

Article

Generalized Fiducial Inference for the Stress–Strength Reliability of Generalized Logistic Distribution

Menghan Li ¹, Liang Yan ^{1,*}, Yaru Qiao ², Xia Cai ² and Khamis K. Said ³

¹ School of Mathematics and Statistics, Hebei University of Economics and Business, Shijiazhuang 050061, China; limenghan1209@163.com

² School of Science, Hebei University of Science and Technology, Shijiazhuang 050018, China; 18833180451@163.com (Y.Q.); caixiasjz@163.com (X.C.)

³ Department of Science, Karume Institute of Science and Technology, Zanzibar P.O. Box 467, Tanzania; khakhasa@kist.ac.tz

* Correspondence: yanliang@hueb.edu.cn

Abstract: Generalized logistic distribution, as the generalized form of the symmetric logistic distribution, plays an important role in reliability analysis. This article focuses on the statistical inference for the stress–strength parameter $R = P(Y < X)$ of the generalized logistic distribution with the same and different scale parameters. Firstly, we use the frequentist method to construct asymptotic confidence intervals, and adopt the generalized inference method for constructing the generalized point estimators as well as the generalized confidence intervals. Then the generalized fiducial method is applied to construct the fiducial point estimators and the fiducial confidence intervals. Simulation results demonstrate that the generalized fiducial method outperforms other methods in terms of the mean square error, average length, and empirical coverage. Finally, three real datasets are used to illustrate the proposed methods.

Keywords: generalized fiducial inference; stress–strength; generalized logistic distribution; point estimation; interval estimation



Citation: Li, M.; Yan, L.; Qiao, Y.; Cai, X.; Said, K.K. Generalized Fiducial Inference for the Stress–Strength Reliability of Generalized Logistic Distribution. *Symmetry* **2023**, *15*, 1365. <https://doi.org/10.3390/sym15071365>

Academic Editors: Xinmin Li, Daojiang He and Weizhong Tian

Received: 14 June 2023

Revised: 1 July 2023

Accepted: 3 July 2023

Published: 5 July 2023



Copyright: © 2023 by the authors. 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/>).

1. Introduction

The stress–strength, which was initially proposed by Birnbaum [1] and developed by Birnbaum and McCarty [2], plays an important role in reliability analysis. For two independent random variables, X and Y , the stress–strength parameter is defined as $R = P(Y < X)$. If stress Y is greater than strength X , it may result in component failure or system malfunction. The stress–strength parameter is originally used in the industrial field to calculate the reliability of the products [3,4]. It is also increasingly used to estimate the probability that one variable exceeds another [5,6], which is of great significance in practical application and has been widely used in various fields, such as electrical cable failure analysis, leukemia treatment, and jute fiber testing. See more details for [5,7–10].

In the literature, there are many life distributions that can be used to estimate R , such as Weibull [5], Pareto [6,11], generalized Pareto [12], exponential [8,13], generalized exponential [14], Lomax [15], unit-half-normal [16], unit-Gompertz [17], and generalized logistic (GL) [18–21] distributions. The logistic distribution is a symmetric heavy-tailed distribution. However, it is not suitable for handling asymmetric or thin-tailed data. Therefore, it is necessary to further extend the logistic distribution according to practical problems, which can handle the data including symmetric, heavy-tailed, asymmetric, and thin-tailed. The GL distribution, as defined by Balakrishnan and Leung [22], is one of the generalized forms of the standard logistic distribution. By introducing a shape parameter to the distribution, the GL distribution expands the range of values for the skewness coefficient and tail index, which allows a wider range of data fitting capabilities. It has attracted extensive attention and is widely used in various fields, including demography,

biology, finance, and neural network, as detailed in [23]. Therefore, we select the GL distribution with the following probability density function (PDF)

$$f(x; \sigma, \alpha) = \frac{\alpha e^{-\frac{x}{\sigma}}}{\sigma(1 + e^{-\frac{x}{\sigma}})^{\alpha+1}}, \quad -\infty < x < \infty, \quad (1)$$

and the corresponding cumulative distribution function (CDF) is

$$F(x; \sigma, \alpha) = \frac{1}{(1 + e^{-\frac{x}{\sigma}})^{\alpha}}, \quad -\infty < x < \infty, \quad (2)$$

where $\sigma > 0$ and $\alpha > 0$ are the scale and shape parameters, respectively. The GL distribution exhibits a negative skew when $\alpha > 1$ and a positive skew when $0 < \alpha < 1$, and it becomes the standard logistic distribution (it is symmetric) when $\alpha = 1$. Meanwhile, the PDF of the GL distribution is unimodal and log-concave, making it suitable for modeling data with both left and right skewness [18]. The expectation and variance of X can be calculated from the moment-generating function of the GL distribution [24]; that is,

$$E(X) = \sigma(\psi(\alpha) - \psi(1)) \quad \text{and} \quad \text{Var}(X) = \sigma^2(\psi'(1) + \psi'(\alpha)), \quad (3)$$

where $\psi(z) = \Gamma'(z)/\Gamma(z)$ is the digamma function and $\psi'(z) = d\psi(z)/dz$ is the trigamma function, with the gamma function $\Gamma(z) = \int_0^\infty x^{z-1}e^{-x}dx$ for $z > 0$. From Formula (3), the coefficient of skewness for X , corresponding to the third standardized moment, is expressed as

$$\text{Skew}(X) = \frac{\psi''(\alpha) - \psi''(1)}{(\psi'(\alpha) + \psi'(1))^{3/2}}, \quad (4)$$

which implies that the expression does not depend on the parameter σ .

Statisticians have conducted numerous kinds of research on R based on the GL distribution, most focus on frequentist and Bayesian inference. For the single component, Asgharzadeh et al. [18], Babayi et al. [19], and Okasha [20] considered the estimation almost at the same time. Asgharzadeh et al. [18] considered the estimation of R for GL distribution under three different cases, and obtained the estimators and confidence intervals based on maximum likelihood (ML), bootstrap, and Bayesian methods. Babayi et al. [19] used ML and Bayes methods to obtain the point estimations and confidence intervals of R for GL distribution with the same and different scale parameters. When the scale parameters were the same, Okasha [20] obtained the point and interval estimations of R using ML and Bayes methods. For the multicomponent stress–strength reliability, Rasekhi et al. [21] discussed the point and interval estimations under Bayesian and ML methods.

Based on the above research, it was found that the empirical coverage of ML estimation sometimes fails to reach the nominal level, while the choice of the prior distribution is improper or subjective in Bayesian inference. Furthermore, Tao [25] stated that the Jeffreys prior and reference prior are improper in the GL distribution, which leads to the improper posterior distribution of the parameter. When the exact pivotal quantity is not available, the generalized inference (GI) proposed by Weerahandi [26] provides us with another way of thinking, and Wang et al. [27] have successfully estimated the R of the generalized exponential distribution based on the GI method. Moreover, Hannig et al. [28] stated that the posterior of generalized fiducial distribution (GFD) is always proper and the confidence intervals of generalized fiducial inference (GFI) intend to maintain stated coverage (or be conservative) while having an average length comparable to or shorter than other methods. Yan and Liu [29], Yan et al. [30], and Cai et al. [31] used the above fiducial approach to consider the estimation of the parameters of the generalized exponential distribution, Lomax distribution, and Weibull distribution, respectively, where GFI often provides better estimation results than the traditional methods. See [32,33] for more applications of the GFI method. For the above reasons, the research objective of this article is to find a more appropriate method among the existing methods to estimate the stress–strength of the GL distribution with the same and different scale parameters. Our original contribution

is mainly to introduce the GI and GFI methods to the estimation of R and compare their performance with the frequentist method. Furthermore, we show the advantages of the GFI method in terms of mean square error, average length, and empirical coverage.

The structure of the rest paper is as follows. For the hypothesis of the same and different scale parameters, Sections 2 and 3 develop the point and interval estimations of R based on the ML, GI, and GFI methods. Section 4 simulates and compares the above methods. Section 5 demonstrates the proposed estimations by providing three real data examples. The implications of our findings are discussed in Section 6. The conclusions based on the research results are drawn in Section 7.

2. Estimation of R with the Same Scale and Different Shape Parameters

Suppose $X \sim GL(\sigma, \alpha_1)$ and $Y \sim GL(\sigma, \alpha_2)$ are independent random variables with the same scale parameter σ , then $R = P(Y < X)$ can be calculated as follows

$$\begin{aligned} R = P(Y < X) &= \int_{-\infty}^{\infty} \int_{-\infty}^x f_Y(y) f_X(x) dy dx \\ &= \int_{-\infty}^{\infty} \frac{1}{(1 + e^{-\frac{x}{\sigma}})^{\alpha_2}} \cdot \frac{\alpha_1 e^{-\frac{x}{\sigma}}}{\sigma (1 + e^{-\frac{x}{\sigma}})^{\alpha_1+1}} dx \\ &= \frac{\alpha_1}{\alpha_1 + \alpha_2}. \end{aligned} \quad (5)$$

2.1. Maximum Likelihood Estimation of R

Given the observed data, $\mathbf{x} = (x_1, \dots, x_n)^T$ and $\mathbf{y} = (y_1, \dots, y_m)^T$, the log-likelihood function of GL distribution is

$$\begin{aligned} L(\sigma, \alpha_1, \alpha_2 | \mathbf{x}, \mathbf{y}) &= n \log \alpha_1 + m \log \alpha_2 - (n + m) \log \sigma - \frac{\sum_{i=1}^n x_i + \sum_{j=1}^m y_j}{\sigma} \\ &\quad - (\alpha_1 + 1) \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}}) - (\alpha_2 + 1) \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}}). \end{aligned} \quad (6)$$

The corresponding ML estimators of σ , α_1 , and α_2 can be derived from

$$\begin{aligned} \frac{\partial L}{\partial \sigma} &= -\frac{n + m}{\sigma} + \frac{\sum_{i=1}^n x_i + \sum_{j=1}^m y_j}{\sigma^2} \\ &\quad - (\alpha_1 + 1) \sum_{i=1}^n \frac{x_i e^{-\frac{x_i}{\sigma}}}{\sigma^2 (1 + e^{-\frac{x_i}{\sigma}})} - (\alpha_2 + 1) \sum_{j=1}^m \frac{y_j e^{-\frac{y_j}{\sigma}}}{\sigma^2 (1 + e^{-\frac{y_j}{\sigma}})} = 0, \end{aligned} \quad (7)$$

$$\frac{\partial L}{\partial \alpha_1} = \frac{n}{\alpha_1} - \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}}) = 0, \quad (8)$$

$$\frac{\partial L}{\partial \alpha_2} = \frac{m}{\alpha_2} - \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}}) = 0. \quad (9)$$

