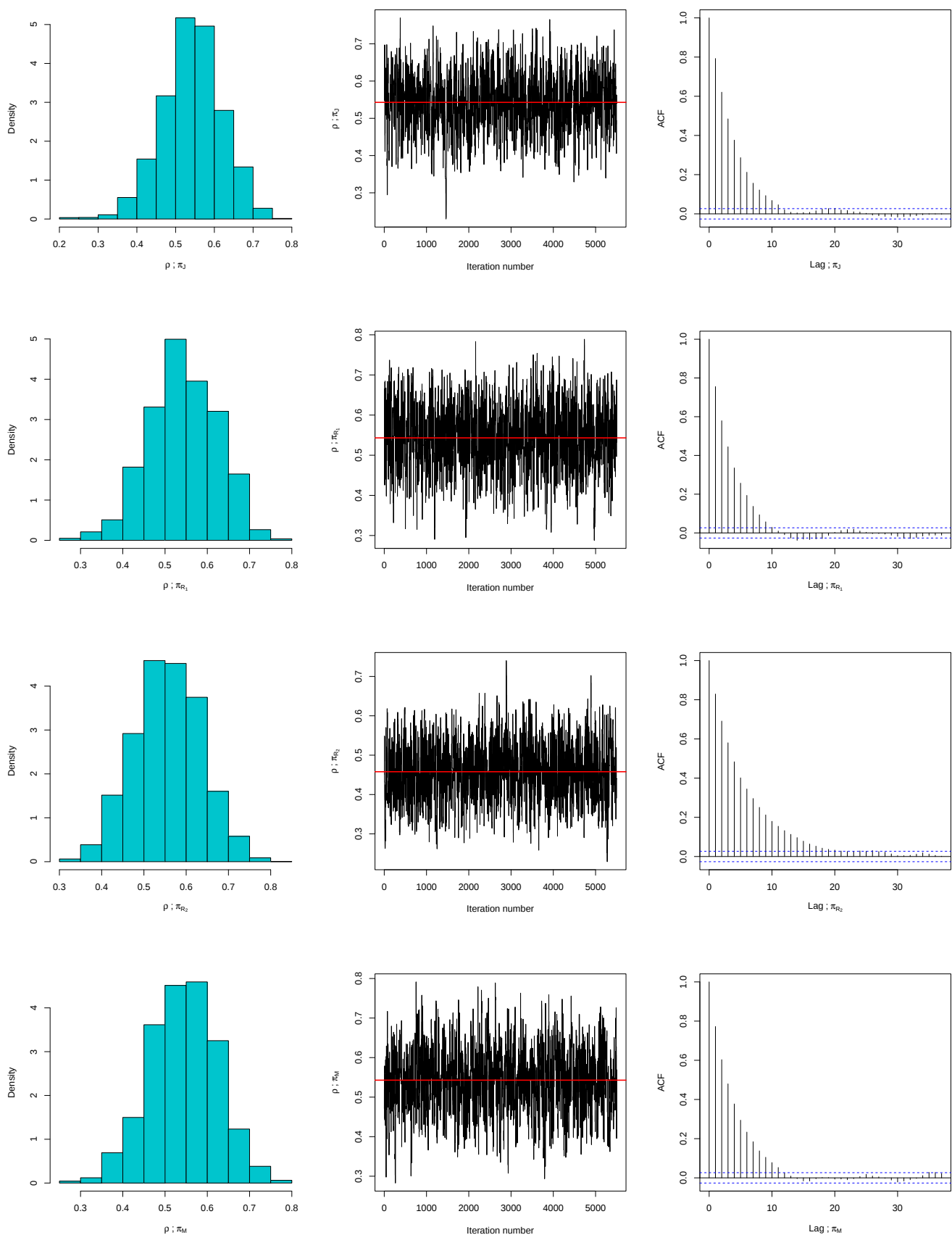


Figure 3. Diagnostic plots of the random walk MH algorithm for ρ under π_J , π_{R_1} , π_{R_2} , and π_M for the second data set.

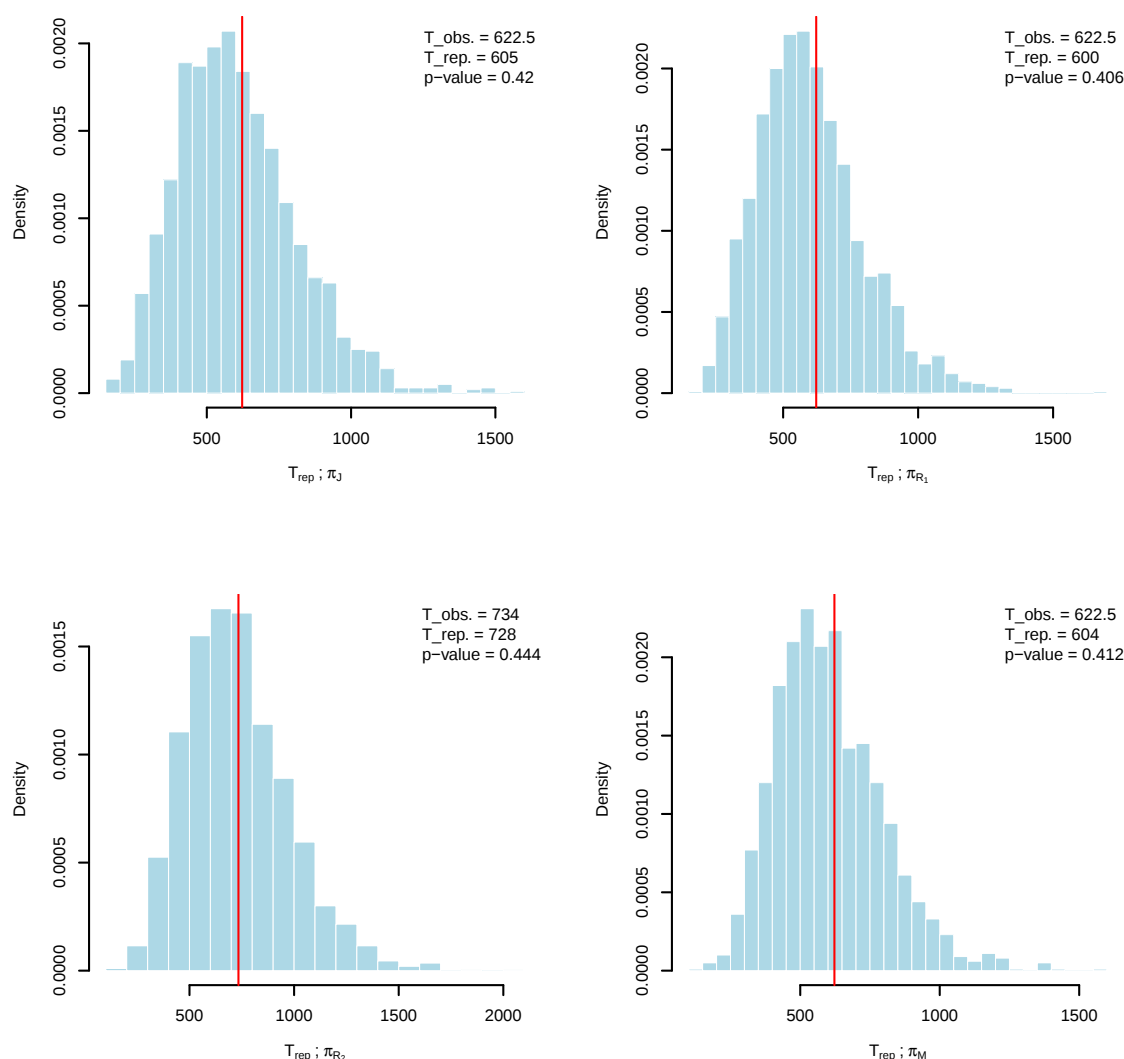


Figure 4. Histogram and kernel density of $T(x, y)$ along with its mean (red line) using π_J , π_{R_1} , π_{R_2} , and π_M (clockwise from the top left) for the second data set.