From Formulas (8) and (9), the ML estimators of α_1 and α_2 as the functions of σ , say $\hat{\alpha}_1(\sigma)$ and $\hat{\alpha}_2(\sigma)$, respectively, can be obtained as

$$\hat{\alpha}_1(\sigma) = \frac{n}{\sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}})} \quad \text{and} \quad \hat{\alpha}_2(\sigma) = \frac{m}{\sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}})}. \quad (10)$$

From Formula (7), the ML estimator of σ can be determined by the following nonlinear equation

$$h(\sigma) = \sigma, \quad (11)$$

where

$$h(\sigma) = \frac{1}{n+m} \left[\sum_{i=1}^n x_i + \sum_{j=1}^m y_j - \left(\frac{n}{\sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}})} + 1 \right) \sum_{i=1}^n \frac{x_i e^{-\frac{x_i}{\sigma}}}{1 + e^{-\frac{x_i}{\sigma}}} \right. \\ \left. - \left(\frac{m}{\sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}})} + 1 \right) \sum_{j=1}^m \frac{y_j e^{-\frac{y_j}{\sigma}}}{1 + e^{-\frac{y_j}{\sigma}}} \right].$$

Since $\hat{\sigma}$ is a fixed point solution of Equation (11), it can be obtained by the iterative algorithm $h(\sigma_{(k)}) = \sigma_{(k+1)}$, where $\sigma_{(k)}$ is the k th iteration of $\hat{\sigma}$. The iteration procedure will stop when $|\sigma_{(k)} - \sigma_{(k+1)}|$ is small enough. Substituting $\hat{\sigma}$ into (10), we can have $\hat{\alpha}_1$ and $\hat{\alpha}_2$. Accordingly, the ML estimator of R is

$$\hat{R}_{ML} = \frac{\hat{\alpha}_1}{\hat{\alpha}_1 + \hat{\alpha}_2}, \quad (12)$$

where $\hat{\alpha}_1$ and $\hat{\alpha}_2$ are the ML estimators of α_1 and α_2 , respectively.

The confidence interval of R can be derived by the following asymptotic distribution; that is,

$$(\hat{\sigma}, \hat{\alpha}_1, \hat{\alpha}_2)^T \xrightarrow{L} N\left((\sigma, \alpha_1, \alpha_2)^T, I_0^{-1}\right), \quad (13)$$

where I_0 is the observed Fisher information matrix, i.e.,

$$I_0^{-1} = \left(\begin{array}{ccc} -\frac{\partial^2 L}{\partial \sigma^2} & -\frac{\partial^2 L}{\partial \sigma \partial \alpha_1} & -\frac{\partial^2 L}{\partial \sigma \partial \alpha_2} \\ -\frac{\partial^2 L}{\partial \alpha_1 \partial \sigma} & -\frac{\partial^2 L}{\partial \alpha_1^2} & -\frac{\partial^2 L}{\partial \alpha_1 \partial \alpha_2} \\ -\frac{\partial^2 L}{\partial \alpha_2 \partial \sigma} & -\frac{\partial^2 L}{\partial \alpha_2 \partial \alpha_1} & -\frac{\partial^2 L}{\partial \alpha_2^2} \end{array} \right)^{-1} \Big|_{(\sigma, \alpha_1, \alpha_2)^T = (\hat{\sigma}, \hat{\alpha}_1, \hat{\alpha}_2)^T} \\ \triangleq \begin{pmatrix} \text{Var}(\hat{\sigma}) & \text{Cov}(\hat{\sigma}, \hat{\alpha}_1) & \text{Cov}(\hat{\sigma}, \hat{\alpha}_2) \\ \text{Cov}(\hat{\alpha}_1, \hat{\sigma}) & \text{Var}(\hat{\alpha}_1) & \text{Cov}(\hat{\alpha}_1, \hat{\alpha}_2) \\ \text{Cov}(\hat{\alpha}_2, \hat{\sigma}) & \text{Cov}(\hat{\alpha}_2, \hat{\alpha}_1) & \text{Var}(\hat{\alpha}_2) \end{pmatrix}.$$

By using the Delta method, the asymptotic variance of $\hat{R}_{ML}$ is given by

$$\text{Var}(\hat{R}_{ML}) = \left(\frac{\partial R}{\partial \sigma}, \frac{\partial R}{\partial \alpha_1}, \frac{\partial R}{\partial \alpha_2} \right) I_0^{-1} \left(\frac{\partial R}{\partial \sigma}, \frac{\partial R}{\partial \alpha_1}, \frac{\partial R}{\partial \alpha_2} \right)^T. \quad (14)$$

Consequently, the $100(1 - \gamma)\%$ asymptotic confidence interval of R is

$$\hat{R}_{ML} - z_{1-\gamma/2} \sqrt{\text{Var}(\hat{R}_{ML})} \leq R \leq \hat{R}_{ML} + z_{1-\gamma/2} \sqrt{\text{Var}(\hat{R}_{ML})}, \quad (15)$$

where $z_{1-\gamma/2}$ is the $1 - \gamma/2$ quantile of standard normal distribution.

2.2. Generalized Inference of R

Since Wang et al. [27] successfully estimated the generalized exponential distribution by the GI method, we introduce the GI method to estimate R under GL distribution, which is formally similar to the generalized exponential distribution.

Lemma 1 (Wang et al. [27] and Yu et al. [34]). *Let $Z_1, \dots, Z_n$ be a random sample from the exponential distribution with mean θ and $Z_{(1)} < Z_{(2)} < \dots < Z_{(n)}$ be the corresponding order statistics. Let*

$$S_i = \sum_{j=1}^i Z_{(j)} + (n-i)Z_{(i)}, \quad i = 1, \dots, n, \\ T = 2 \sum_{i=1}^{n-1} \log(S_n/S_i).$$

Then (1) T and S_n are independent; (2) $T \sim \chi^2(2n-2)$ and $2S_n/\theta \sim \chi^2(2n)$.

Lemma 2. Let

$$f(\sigma) = \frac{\log(1 + e^{-\frac{b}{\sigma}})}{\log(1 + e^{-\frac{a}{\sigma}})},$$

where $b > a > 0$ are constants. Thus, $f(\sigma)$ is strictly decreasing on $(0, +\infty)$.

Proof of Lemma 2. From the function

$$f(\sigma) = \frac{\log(1 + e^{-\frac{b}{\sigma}})}{\log(1 + e^{-\frac{a}{\sigma}})}, \quad b > a > 0,$$

we can calculate that

$$f'(\sigma) = \frac{b(e^{\frac{a}{\sigma}} + 1) \log(1 + e^{-\frac{a}{\sigma}}) - a(e^{\frac{b}{\sigma}} + 1) \log(1 + e^{-\frac{b}{\sigma}})}{\sigma^2 (e^{\frac{a}{\sigma}} + 1)(e^{\frac{b}{\sigma}} + 1)(\log(1 + e^{-\frac{a}{\sigma}}))^2}.$$

It is obvious that the denominator of $f'(\sigma)$ is greater than 0, so we mainly focus on the numerator. Let

$$g(\sigma) = b(e^{\frac{a}{\sigma}} + 1) \log(1 + e^{-\frac{a}{\sigma}}) - a(e^{\frac{b}{\sigma}} + 1) \log(1 + e^{-\frac{b}{\sigma}}),$$

then

$$g'(\sigma) = \frac{ab}{\sigma^2} \left[\frac{\log(1 + e^{-\frac{b}{\sigma}})}{e^{-\frac{b}{\sigma}}} - \frac{\log(1 + e^{-\frac{a}{\sigma}})}{e^{-\frac{a}{\sigma}}} \right].$$

Because $\frac{\log(1+x)}{x}$ is strictly decreasing on $(0, +\infty)$ and the $e^{-\frac{x}{\sigma}}$ is strictly increasing in $x > 0$ for $\sigma > 0$, we have that $\frac{\log(1+e^{-\frac{x}{\sigma}})}{e^{-\frac{x}{\sigma}}}$ is strictly decreasing in $x > 0$ for $\sigma > 0$. Thus $g'(\sigma) < 0$ on $(0, +\infty)$ and $g(\sigma)$ is strictly decreasing on $(0, +\infty)$. Therefore, for $\sigma > 0$, we can obtain

$$g(\sigma) < \lim_{\sigma \rightarrow 0^+} g(\sigma) = 0.$$

Finally, we have that $f'(\sigma) < 0$ on $(0, +\infty)$; thus, $f(\sigma)$ is strictly decreasing on $(0, +\infty)$. $\square$

Let $X_1, \dots, X_n$ be a random sample from $GL(\sigma, \alpha_1)$ and $\mathbf{X} = (X_{(1)}, \dots, X_{(n)})^T$ be the corresponding order statistics. If a random variable X follows the standard uniform distribution, then $-\log X$ follows the standard exponential distribution $Exp(1)$. Obviously, we can find that $(1 + e^{-\frac{X_{(1)}}{\sigma}})^{-\alpha_1}, \dots, (1 + e^{-\frac{X_{(n)}}{\sigma}})^{-\alpha_1}$ are the order statistics from the standard uniform distribution. Thus, $Z_{(i)} = \log(1 + e^{-\frac{X_{(n-i+1)}}{\sigma}}), i = 1, \dots, n$ are the order statistics from the exponential distribution with mean $1/\alpha_1$. Similarly, we can easily obtain the order statistics $Z_{(j)} = \log(1 + e^{-\frac{Y_{(m-j+1)}}{\sigma}}), j = 1, \dots, m$ from the exponential distribution with mean $1/\alpha_2$, where the random sample $Y_1, \dots, Y_m$ follows $GL(\sigma, \alpha_2)$. Let

$$T(\sigma) = 2 \left(\sum_{i=1}^{n-1} \log(S_n/S_i) + \sum_{j=1}^{m-1} \log(S_m/S_j) \right), \quad (16)$$

where $S_i = \sum_{u=1}^i Z_{(u)} + (n-i)Z_{(i)}$ and $S_j = \sum_{v=1}^j Z_{(v)} + (m-j)Z_{(j)}$. Then from Lemma 1, we have $T(\sigma) \sim \chi^2(2n+2m-4)$. Together, Lemma 2 with

$$\frac{S_n}{S_i} = 1 + \frac{S_n - S_i}{S_i} = 1 + \frac{\frac{Z_{(i+1)}}{Z_{(i)}} + \dots + \frac{Z_{(n)}}{Z_{(i)}} - (n-i)}{\frac{Z_{(1)}}{Z_{(i)}} + \dots + \frac{Z_{(i-1)}}{Z_{(i)}} + (n-i+1)}, \quad (17)$$

$$\frac{S_m}{S_j} = 1 + \frac{S_m - S_j}{S_j} = 1 + \frac{\frac{Z_{(j+1)}}{Z_{(j)}} + \dots + \frac{Z_{(m)}}{Z_{(j)}} - (m-j)}{\frac{Z_{(1)}}{Z_{(j)}} + \dots + \frac{Z_{(j-1)}}{Z_{(j)}} + (m-j+1)}, \quad (18)$$

$T(\sigma)$ is strictly increasing on $(0, +\infty)$. Notice that

$$\lim_{\sigma \rightarrow 0^+} T(\sigma) = 0 \quad \text{and} \quad \lim_{\sigma \rightarrow +\infty} T(\sigma) = +\infty. \quad (19)$$

Therefore, equation $T(\sigma) = T$ has the unique solution $g(T, \mathbf{X}, \mathbf{Y})$ when n and m are given. The solution of equation $T(\sigma) = T$ can be obtained by the bisection method.