7. Conclusions

In this paper, we investigated maximum likelihood (ML) and Bayesian estimation procedures for the stress–strength parameter of the geometric distribution. We derived the ML estimator of ρ and established its asymptotic distribution. For the Bayesian approach, we developed objective priors for ρ , namely Jeffreys, approximate reference, and matching priors. A key advantage of the Jeffreys prior is its invariance: the Bayes estimator of ρ remains the same whether it is obtained directly in terms of ρ or via a one-to-one transformation of the original model parameters. By contrast, reference and matching priors cannot in general be derived from those of the original parameters, since ρ is a function of θ_1 and θ_2 . This underscores the importance of constructing these priors directly in terms of ρ , ensuring that invariance is preserved upon transformation back to the original parameters.

We also examined the posterior distribution under these priors, noting that they are improper but lead to valid posteriors. Since closed-form Bayes estimates under the squared error loss function (SELF) are not available, we employed MCMC methods to compute them. Simulation studies demonstrated that the proposed Bayes estimators perform very well, with the matching prior generally outperforming both the other Bayesian approaches.

and the ML estimator, making it particularly suitable for practical use. Finally, two real data applications confirmed these findings: Bayesian estimates consistently outperformed ML estimates, with the matching prior yielding the best results overall, although the reference prior ρ_{R_2} provided competitive performance in the first data set.

As in any study, a limitation of this work is that the proposed noninformative priors, including Jeffreys, reference, and matching priors, for the geometric stress–strength reliability may suffer from stability issues when the parameters lie close to the boundaries of the parameter space. This motivates us to further investigate such cases in future research. In addition, the methodology developed in this paper can be extended to other settings, such as record values or type-II censored data, which are frequently encountered in real applications, particularly in hydrology and lifetime studies. Under these two schemes, the likelihood remains tractable, making the use of the developed priors feasible.

Funding: This research received no external funding.

Data Availability Statement: The original contributions presented in this study are included in the article. Further inquiries can be directed to the corresponding author.

Acknowledgments: The author thanks the editor and the three reviewers for their constructive comments and suggestions, which greatly enhanced the earlier revision of this article.

Conflicts of Interest: The author declares that there are no conflicts of interest.