According to Lemma 1, we find that $U_1 = 2\alpha_1 S_n \sim \chi^2(2n)$, then $\alpha_1 = U_1/(2S_n)$. $\alpha_2 = U_2/(2S_m)$ can be calculated in the same way, so the generalized pivotal quantity is

$$R_{GI} = \frac{U_2/s_m}{U_1/s_n + U_2/s_m}, \quad (20)$$

where $s_n = \sum_{i=1}^n \log(1 + e^{-x_{(i)}/g(T, \mathbf{x}, \mathbf{y})})$, $s_m = \sum_{j=1}^m \log(1 + e^{-y_{(j)}/g(T, \mathbf{x}, \mathbf{y})})$, $\mathbf{x} = (x_{(1)}, \dots, x_{(n)})^T$ and $\mathbf{y} = (y_{(1)}, \dots, y_{(m)})^T$ are the observed values of $\mathbf{X} = (X_{(1)}, \dots, X_{(n)})^T$ and $\mathbf{Y} = (Y_{(1)}, \dots, Y_{(m)})^T$, respectively. The generalized point estimation and generalized confidence interval of R can be obtained by using the following algorithm.

1. Generate a realization t of T from $\chi^2(2n + 2m - 4)$. Then for given samples $\mathbf{x}$ and $\mathbf{y}$, one can obtain a realization of $g(T, \mathbf{x}, \mathbf{y})$ from the equation $T(\sigma) = t$.
2. Derive a realization of U_1 and U_2 from $\chi^2(2n)$ and $\chi^2(2m)$, respectively. Compute $\hat{R}_{GI}^{(1)}$ on the basis of (20).
3. Perform Step 1 and Step 2 for N times, iteratively. The value of N is equal to 1000.
4. The generalized point estimator of R is $\hat{R}_{GI} = \frac{1}{N} \sum_{l=1}^N \hat{R}_{GI}^{(l)}$. If $\hat{R}_{GI, \gamma/2}$ and $\hat{R}_{GI, 1-\gamma/2}$ denote the $\gamma/2$ and $1 - \gamma/2$ percentile of $\hat{R}_{GI}$, the generalized confidence interval of R is $[\hat{R}_{GI, \gamma/2}, \hat{R}_{GI, 1-\gamma/2}]$.

2.3. Generalized Fiducial Inference of R

Let the data-generating equation be

$$\mathbf{x} = \mathbf{G}(\mathbf{U}, \boldsymbol{\theta}), \quad (21)$$

where $\mathbf{x}$ denotes the data, $\boldsymbol{\theta}$ is the parameter vector, and $\mathbf{U}$ is a random vector with 0–1 uniform distribution $U(0, 1)$ in each dimension. Under some differentiability conditions, Hannig et al. [28] provided a user-friendly formula to compute the GFD of $\boldsymbol{\theta}$, i.e.,

$$f_F(\boldsymbol{\theta}) = \frac{f(\mathbf{x}|\boldsymbol{\theta})J(\mathbf{x}, \boldsymbol{\theta})}{\int f(\mathbf{x}|\boldsymbol{\theta})J(\mathbf{x}, \boldsymbol{\theta})d\boldsymbol{\theta}}, \quad (22)$$

where $f(\mathbf{x}|\boldsymbol{\theta})$ denotes the joint density function of $\mathbf{x}$ and the function $J(\mathbf{x}, \boldsymbol{\theta})$ is a Jacobian determinant. We usually take $J(\mathbf{x}, \boldsymbol{\theta})$ as the infinite norm as follows

$$J(\mathbf{x}, \boldsymbol{\theta}) = \text{Det} \left(\frac{d}{d\boldsymbol{\theta}} \mathbf{G}(\mathbf{u}, \boldsymbol{\theta}) \Big|_{\mathbf{u}=\mathbf{G}^{-1}(\mathbf{x}, \boldsymbol{\theta})} \right). \quad (23)$$

In practice, Hannig et al. [28] recommended using $\text{Det}(A) = \sum_{i=(i_1, \dots, i_p)} |\det(A)_i|$ and the

above sum goes over $\binom{n}{p}$ of p -tuples of indexes $\mathbf{i} = (1 \leq i_1 < \dots < i_p \leq n)$. For any $n \times p$ matrix A , the sub-matrix $(A)_i$ is the $p \times p$ matrix consisting of the rows $\mathbf{i} = (i_1, \dots, i_p)$ of A .

Regarding our concern, we have that

$$U_i = F(x_i; \boldsymbol{\theta}), \quad i = 1, \dots, n, \quad (24)$$

where $\theta = (\sigma, \alpha_1)^T$, $F(x_i; \sigma, \alpha_1) \triangleq (1 + e^{-\frac{x_i}{\sigma}})^{-\alpha_1}$ is the CDF of GL distribution and U_i follows $U(0, 1)$. According to (24), we can obtain the data generating equation $\mathbf{x} = \mathbf{G}(\mathbf{U}, \theta)$ and the i -th $x_i = G_i(u_i, \theta)$ is

$$x_i = -\sigma \log(U_i^{-\frac{1}{\alpha_1}} - 1). \quad (25)$$

Then, we have

$$\left. \frac{\partial G_i}{\partial \sigma} \right|_{u_i=(1+e^{-\frac{x_i}{\sigma}})^{-\alpha_1}} = \frac{x_i}{\sigma} \quad \text{and} \quad \left. \frac{\partial G_i}{\partial \alpha_1} \right|_{u_i=(1+e^{-\frac{x_i}{\sigma}})^{-\alpha_1}} = \frac{\sigma}{\alpha_1} (1 + e^{-\frac{x_i}{\sigma}}) \log(1 + e^{-\frac{x_i}{\sigma}}). \quad (26)$$

Substituting (26) into (23), it follows that

$$J(\mathbf{x}, \sigma, \alpha_1) = \frac{1}{\alpha_1} \sum_{i \neq j} \left| x_i (1 + e^{-\frac{x_j}{\sigma}}) \log(1 + e^{-\frac{x_j}{\sigma}}) - x_j (1 + e^{-\frac{x_i}{\sigma}}) \log(1 + e^{-\frac{x_i}{\sigma}}) \right|. \quad (27)$$

The function $J(\mathbf{y}, \sigma, \alpha_2)$ can be obtained through a similar process. Finally, we can derive the following GFD for $(\sigma, \alpha_1, \alpha_2)$; that is,

$$f_F(\sigma, \alpha_1, \alpha_2 | \mathbf{x}, \mathbf{y}) = \frac{f(\mathbf{x}, \mathbf{y} | \sigma, \alpha_1, \alpha_2) J(\mathbf{x}, \mathbf{y}, \sigma, \alpha_1, \alpha_2)}{\int_0^\infty \int_0^\infty \int_0^\infty f(\mathbf{x}, \mathbf{y} | \sigma, \alpha_1, \alpha_2) J(\mathbf{x}, \mathbf{y}, \sigma, \alpha_1, \alpha_2) d\sigma d\alpha_1 d\alpha_2}, \quad (28)$$

where

$$\begin{aligned} f(\mathbf{x}, \mathbf{y} | \sigma, \alpha_1, \alpha_2) &= f(\mathbf{x} | \sigma, \alpha_1) f(\mathbf{y} | \sigma, \alpha_2) = \prod_{i=1}^n f_i(x_i; \sigma, \alpha_1) \prod_{j=1}^m f_j(y_j; \sigma, \alpha_2) \\ &= \prod_{i=1}^n \frac{\alpha_1 e^{-\frac{x_i}{\sigma}}}{\sigma (1 + e^{-\frac{x_i}{\sigma}})^{\alpha_1 + 1}} \cdot \prod_{j=1}^m \frac{\alpha_2 e^{-\frac{y_j}{\sigma}}}{\sigma (1 + e^{-\frac{y_j}{\sigma}})^{\alpha_2 + 1}}, \\ J(\mathbf{x}, \mathbf{y}, \sigma, \alpha_1, \alpha_2) &= w_1 J(\mathbf{x}, \sigma, \alpha_1) + w_2 J(\mathbf{y}, \sigma, \alpha_2) \\ &= \frac{\binom{n}{2}}{\binom{n}{2} + \binom{m}{2}} J(\mathbf{x}, \sigma, \alpha_1) + \frac{\binom{m}{2}}{\binom{n}{2} + \binom{m}{2}} J(\mathbf{y}, \sigma, \alpha_2). \end{aligned}$$

Specifically,

$$\begin{aligned} f_F(\sigma, \alpha_1, \alpha_2 | \mathbf{x}, \mathbf{y}) &\propto \frac{\alpha_1^n e^{-\sum_{i=1}^n \frac{x_i}{\sigma}}}{\sigma^n \prod_{i=1}^n (1 + e^{-\frac{x_i}{\sigma}})^{\alpha_1 + 1}} \cdot \frac{\alpha_2^m e^{-\sum_{j=1}^m \frac{y_j}{\sigma}}}{\sigma^m \prod_{j=1}^m (1 + e^{-\frac{y_j}{\sigma}})^{\alpha_2 + 1}} \\ &\times \left[\frac{\binom{n}{2}}{\binom{n}{2} + \binom{m}{2}} \cdot \frac{1}{\alpha_1} \sum_{i \neq j} |q(x_i, x_j, \sigma)| \right. \\ &\left. + \frac{\binom{m}{2}}{\binom{n}{2} + \binom{m}{2}} \cdot \frac{1}{\alpha_2} \sum_{i \neq j} |q(y_i, y_j, \sigma)| \right], \quad (29) \end{aligned}$$

where

$$\begin{aligned} q(x_i, x_j, \sigma) &= x_i (1 + e^{-\frac{x_j}{\sigma}}) \log(1 + e^{-\frac{x_j}{\sigma}}) - x_j (1 + e^{-\frac{x_i}{\sigma}}) \log(1 + e^{-\frac{x_i}{\sigma}}), \\ q(y_i, y_j, \sigma) &= y_i (1 + e^{-\frac{y_j}{\sigma}}) \log(1 + e^{-\frac{y_j}{\sigma}}) - y_j (1 + e^{-\frac{y_i}{\sigma}}) \log(1 + e^{-\frac{y_i}{\sigma}}). \end{aligned}$$

On the one hand, the conditional fiducial density function of σ given α_1 and α_2 can be obtained from (29) and it is given by

$$f_F(\sigma|\alpha_1, \alpha_2, \mathbf{x}, \mathbf{y}) \propto \frac{e^{-\sum_{i=1}^n \frac{x_i}{\sigma}}}{\sigma^n \prod_{i=1}^n (1 + e^{-\frac{x_i}{\sigma}})^{\alpha_1+1}} \cdot \frac{e^{-\sum_{j=1}^m \frac{y_j}{\sigma}}}{\sigma^m \prod_{j=1}^m (1 + e^{-\frac{y_j}{\sigma}})^{\alpha_2+1}} \\ \times \left[\frac{\binom{n}{2}}{\binom{n}{2} + \binom{m}{2}} \cdot \sum_{i \neq j} |q(x_i, x_j, \sigma)| + \frac{\binom{m}{2}}{\binom{n}{2} + \binom{m}{2}} \cdot \sum_{i \neq j} |q(y_i, y_j, \sigma)| \right]. \quad (30)$$

On the other hand, we can obtain

$$f_F(\sigma, \alpha_1|\mathbf{x}) = \frac{f(\mathbf{x}|\sigma, \alpha_1)J(\mathbf{x}, \sigma, \alpha_1)}{\int_0^\infty \int_0^\infty f(\mathbf{x}|\sigma, \alpha_1)J(\mathbf{x}, \sigma, \alpha_1)d\sigma d\alpha_1} \\ \propto \frac{\alpha_1^n e^{-\sum_{i=1}^n \frac{x_i}{\sigma}}}{\sigma^n \prod_{i=1}^n (1 + e^{-\frac{x_i}{\sigma}})^{\alpha_1+1}} \cdot \frac{1}{\alpha_1} \sum_{i \neq j} |q(x_i, x_j, \sigma)|. \quad (31)$$

Therefore, the conditional fiducial density functions of α_1 given σ can be obtained as

$$f_F(\alpha_1|\sigma, \mathbf{x}) \propto \alpha_1^{n-1} e^{-\alpha_1 \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}})}, \quad (32)$$

similarly, the conditional fiducial density functions of α_2 given σ are

$$f_F(\alpha_2|\sigma, \mathbf{y}) \propto \alpha_2^{m-1} e^{-\alpha_2 \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}})}, \quad (33)$$

which implies that the conditional density of α_1 and α_2 are $Ga\left(n, \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}})\right)$ and $Ga\left(m, \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}})\right)$, respectively, where Ga stands for the Gamma distribution.

Using the Gibbs sampler to estimate the GFD requires being able to sample from the full conditional distribution for each quantity involved, so this is the case for α_1 and α_2 , but not for σ . Consequently, we introduce the standard Metropolis–Hastings steps into the Gibbs sampler to update σ in (30) while updating α_1 and α_2 from their exact conditional distribution. To reduce the autocorrelation of the Monte Carlo Markov Chains, we introduce a thin parameter T , which is an integer specifying the number of steps between each saved sample. The detailed steps of the algorithm are as follows.

1. Obtain three starting values of $\sigma^{(0)}$, $\alpha_1^{(0)}$ and $\alpha_2^{(0)}$.
2. Let $\sigma^{(l)}$, $\alpha_1^{(l)}$, and $\alpha_2^{(l)}$ denote the values of the l th iteration. Sample a candidate $\sigma^{(l+1)}$ from $f_F(\sigma|\alpha_1, \alpha_2, \mathbf{x}, \mathbf{y})$ by using the Metropolis–Hastings method [35]. Sample the candidate $\alpha_1^{(l+1)}$ and $\alpha_2^{(l+1)}$ from $Ga\left(n, \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma}})\right)$ and $Ga\left(m, \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma}})\right)$, respectively. $\hat{R}_{GFI}^{(l)}$ can be obtained by plugging the values of $\sigma^{(l)}$, $\alpha_1^{(l)}$ and $\alpha_2^{(l)}$ into Formula (5).
3. Conduct Step 2 for $M + TN$ times, iteratively, where M is the burn-in period. The values of M and N are both equal to 1000, and the value of T is equal to 10.
4. The generalized fiducial point estimator of R is $\hat{R}_{GFI} = \frac{1}{N} \sum_{l=M+1}^{M+N} \hat{R}_{GFI}^{(l)}$. Select the $N\gamma/2$ th and $N(1 - \gamma/2)$ th of the permutation as $\hat{R}_{GFI, \gamma/2}$ and $\hat{R}_{GFI, 1-\gamma/2}$, respectively. Then, the $100(1 - \gamma)\%$ fiducial confidence interval of R is $[\hat{R}_{GFI, \gamma/2}, \hat{R}_{GFI, 1-\gamma/2}]$.

3. Estimation of R with Different Scale and Shape Parameters

Suppose $X \sim GL(\sigma_1, \alpha_1)$ and $Y \sim GL(\sigma_2, \alpha_2)$ are independently distributed under different scale parameters, σ_1 and σ_2 , then $R = P(Y < X)$, it can be easily seen that

$$\begin{aligned} R = P(Y < X) &= \int_{-\infty}^{\infty} \int_{-\infty}^x f_Y(y) f_X(x) dy dx \\ &= \int_0^1 \left[1 + \left(t^{-\frac{1}{\alpha_1}} - 1 \right)^{\frac{\sigma_1}{\sigma_2}} \right]^{-\alpha_2} dt. \end{aligned} \quad (34)$$

3.1. Maximum Likelihood Estimation of R

Let $\mathbf{x} = (x_1, \dots, x_n)^T$ be a random sample from $GL(\sigma_1, \alpha_1)$ and let $\mathbf{y} = (y_1, \dots, y_m)^T$ be another independent random sample from $GL(\sigma_2, \alpha_2)$. The log-likelihood function is

$$\begin{aligned} L(\sigma_1, \sigma_2, \alpha_1, \alpha_2 | \mathbf{x}, \mathbf{y}) &= n \log \alpha_1 + m \log \alpha_2 - n \log \sigma_1 - m \log \sigma_2 - \frac{1}{\sigma_1} \sum_{i=1}^n x_i - \frac{1}{\sigma_2} \sum_{j=1}^m y_j \\ &\quad - (\alpha_1 + 1) \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma_1}}) - (\alpha_2 + 1) \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma_2}}). \end{aligned} \quad (35)$$

Similarly, the ML estimators of α_1 as a function of σ_1 and α_2 as a function of σ_2 are

$$\hat{\alpha}_1(\sigma_1) = \frac{n}{\sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma_1}})} \quad \text{and} \quad \hat{\alpha}_2(\sigma_2) = \frac{m}{\sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma_2}})}. \quad (36)$$

The ML estimators of σ_1 and σ_2 can be solved from $h_1(\sigma_1) = \sigma_1$ and $h_2(\sigma_2) = \sigma_2$. Therefore, the ML estimator of R is

$$\hat{R}_{ML} = \int_0^1 \left[1 + \left(t^{-\frac{1}{\hat{\alpha}_1}} - 1 \right)^{\frac{\hat{\sigma}_1}{\hat{\sigma}_2}} \right]^{-\hat{\alpha}_2} dt, \quad (37)$$

and the asymptotic $100(1 - \gamma)\%$ confidence intervals of R can also be obtained.

3.2. Generalized Inference of R

Let

$$T_1(\sigma_1) = 2 \sum_{i=1}^{n-1} \log(S_n/S_i) \quad \text{and} \quad T_2(\sigma_2) = 2 \sum_{j=1}^{m-1} \log(S_m/S_j). \quad (38)$$

Then we have $T_1(\sigma_1) \sim \chi^2(2n - 2)$ and $T_2(\sigma_2) \sim \chi^2(2m - 2)$ from Lemma 1. In addition, we can prove that $T_1(\sigma_1)$ and $T_2(\sigma_2)$ are strictly increasing on $(0, +\infty)$, and

$$\lim_{\sigma_i \rightarrow 0^+} T_i(\sigma_i) = 0 \quad \text{and} \quad \lim_{\sigma_i \rightarrow +\infty} T_i(\sigma_i) = +\infty, \quad i = 1, 2. \quad (39)$$

Furthermore, when $T_1(\sigma_1) \sim \chi^2(2n - 2)$ and $T_2(\sigma_2) \sim \chi^2(2m - 2)$ are given, both $T_1(\sigma_1) = T_1$ and $T_2(\sigma_2) = T_2$ have unique solutions denoted by $\sigma_1 = g_1(T_1, \mathbf{X})$ and $\sigma_2 = g_2(T_2, \mathbf{Y})$.

Since $U_1 = 2\alpha_1 S_n \sim \chi^2(2n)$ and $U_2 = 2\alpha_2 S_m \sim \chi^2(2m)$, we find that $\alpha_1 = U_1/(2S_n)$ and $\alpha_2 = U_2/(2S_m)$. The generalized pivotal quantity of R is

$$R_{GI} = \int_0^1 \left[1 + \left(t^{-\frac{2s_n}{U_1}} - 1 \right)^{\frac{g_1(T_1, \mathbf{X})}{g_2(T_2, \mathbf{Y})}} \right]^{-\frac{U_2}{2s_m}} dt, \quad (40)$$

where $s_n = \sum_{i=1}^n \log(1 + e^{-x_{(i)}/g_1(T_1, \mathbf{X})})$ and $s_m = \sum_{j=1}^m \log(1 + e^{-y_{(j)}/g_2(T_2, \mathbf{Y})})$. The steps to calculate the generalized point and interval estimations of R are similar to those in Section 2.2.