References

1. Kundu, D.; Gupta, R.D. Estimation of $P(Y < X)$ for generalized exponential distribution. *Metrika* **2005**, *61*, 291–308.
2. Kumar, I.; Kumar, K. On estimation of $P(V < U)$ for inverse pareto distribution under progressively censored data. *Int. J. Syst. Assur. Eng. Manag.* **2022**, *13*, 189–202.
3. Khan, M.; Khatoon, B. Statistical inferences of $R = P(X < Y)$ for exponential distribution based on generalized order statistics. *Ann. Data Sci.* **2020**, *7*, 525–545.
4. Yadav, A.S.; Singh, S.; Singh, U. Bayesian estimation of stress–strength reliability for lomax distribution under type-II hybrid censored data using asymmetric loss function. *Life Cycle Reliab. Saf. Eng.* **2019**, *8*, 257–267. [[CrossRef](#)]
5. Sun, D.; Ghosh, M.; Basu, A.P. Bayesian analysis for a stress-strength system under noninformative priors. *Can. J. Stat.* **1998**, *26*, 323–332.
6. Kang, S.G.; Lee, W.D.; Kim, Y. Objective bayesian analysis for generalized exponential stress–strength model. *Comput. Stat.* **2021**, *36*, 2079–2109. [[CrossRef](#)]
7. Abbas, K.; Tang, Y. Objective bayesian analysis of the Fréchet stress–strength model. *Stat. Probab. Lett.* **2014**, *84*, 169–175. [[CrossRef](#)]
8. Barbiero, A. Inference on reliability of stress-strength models for Poisson data. *J. Qual. Reliab. Eng.* **2013**, *2013*, 530530. [[CrossRef](#)]
9. Obradovic, M.; Jovanovic, M.; Milošević, B.; Jevremović, V. Estimation of $P(X \leq Y)$ for geometric-poisson model. *Hacet. J. Math. Stat.* **2015**, *44*, 949–964.
10. Ahmad, K.E.; Fakhry, M.E.; Jaheen, Z.F. Bayes estimation of $P(Y > X)$ in the geometric case. *Microelectron. Reliab.* **1995**, *35*, 817–820. [[CrossRef](#)]
11. Maiti, S.S. Estimation of $P(X \leq Y)$ in the geometric case. *J. Indian Stat. Assoc.* **1995**, *33*, 87–91.
12. Mohamed, M. Inference for reliability and stress-strength for geometric distribution. *Sylwan* **2015**, *159*, 281–289.
13. Mohamed, M. Estimation of R for geometric distribution under lower record values. *J. Appl. Res. Technol.* **2020**, *18*, 368–375. [[CrossRef](#)]
14. Datta, G.S.; Ghosh, M. On the invariance of noninformative priors. *Ann. Stat.* **1996**, *24*, 141–159. [[CrossRef](#)]
15. Dong, G.; Shakhatareh, M.K.; He, D. Bayesian analysis for the shannon entropy of the lomax distribution using noninformative priors. *J. Stat. Comput. Simul.* **2024**, *94*, 1317–1338. [[CrossRef](#)]
16. Sen, P.K.; Singer, J.M. *Large Sample Methods in Statistics: An Introduction with Applications*; Chapman & Hall: London, UK, 1993.
17. Berger, J.O.; Bernardo, J.M. Ordered group reference priors with application to the multinomial problem. *Biometrika* **1992**, *79*, 25–37. [[CrossRef](#)]
18. Bernardo, J.M. Reference analysis. *Handb. Stat.* **2005**, *25*, 17–90.
19. Bernardo, J.M.; Smith, A.F. *Bayesian Theory*; John Wiley & Sons: Hoboken, NJ, USA, 2009; Volume 405.

20. Welch, B.L.; Peers, H. On formulae for confidence points based on integrals of weighted likelihoods. *J. R. Stat. Soc. Ser. B (Methodol.)* **1963**, *25*, 318–329. [[CrossRef](#)]
21. Datta, G.S.; Ghosh, J.K. On priors providing frequentist validity for Bayesian inference. *Biometrika* **1995**, *82*, 37–45. [[CrossRef](#)]
22. Neal, P.; Roberts, G. Optimal scaling for random walk metropolis on spherically constrained target densities. *Methodol. Comput. Appl. Probab.* **2008**, *10*, 277–297. [[CrossRef](#)]
23. Guttman, I. The use of the concept of a future observation in goodness-of-fit problems. *J. R. Stat. Soc. Ser. B (Methodol.)* **1967**, *29*, 83–100. [[CrossRef](#)]
24. Meng, X.-L. Posterior predictive p -values. *Ann. Stat.* **1994**, *22*, 1142–1160. [[CrossRef](#)]
25. Pedersen, M.M.; Mouritsen, O.Ø.; Hansen, M.R.; Andersen, J.G.; Wenderby, J. Comparison of post-weld treatment of high-strength steel welded joints in medium cycle fatigue. *Weld. World* **2010**, *54*, R208–R217. [[CrossRef](#)]
26. Nayal, A.S.; Singh, B.; Tyagi, A.; Chesneau, C. Classical and bayesian inferences on the stress-strength reliability $R = P[Y < X < Z]$ in the geometric distribution setting. *AIMS Math.* **2023**, *8*, 20679–20699.
27. Brooks, S.P.; Gelman, A. General Methods for Monitoring Convergence of Iterative Simulations. *J. Comput. Graph. Stat.* **1998**, *7*, 434–455. [[CrossRef](#)]
28. Kimber, A. Exploratory data analysis for possibly censored data from skewed distributions. *J. R. Stat. Soc. Ser. C Appl. Stat.* **1990**, *39*, 21–30. [[CrossRef](#)]
29. Crowder, M. Tests for a family of survival models based on extremes. In *Recent Advances in Reliability Theory*; Birkhäuser: Boston, MA, USA, 2000.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.