3.3. Generalized Fiducial Inference of R

For the observed value $\mathbf{x} = (x_1, \dots, x_n)^T$, we have

$$U_i = F(x_i; \boldsymbol{\eta}), \quad i = 1, \dots, n, \quad (41)$$

where $\boldsymbol{\eta} = (\sigma_1, \alpha_1)^T$ and $F(x_i; \sigma_1, \alpha_1) \triangleq (1 + e^{-\frac{x_i}{\sigma_1}})^{-\alpha_1}$ is the CDF of GL distribution. Based on (41), we can obtain the i -th $x_i = G_i(u_i, \boldsymbol{\theta})$ is $x_i = -\sigma_1 \log(U_i^{-\frac{1}{\alpha_1}} - 1)$. Then, we have

$$\left. \frac{\partial G_i}{\partial \sigma_1} \right|_{u_i=(1+e^{-\frac{x_i}{\sigma_1}})^{-\alpha_1}} = \frac{x_i}{\sigma_1} \quad \text{and} \quad \left. \frac{\partial G_i}{\partial \alpha_1} \right|_{u_i=(1+e^{-\frac{x_i}{\sigma_1}})^{-\alpha_1}} = \frac{\sigma_1}{\alpha_1} (1 + e^{-\frac{x_i}{\sigma_1}}) \log(1 + e^{-\frac{x_i}{\sigma_1}}). \quad (42)$$

It can be calculated that

$$J(\mathbf{x}, \sigma_1, \alpha_1) = \frac{1}{\alpha_1} \sum_{i \neq j} \left| x_i (1 + e^{-\frac{x_j}{\sigma_1}}) \log(1 + e^{-\frac{x_j}{\sigma_1}}) - x_j (1 + e^{-\frac{x_i}{\sigma_1}}) \log(1 + e^{-\frac{x_i}{\sigma_1}}) \right|. \quad (43)$$

Finally, we obtain the following GFD for (σ_1, α_1) , i.e.,

$$f_F(\sigma_1, \alpha_1 | \mathbf{x}) = \frac{f(\mathbf{x} | \sigma_1, \alpha_1) J(\mathbf{x}, \sigma_1, \alpha_1)}{\int_0^\infty \int_0^\infty f(\mathbf{x} | \sigma_1, \alpha_1) J(\mathbf{x}, \sigma_1, \alpha_1) d\sigma_1 d\alpha_1}, \quad (44)$$

where $f(\mathbf{x} | \sigma_1, \alpha_1) = \prod_{i=1}^n f_i(x_i; \sigma_1, \alpha_1)$. Specifically,

$$f_F(\sigma_1, \alpha_1 | \mathbf{x}) \propto \frac{\alpha_1^n e^{-\sum_{i=1}^n \frac{x_i}{\sigma_1}}}{\sigma_1^n \prod_{i=1}^n (1 + e^{-\frac{x_i}{\sigma_1}})^{\alpha_1 + 1}} \cdot \frac{1}{\alpha_1} \sum_{i \neq j} |q(x_i, x_j, \sigma_1)|, \quad (45)$$

where

$$q(x_i, x_j, \sigma_1) = x_i (1 + e^{-\frac{x_j}{\sigma_1}}) \log(1 + e^{-\frac{x_j}{\sigma_1}}) - x_j (1 + e^{-\frac{x_i}{\sigma_1}}) \log(1 + e^{-\frac{x_i}{\sigma_1}}).$$

From Formula (45), the conditional fiducial density function of σ_1 given α_1 is given by

$$f_F(\sigma_1 | \alpha_1, \mathbf{x}) \propto \sigma_1^{-n} e^{-\sum_{i=1}^n \frac{x_i}{\sigma_1} - (\alpha_1 + 1) \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma_1}})} \sum_{i=1}^n |q(x_i, x_j, \sigma_1)|. \quad (46)$$

Then, the conditional fiducial density function of α_1 given σ_1 is

$$f_F(\alpha_1 | \sigma_1, \mathbf{x}) \propto \alpha_1^{n-1} e^{-\alpha_1 \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma_1}})}, \quad (47)$$

which implies that the conditional density of α_1 is $Ga\left(n, \sum_{i=1}^n \log(1 + e^{-\frac{x_i}{\sigma_1}})\right)$. Based on the same method, the conditional fiducial density function of σ_2 and α_2 can be obtained as follows

$$f_F(\sigma_2 | \alpha_2, \mathbf{y}) \propto \sigma_2^{-m} e^{-\sum_{j=1}^m \frac{y_j}{\sigma_2} - (\alpha_2 + 1) \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma_2}})} \sum_{j=1}^m |q(y_i, y_j, \sigma_2)|, \quad (48)$$

$$f_F(\alpha_2 | \sigma_2, \mathbf{y}) \propto \alpha_2^{m-1} e^{-\alpha_2 \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma_2}})}, \quad (49)$$

which means the conditional density of α_2 is $Ga\left(m, \sum_{j=1}^m \log(1 + e^{-\frac{y_j}{\sigma_2}})\right)$.

We still introduce standard Metropolis–Hastings steps into the Gibbs sampler to update σ_1 and σ_2 while updating α_1 and α_2 from their exact conditional distributions. The detailed steps of the algorithm are similar to Section 2.3.

4. Simulations

Let $\hat{R}_{ML}$ represent the ML estimators, $\hat{R}_{GI}$ denote the point estimators via the GI method, and $\hat{R}_{GFI}$ denotes the point estimators by the GFI method. ACI refers to the asymptotic confidence interval, GCI denotes the generalized confidence interval, and FCI is the fiducial confidence interval. To compare the above point estimators, 1000 simulations are conducted by using the mean square error (MSE) and relative mean square error (RMSE). The RMSE is calculated as the MSE obtained from ML and GI methods divided by the MSE of the GFI method. For example, the RMSE of $\hat{R}_{ML}$ is given by the MSE of $\hat{R}_{ML}$ divided by the MSE of $\hat{R}_{GFI}$, where the GFI method is always the benchmark method. Meanwhile, we calculate the performance of the above confidence intervals with average length and empirical coverage. The relative length is the ratio of the average length gained by the ML and GI methods to the average length obtained by the GFI method. Different combinations of $(n, m, \sigma, \alpha_1, \alpha_2)$ and $(n, m, \sigma_1, \sigma_2, \alpha_1, \alpha_2)$ are provided at a nominal level $1 - \gamma = 0.95$. We have the following conclusions.

4.1. Analysis of Point Estimates

- The case with the same scale parameter.

Table 1 provides the MSEs of R for different parameter combinations, and Figure 1 shows the boxplots of RMSEs of R . The detailed information is shown as follows.

Table 1. The MSE of the point estimations for R with the same scale parameter.

$(\sigma, \alpha_1, \alpha_2)$	n	m	MSE for R		
			$\hat{R}_{ML}$	$\hat{R}_{GI}$	$\hat{R}_{GFI}$
(1.0, 1.5, 2.0)	15	15	0.008158	0.007170	0.007172
	15	25	0.006360	0.005731	0.005707
	15	50	0.005174	0.004761	0.004743
	25	15	0.006392	0.005827	0.005788
	25	25	0.005120	0.004749	0.004730
	25	50	0.003560	0.003364	0.003337
	50	15	0.005224	0.004895	0.004892
	50	25	0.003372	0.003231	0.003230
	50	50	0.002353	0.002269	0.002250
(1.0, 2.0, 1.5)	15	15	0.009041	0.007948	0.007925
	15	25	0.007053	0.006388	0.006369
	15	50	0.005025	0.004639	0.004625
	25	15	0.006205	0.005595	0.005577
	25	25	0.005155	0.004742	0.004749
	25	50	0.003872	0.003617	0.003631
	50	15	0.005104	0.004693	0.004661
	50	25	0.003676	0.003468	0.003457
	50	50	0.002316	0.002232	0.002221

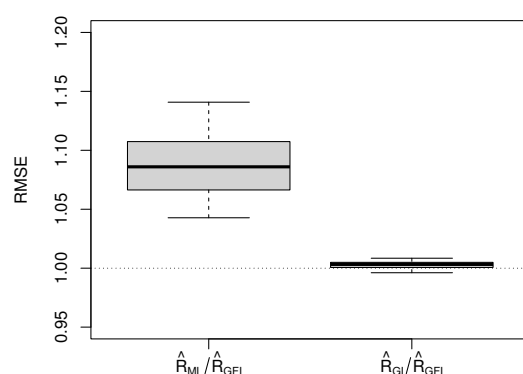


Figure 1. Summary of RMSE of the point estimations for R with the same scale parameter.

From Table 1 and Figure 1, the RMSEs of $\hat{R}_{ML}$ are larger than 1 while the RMSEs of $\hat{R}_{GI}$ are close to 1. Specifically, the MSEs of $\hat{R}_{GFI}$ are often smaller than those of $\hat{R}_{ML}$, and the gap is significant in small and moderate samples, such as $n, m \leq 25$. Meanwhile, the difference between the MSEs of $\hat{R}_{GFI}$ and $\hat{R}_{GI}$ is trivial.

- The case with different scale parameters.

The MSEs of R and the boxplots of RMSEs of R under different parameter combinations are shown in Table 2 and Figure 2.

Table 2. The MSEs of the point estimations for R with different scale parameters.

$(\sigma_1, \alpha_1, \sigma_2, \alpha_2)$	n	m	MSE for R		
			$\hat{R}_{ML}$	$\hat{R}_{GI}$	$\hat{R}_{GFI}$
(1.0, 1.5, 2.0, 2.0)	15	15	0.009901	0.018722	0.008245
	15	25	0.006940	0.014136	0.006285
	15	50	0.004604	0.009251	0.004256
	25	15	0.009201	0.018626	0.007926
	25	25	0.005662	0.013187	0.005149
	25	50	0.003431	0.007822	0.003238
	50	15	0.007747	0.016616	0.006800
	50	25	0.004920	0.011690	0.004507
	50	50	0.002845	0.007615	0.002739
(2.0, 2.0, 1.0, 1.5)	15	15	0.010145	0.007937	0.008674
	15	25	0.008758	0.007502	0.007744
	15	50	0.008042	0.007036	0.007151
	25	15	0.007037	0.005460	0.006330
	25	25	0.005563	0.004669	0.005115
	25	50	0.005155	0.004710	0.004744
	50	15	0.004525	0.003373	0.004208
	50	25	0.003411	0.002811	0.003236
	50	50	0.003105	0.002765	0.002963

In Figure 2, it is shown that the RMSEs of $\hat{R}_{ML}$ and $\hat{R}_{GI}$ are often larger than 1. From Table 2, the MSEs of $\hat{R}_{GFI}$ are smaller than those of $\hat{R}_{ML}$, and the MSEs of $\hat{R}_{GI}$ exhibit lower stability. At the same time, the MSEs of the three methods decrease with the increase in the sample size.

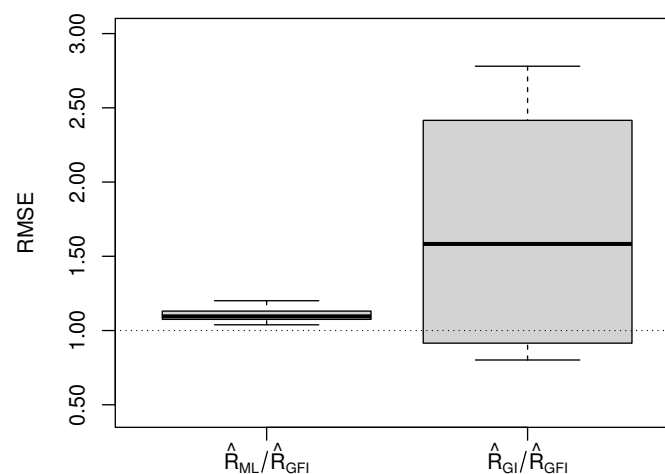


Figure 2. Summary of RMSEs of the point estimations for R with different scale parameters.

4.2. Analysis of Interval Estimates

- The case with the same scale parameter.

Table 3 provides the average length and empirical coverage of R , and Figure 3 shows the boxplots of relative length and empirical coverage. The details are as follows.

Table 3. The average length and empirical coverage of 95% two-sided confidence intervals for R with the same scale parameter.

$(\sigma, \alpha_1, \alpha_2)$	n	m	Average Length			Empirical Coverage		
			ACI	GCI	FCI	ACI	GCI	FCI
(1.0, 1.5, 2.0)	15	15	0.350	0.338	0.338	0.911	0.957	0.958
	15	25	0.317	0.305	0.305	0.930	0.957	0.960
	15	50	0.283	0.275	0.275	0.950	0.960	0.962
	25	15	0.312	0.305	0.305	0.899	0.954	0.955
	25	25	0.277	0.265	0.266	0.911	0.954	0.953
	25	50	0.240	0.231	0.231	0.949	0.957	0.959
	50	15	0.287	0.277	0.277	0.885	0.952	0.951
	50	25	0.247	0.232	0.233	0.919	0.955	0.953
	50	50	0.197	0.190	0.190	0.949	0.957	0.961
	50	50	0.197	0.190	0.190	0.949	0.957	0.961
(1.0, 2.0, 1.5)	15	15	0.348	0.336	0.337	0.912	0.949	0.953
	15	25	0.310	0.305	0.305	0.942	0.961	0.958
	15	50	0.278	0.277	0.277	0.935	0.958	0.958
	25	15	0.313	0.304	0.304	0.944	0.964	0.962
	25	25	0.273	0.265	0.265	0.939	0.953	0.955
	25	50	0.233	0.231	0.232	0.936	0.956	0.957
	50	15	0.284	0.275	0.276	0.947	0.954	0.955
	50	25	0.236	0.231	0.231	0.947	0.956	0.958
	50	50	0.192	0.190	0.190	0.948	0.958	0.955
	50	50	0.192	0.190	0.190	0.948	0.958	0.955

Table 3 and Figure 3 show that the relative lengths of ACIs are greater than 1 and the ACIs are too liberal. The difference between GCIs and FCIs is small and both of them are conservative. When the sample size is small, GCIs and FCIs are better than ACIs. Meanwhile, the average lengths of the three methods tend to decrease with the increase in sample size.

- The case with different scale parameters.

The average length and empirical coverage of R , the boxplots of relative length, and empirical coverage are shown in Table 4 and Figure 4.

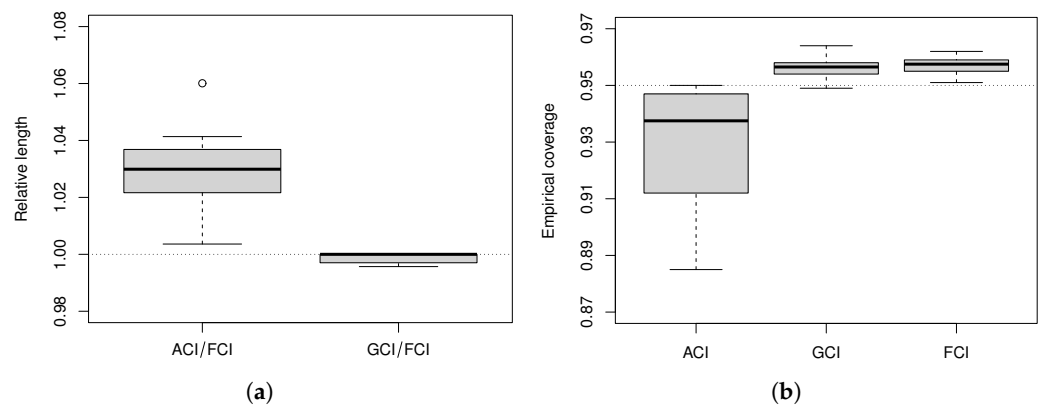
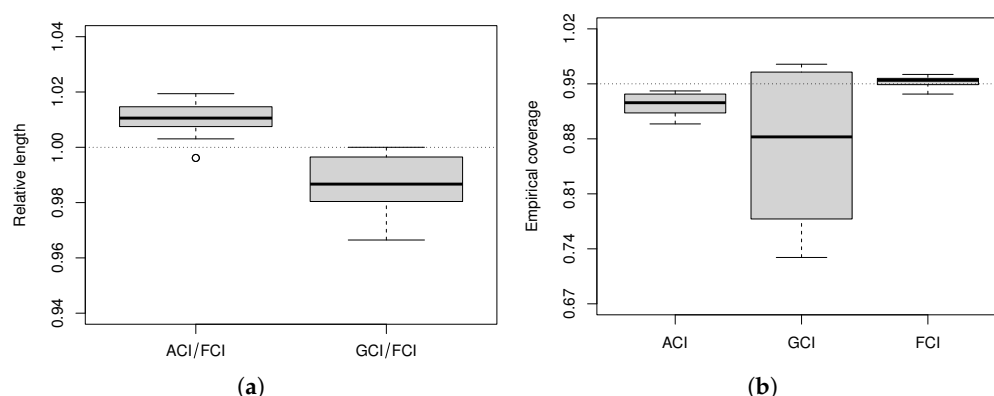


Figure 3. Summary of the relative length and empirical coverage for R with the same scale parameter. (a) Display of the relative length; (b) display of the empirical coverage.

Table 4. The average length and empirical coverage of 95% two-sided confidence intervals for R with different scale parameters.

$(\sigma_1, \alpha_1, \sigma_2, \alpha_2)$	n	m	Average Length			Empirical Coverage		
			ACI	GCI	FCI	ACI	GCI	FCI
(1.0, 1.5, 2.0, 2.0)	15	15	0.365	0.352	0.359	0.899	0.818	0.957
	15	25	0.313	0.304	0.309	0.928	0.804	0.959
	15	50	0.258	0.254	0.259	0.923	0.821	0.953
	25	15	0.346	0.332	0.341	0.900	0.778	0.947
	25	25	0.289	0.278	0.284	0.929	0.765	0.962
	25	50	0.230	0.222	0.228	0.941	0.793	0.961
	50	15	0.329	0.317	0.328	0.913	0.765	0.955
	50	25	0.268	0.261	0.265	0.932	0.761	0.957
	50	50	0.206	0.200	0.204	0.940	0.729	0.954
(2.0, 2.0, 1.0, 1.5)	15	15	0.368	0.360	0.361	0.901	0.969	0.957
	15	25	0.349	0.343	0.343	0.918	0.959	0.949
	15	50	0.330	0.328	0.328	0.908	0.957	0.952
	25	15	0.312	0.306	0.308	0.937	0.966	0.956
	25	25	0.289	0.284	0.285	0.937	0.965	0.956
	25	50	0.269	0.267	0.267	0.924	0.944	0.940
	50	15	0.259	0.257	0.260	0.937	0.975	0.953
	50	25	0.231	0.227	0.229	0.939	0.971	0.949
	50	50	0.206	0.203	0.204	0.922	0.946	0.937

Figure 4 demonstrates that the relative lengths of ACIs are greater than 1 while those of GCIs are smaller than 1. FCIs are close to the nominal level while ACIs and GCIs are obviously liberal. To be specific, Table 4 shows that the average lengths of ACIs are long while the empirical coverages of ACIs are often less than 0.95. The average lengths of GCIs are short, but the empirical coverages of GCIs exhibit instability. The average lengths of FCIs and GCIs are comparable when the sample size is large, and FCIs can reach the nominal level.

**Figure 4.** Summary of relative length and empirical coverage for R with different scale parameters. (a) Display of the relative length; (b) display of the empirical coverage.

5. Real Data Example

5.1. The Breaking Strengths of Jute Fibers

The first dataset was originally introduced by Xia et al. [36]. It consists of the breaking strengths of jute fibers at 4 different gauge lengths: 5 mm, 10 mm, 15 mm, and 20 mm. The breaking strengths of jute fibers at 10 mm and 20 mm are presented in Table 5.

Table 5. The breaking strengths of jute fibers at different gauge lengths.

Gauge Lengths	Data	Sample Size
10 mm	693.73, 704.66, 323.83, 778.17, 123.06, 637.66, 383.43, 151.48, 108.94, 50.16, 671.49, 183.16, 257.44, 727.23, 291.27, 101.15, 376.42, 163.40, 141.38, 700.74, 262.90, 353.24, 422.11, 43.93, 590.48, 212.13, 303.90, 506.60, 530.55, 177.25	30
20 mm	71.46, 419.02, 284.64, 585.57, 456.60, 113.85, 187.85, 688.16, 662.66, 45.58, 578.62, 756.70, 594.29, 166.49, 99.72, 707.36, 765.14, 187.13, 145.96, 350.70, 547.44, 116.99, 375.81, 581.60, 119.86, 48.01, 200.16, 36.75, 244.53, 83.55	30

The breaking strengths of jute fibers at two different gauge lengths are fitted with GL distribution, respectively. The estimated scale parameters, shape parameters, Kolmogorov–Smirnov (K-S) distances, and the corresponding p -values are shown in Table 6.

Table 6. The scale parameter, shape parameter, K-S, and p -values of the breaking strengths of jute fibers.

Gauge Lengths	Scale Parameter	Shape Parameter	K-S	p -Value
10 mm	170.826	5.190	0.118	0.756
20 mm	178.278	4.210	0.161	0.376

Referring to Table 6, the p -values obtained from the K-S test indicate that the GL distribution shows good agreement with the jute fiber data. Since the difference is significant between the scale parameters estimated at the two gauge lengths, it is reasonable to assume the scale parameters are different. The point and interval estimations of R are shown in Table 7, which implies that the ACI is the shortest while the GCI is the longest. Because the empirical coverage of ACI tends to be liberal in the simulation, we prefer to recommend the GFI method.

Table 7. The result of R for the breaking strengths of jute fibers.

	10 mm and 20 mm		
	Point	Interval	Length
ACI	0.535	[0.390, 0.680]	0.290
GCI	0.526	[0.338, 0.702]	0.364
FCI	0.531	[0.374, 0.691]	0.317

5.2. The Sulfur Dioxide Concentration Data

To illustrate the methods developed in Sections 2 and 3, the second real dataset provided by Roberts [37] is considered. It consists of data on the monthly and annual maxima of one-hour mean concentrations of sulfur dioxide (pphm) for Long Beach, California from 1956 to 1974. In this paper, the average hourly concentrations of sulfur dioxide in January, March, and August are shown in Table 8.

Table 8. The sulfur dioxide concentration data under different months.

Months	Data	Sample Size
January	47, 22, 15, 20, 22, 25, 20, 12, 16, 16, 27, 30, 51, 37, 23, 22, 30, 10, 8	19
March	44, 20, 20, 20, 23, 20, 15, 27, 3, 9, 25, 32, 18, 55, 10, 20, 18, 8, 9	19
August	21, 16, 20, 15, 9, 10, 10, 4, 25, 18, 18, 26, 25, 17, 40, 55, 19, 16, 9	19

The sulfur dioxide concentration data of three months are fitted with GL distributions separately. We present the estimated scale parameters, shape parameters, K-S distances, and corresponding p -values in Table 9.

Table 9. The scale parameter, shape parameter, K-S, and p -values of sulfur dioxide concentration data.

Months	Scale Parameter	Shape Parameter	K-S	p -Value
January	8.070	11.009	0.096	0.995
March	8.417	6.991	0.147	0.807
August	7.376	8.154	0.109	0.978

From Table 9, the p -values of the K-S test are pretty good (p -values of 0.995, 0.804, and 0.978, respectively), which means that the GL distribution fits well with the sulfur dioxide concentration data. In addition, the p -values of the K-S test for the GL distribution are larger than those of the Weibull and generalized exponential distributions, which means the GL distribution provides a better fit than other distributions. Hence, the GL distribution is adopted in this real dataset and we consider the following two cases.

- The case with the same scale parameter.

In this case, the average hourly concentrations of sulfur dioxide in January and March are chosen. Since the two estimated scale parameters are not very different, it is natural to assume that the two scale parameters are equal. The ML estimations for the parameters σ , α_1 , and α_2 are given by $\hat{\sigma} = 8.247$, $\hat{\alpha}_1 = 10.593$, and $\hat{\alpha}_2 = 7.179$, respectively. Using the three methods in Section 2, the point and interval estimations for R are shown in Table 10.

Table 10. The results of R of sulfur dioxide concentration data.

	January and March			January and August		
	Point	Interval	Length	Point	Interval	Length
ACI	0.596	[0.444, 0.748]	0.304	0.626	[0.453, 0.799]	0.346
GCI	0.586	[0.412, 0.731]	0.319	0.621	[0.327, 0.840]	0.513
FCI	0.590	[0.431, 0.726]	0.295	0.615	[0.401, 0.808]	0.407

From Table 10, it can be concluded that the difference between the point estimates of R is small while the FCI is shorter than the ACI and GCI.

- The case with different scale parameters.

In this case, the average hourly concentration of sulfur dioxide in January and August are selected. Table 9 shows that the differences between the two estimated scale and shape parameters are significant, so it is reasonable to assume the parameters are different. The ML estimations for the parameters σ_1 , α_1 , σ_2 , and α_2 are given by $\hat{\sigma}_1 = 8.070$, $\hat{\alpha}_1 = 11.009$, $\hat{\sigma}_2 = 7.376$, and $\hat{\alpha}_2 = 8.154$, respectively. The point and interval estimations for R are shown in Table 10 by using the three methods in Section 3. It is found that the difference between the point estimates of R is small, and the ACI is the shortest while the GCI is the longest.

In general, the FCI performs better in the first case while the ACI performs better in the second case. However, the FCI is recommended, considering that the empirical coverage of ACI in the simulation is often lower than the nominal level.

5.3. The Insulating Fluid Data

The Ln times to breakdown of insulating fluid in an accelerated test reported by Nelson [38] is chosen as the third real data example. The Ln times to breakdown for insulating fluid were reported at different voltages of 26, 28, 30, 32, 34, 36, and 38 kV. The Ln times to breakdown 32, 34, and 36 kV are demonstrated in Table 11.

Table 11. The insulating fluid data at different voltages.

Voltages	Data	Sample Size
32 kV	−1.3094, −0.9163, −0.3711, −0.2358, 1.0116, 1.3635, 2.2905, 2.6354, 2.7682, 3.3250, 3.9748, 4.4170, 4.4918, 4.6109, 5.3711	15
34 kV	−1.6608, −0.2485, −0.0409, 0.2700, 1.0224, 1.1505, 1.4231, 1.5411, 1.5789, 1.8718, 1.9947, 2.0806, 2.1126, 2.4898, 3.4578, 3.4818, 3.5237, 3.6030, 4.2889	19
36 kV	−1.0499, −0.5277, −0.0409, −0.0101, 0.5247, 0.6780, 0.7275, 0.9477, 0.9969, 1.0647, 1.3001, 1.3837, 1.6770, 2.6224, 3.2386	15

The data of Ln times to breakdown at three different voltages are fitted with GL distributions separately. Table 12 demonstrates the estimated scale parameters, shape parameters, K-S distances, and the corresponding p -values.

Table 12. The scale parameter, shape parameter, K-S, and p -values of insulating fluid data.

Voltages	Scale Parameter	Shape Parameter	K-S	p -Value
32 kV	1.690	2.680	0.141	0.888
34 kV	1.211	3.067	0.124	0.900
36 kV	0.755	2.354	0.119	0.967

Based on Table 12, the GL distribution fits quite well with the data of Ln times to breakdown 32, 34, and 36 kV. Therefore, it is reasonable for us to apply the GL distribution to this real dataset, and we still consider the following two cases.

- The case with the same scale parameter.

The data of Ln times to breakdown at 32 and 34 kV are selected in this case. If we suppose the two scale parameters are equal, the ML estimations for the parameters σ , α_1 and α_2 are given by $\hat{\sigma} = 1.424$, $\hat{\alpha}_1 = 2.769$, and $\hat{\alpha}_2 = 2.882$, respectively. The point and interval estimations for R are shown in Table 13, illustrating that the ACI is the shortest while the GCI is the longest.

Table 13. The result of R of insulating fluid data.

	32 kV and 34 kV			34 kV and 36 kV		
	Point	Interval	Length	Point	Interval	Length
ACI	0.490	[0.329, 0.651]	0.322	0.674	[0.499, 0.848]	0.349
GCI	0.489	[0.323, 0.661]	0.338	0.666	[0.447, 0.829]	0.382
FCI	0.487	[0.319, 0.650]	0.331	0.652	[0.466, 0.817]	0.351

- The case with different scale parameters.

In the second case, the data of Ln times to breakdown at 34 and 36 kV are considered. When we assume all the parameters are different, the ML estimations for the parameters σ_1 , α_1 , σ_2 , and α_2 can be given by $\hat{\sigma}_1 = 1.2106$, $\hat{\alpha}_1 = 3.0672$, $\hat{\sigma}_2 = 0.7552$, and $\hat{\alpha}_2 = 2.3542$, respectively. Table 13 states the point and interval estimations for R , and it shows that the ACI is the shortest and the GCI is the longest.

According to the above two cases, it seems that the effect of ACI is better. Since the simulation results show that the empirical coverage of ACI is often lower than the nominal level, we still prefer to select FCI with higher reliability.

6. Discussion

The estimation of R in the GL distribution is an important research problem. Most of the existing literature studies focus on the ML estimation and Bayesian inference. However, the ML estimation cannot obtain the exact pivotal quantity and its empirical coverage

sometimes fails to reach the nominal level. In Bayesian inference, the choice of the prior distribution is improper or subjective. Therefore, we introduce two novel methods to estimate R in the GL distribution.

On the one hand, there are two theoretical implications worth noting. First, the GFI method is applied to estimate R . The prior of the GFI is based on actual data, which makes the posterior distribution more objective. In addition, the weighted prior is applied when the scale parameters are the same. Our findings suggest that this approach of constructing the prior is suitable for estimating R in the two-parameter GL distribution and can be extended to other distributions as well. Second, the GI method offers another way when the conventional pivotal quantity is not available. By developing two lemmas, the generalized point estimation and generalized confident interval of R can be given.

On the other hand, this article has three practical implications. First, the simulation results indicate that the generalized fiducial method is better for the point estimation of R with the comparisons of the MSE. Moreover, it can be concluded that the GFI method often outperforms the ML and GI methods for the interval estimation of R , which presents more advantages in average length and empirical coverage. Second, the results of the three real data example state that the estimation of R can be applied in many different fields. Third, the two-parameter GL distribution without a location parameter is particularly useful in estimating R , where the dataset contains values less than zero. This characteristic expands its applicability to a wider range of datasets and deserves more attention in the scale-shape life distribution.

There are some limitations in our study. Due to encountering censored data in numerous survival analyses, such as the research of Rao [39], Babayi and Khorram [40], and Wang et al. [41], the statistical inference of parameters, reliability, and stress–strength based on censored samples under the GL distribution would be an interesting direction for future works.

7. Conclusions

This article considers the statistical inference of R for the generalized logistic distribution with either the same or different scale parameters. Based on the simulation of the point estimation, the MSE of the GI and GFI methods is often smaller for the same scale parameter. However, the GFI has the smallest MSE when the scale parameters are different. According to the simulation of the interval estimation, both the GI and GFI methods exhibit shorter average lengths and more conservative empirical coverage for the same scale parameter. When the scale parameters are different, the GFI method performs better in the length and coverage. Therefore, we believe that the GFI method is more suitable for estimating R in the GL distribution and many other issues related to it.

Author Contributions: Conceptualization, L.Y. and X.C.; methodology, M.L., L.Y., Y.Q. and X.C.; software, M.L. and L.Y.; validation, M.L. and L.Y.; formal analysis, M.L.; investigation, Y.Q. and X.C.; resources, L.Y., Y.Q. and X.C.; data curation, M.L.; writing—original draft preparation, M.L.; writing—review and editing, K.K.S.; visualization, M.L.; supervision, L.Y. and K.K.S.; project administration, L.Y. and X.C.; funding acquisition, L.Y. and X.C. All authors have read and agreed to the published version of the manuscript.

Funding: This work was supported by the Natural Science Foundation of Hebei Province under grant No. A2020207006 and No. A2022208001, the Foundation of Hebei Educational Department under grant No. CXZZSS2023103, and the National Natural Science Foundation of China under grant No. 12001155.

Data Availability Statement: No data were used to support this study.

Acknowledgments: The work is based on research supported by the Natural Science Foundation of Hebei Province, the Foundation of Hebei Educational Department, and the National Natural Science Foundation of China. Any opinion, finding, conclusion, or recommendation expressed in this material is that of the authors, and the NRF does not accept any liability in this regard.

Conflicts of Interest: The authors declare no conflict of interest.

References

- Birnbaum, Z.W. On a use of the mann-whitney statistics. In Proceedings of the Third Berkeley Symposium on Mathematical Statistics and Probability, Berkeley, CA, USA, 26–31 December 1954; Volume 1, pp. 13–17.
- Birnbaum, Z.W.; McCarty, R.C. A distribution-free upper confidence bound for $Pr(Y < X)$, based on independent samples of X and Y . *Ann. Math. Stat.* **1958**, *29*, 558–562.
- Hall, I.J. Approximate one-sided tolerance limits for the difference or sum of two independent normal variates. *J. Qual. Technol.* **1984**, *16*, 15–19. [\[CrossRef\]](#)
- Guttman, I.; Johnson, R.A.; Bhattacharyya, G.K.; Reiser, B. Confidence limits for stress-strength models with explanatory variables. *Technometrics* **1988**, *30*, 161–168. [\[CrossRef\]](#)
- Krishnamoorthy, K.; Lin, Y. Confidence limits for stress–strength reliability involving Weibull models. *J. Stat. Plan. Infer.* **2010**, *140*, 1754–1764. [\[CrossRef\]](#)
- Nooghabi, M.J.; Naderi, M. Stress–strength reliability inference for the Pareto distribution with outliers. *J. Comput. Appl. Math.* **2022**, *404*, 113911. [\[CrossRef\]](#)
- Surles, J.G.; Padgett, W.J. Inference for $P(Y < X)$ in the Burr Type X model. *J. Appl. Stat. Sci.* **1998**, *7*, 225–238.
- Jafari, A.A.; Bafekri, S. Inference on stress-strength reliability for the two-parameter exponential distribution based on generalized order statistics. *Math. Popul. Stud.* **2021**, *28*, 201–227. [\[CrossRef\]](#)
- Yousef, M.M.; Almetwally, E.M. Multi stress-strength reliability based on progressive first failure for Kumaraswamy model: Bayesian and non-Bayesian estimation. *Symmetry* **2021**, *13*, 2120. [\[CrossRef\]](#)
- Yousef, M.M.; Fayomi, A.; Almetwally, E.M. Simulation techniques for strength component partially accelerated to analyze stress–strength model. *Symmetry* **2023**, *15*, 1183. [\[CrossRef\]](#)
- Gunasekera, S. Generalized inferences of $R = P(X > Y)$ for Pareto distribution. *Stat. Pap.* **2014**, *56*, 333–351. [\[CrossRef\]](#)
- Rezaei, S.; Tahmasbi, R.; Mahmoodi, M. Estimation of $P(Y < X)$ for generalized Pareto distribution. *J. Stat. Plan. Infer.* **2010**, *140*, 480–494.
- Krishnamoorthy, K.; Mukherjee, S.; Guo, H. Inference on reliability in two-parameter exponential stress–strength model. *Metrika* **2007**, *65*, 261–273. [\[CrossRef\]](#)
- Kundu, D.; Gupta, R.D. Estimation of $P(Y < X)$ for generalized exponential distribution. *Metrika* **2005**, *61*, 291–308.
- Mahmoud, M.A.W.; El-Sagheer, R.M.; Soliman, A.A.; Ellah, A.H.A. Bayesian estimation of $P(Y < X)$ based on record values from the Lomax distribution and MCMC technique. *J. Mod. Appl. Stat. Meth.* **2016**, *15*, 488–510.
- Cruz, R.D.L.; Salinas, H.S.; Meza, C. Reliability estimation for stress-strength model based on Unit-half-normal distribution. *Symmetry* **2022**, *14*, 837. [\[CrossRef\]](#)
- Alsadat, N.; Hassan, A.S.; Elgarhy, M.; Chesneau, C.; Mohamed, R.E. An efficient stress–Strength reliability estimate of the Unit gompertz distribution using ranked set sampling. *Symmetry* **2023**, *15*, 1121. [\[CrossRef\]](#)
- Asgharzadeh, A.; Valiollahi, R.; Raqab, M.Z. Estimation of the stress–strength reliability for the generalized logistic distribution. *Stat. Methodol.* **2013**, *15*, 73–94. [\[CrossRef\]](#)
- Babayi, S.; Khorram, E.; Tondro, F. Inference of $R = P(X < Y)$ for generalized logistic distribution. *J. Theor. Appl. Stat.* **2014**, *48*, 862–871.
- Okasha, H.M. Estimation of $P(Y < X)$ for generalized logistic distribution. *J. Appl. Stat. Sci.* **2014**, *19*, 43–56.
- Rasekhi, M.; Saber, M.M.; Yousof, H.M. Bayesian and classical inference of reliability in multicomponent stress-strength under the generalized logistic model. *Commun. Stat. Theory Meth.* **2020**, *50*, 5114–5125. [\[CrossRef\]](#)
- Balakrishnan, N.; Leung, M.Y. Order statistics from the type I generalized logistic distribution. *Commun. Stat. Simul. Comput.* **1988**, *17*, 25–50. [\[CrossRef\]](#)
- Balakrishnan, N. *Handbook of the Logistic Distribution*, 2nd ed.; Marcel Dekker: New York, NY, USA, 2010; pp. 50–77.
- Lagos-Álvarez, B.; Jerez-Lillo, N.; Navarrete, J.P.; Figueroa-Zúñiga, J.; Leiva, V. A type I generalized logistic distribution: Solving its estimation problems with a bayesian approach and numerical applications based on simulated and engineering data. *Symmetry* **2022**, *14*, 655. [\[CrossRef\]](#)
- Tao, M. Objective Bayesian Analysis for the Generalized Logistic Distribution and Doubly Accelerated Degradation Model. Master’s Thesis, Anhui Normal University, Wuhu, China, 2019.
- Weerahandi, S. Generalized confidence intervals. *J. Am. Stat. Assoc.* **1993**, *88*, 899–905. [\[CrossRef\]](#)
- Wang, B.; Geng, Y.; Zhou, J. Inference for the generalized exponential stress-strength model. *Appl. Math. Modell.* **2017**, *53*, 267–275. [\[CrossRef\]](#)
- Hannig, J.; Iyer, H.; Lai, R.C.S.; Lee, T.C.M. Generalized fiducial inference: A review and new results. *J. Am. Stat. Assoc.* **2016**, *111*, 1346–1361. [\[CrossRef\]](#)
- Yan, L.; Liu, X. Generalized fiducial inference for generalized exponential distribution. *J. Stat. Comput. Sim.* **2018**, *88*, 1369–1381. [\[CrossRef\]](#)
- Yan, L.; Geng, J.; Wang, L.; He, D. Generalized fiducial inference for the Lomax distribution. *J. Stat. Comput. Sim.* **2021**, *91*, 2402–2413. [\[CrossRef\]](#)
- Cai, X.; Feng, S.; Yan, L. Generalized fiducial inference for the lower confidence limit of reliability based on Weibull distribution. *Commun. Stat. Simul. Comput.* **2022**, 1–11. [\[CrossRef\]](#)

32. Wang, X.; Zou, C.; Yi, Y.; Wang, J.; Li, X. Fiducial inference for gamma distributions: Two-sample problems. *Commun. Stat. Simul. Comput.* **2021**, *50*, 811–821. [[CrossRef](#)]
33. Wang, X.; Li, M.; Sun, W.; Gao, Z.; Li, X. Confidence intervals for zero-inflated gamma distribution. *Commun. Stat. Simul. Comput.* **2022**, 1–18. [[CrossRef](#)]
34. Yu, K.; Wang, B.; Patilea, V. New estimating equation approaches with application in lifetime data analysis. *Ann. Inst. Stat. Math.* **2013**, *65*, 589–615. [[CrossRef](#)]
35. Chib, S.; Greenberg, E. Understanding the Metropolis-Hastings algorithm. *Am. Stat.* **1995**, *49*, 327–335.
36. Xia, Z.P.; Yu, J.Y.; Cheng, L.D.; Liu, L.F.; Wang, W.M. Study on the breaking strength of jute fibres using modified Weibull distribution. *Compos. Part A Appl. Sci. Manuf.* **2009**, *40*, 54–59. [[CrossRef](#)]
37. Roberts, E.M. Review of statistics of extreme values with applications to air quality data: Part II. Applications. *J. Air. Pollut. Control Assoc.* **1979**, *29*, 733–740. [[CrossRef](#)]
38. Nelson, W. Analysis of accelerated life test Data-Least squares methods for the inverse power law model. *IEEE. Trans. Reliab.* **1979**, *24*, 103–107. [[CrossRef](#)]
39. Rao, G.S. Estimation of stress-strength reliability from truncated type-I generalised logistic distribution. *Int. J. Math. Oper. Res.* **2015**, *7*, 372–382. [[CrossRef](#)]
40. Babayi, S.; Khorram, E. Inference of stress-strength for the Type-II generalized logistic distribution under progressively Type-II censored samples. *Commun. Stat. Simul. Comput.* **2017**, *47*, 1975–1995. [[CrossRef](#)]
41. Wang, X.; Li, X.; Zhang, L.; Liu, Z.; Li, M. Fiducial inference on gamma distributions: Two-sample problems with multiple detection limits. *Environ. Ecol. Stat.* **2022**, *29*, 453–475. [[CrossRef](#)]

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.