

Gaussian-Bootstrap-Tuned Logit-Shrinkage Estimation of the Log-Logistic Reliability Function under Right Censoring and Precautionary Loss

Rusul Ismail Khalil¹*, Hazim Tallat Hazim²

¹ *Department of Mathematics, College of Basic Education, Diyala University, Diyala, Iraq*

² *Department of Medical Devices Technology, Polytechnic College Mosul, Northern Technical University, Mosul, Iraq*

Abstract

Reliability estimates from small right-censored lifetime samples can be unstable when both parameters of the log-logistic model are estimated. This study shrinks the logit of the mission-time reliability toward an external target and selects the shrinkage coefficient by minimizing a deterministic stratified Gaussian approximation to precautionary risk. The method is compared primarily with the classical plug-in estimator and a positive-part adaptive logit rule, while fixed shrinkage rules are retained as secondary oracle/reference comparators. The Monte Carlo study used 10,000 replications per core condition. With an exact target, the proposed estimator achieved relative risks of approximately 2.09–2.15 against the MLE. Its benefit was conditional on target quality: in the baseline setting, a local sensitivity study placed the approximate break-even target-shift interval at $-0.081 < R_0 - R < 0.122$. Under exponential random censoring, relative risk remained about 2.10. A quadrature-size study showed negligible change between $B = 500$ and $B = 1000$, and a computational benchmark found the stratified Gaussian tuning close to random-Gaussian and nonparametric pair-bootstrap tuning but substantially faster because it avoids repeated model refitting. In 500 repeated heart-transplant pilot/main splits, the internally estimated Gaussian-risk ratio was 1.807 and median k_G was 0.631. The method is therefore useful when target information is credible, but it should use $k = 1$ and reduce to the MLE when no defensible target is available; it is not uniformly superior under target misspecification.

Keywords Gaussian bootstrap, log-logistic distribution, logit shrinkage, precautionary loss, reliability function, right censoring

AMS 2010 subject classifications 62N05, 62F10, 62P30

DOI: 10.19139/soic-2310-5070-4589

1. Introduction

Many reliability and survival data sets contain positive, skewed times together with observations that end before failure. The log-logistic model is useful in this setting because its distribution and survival functions are available in closed form, while its hazard can be decreasing or unimodal according to the shape parameter [1]. Small samples can still produce biased or numerically sensitive maximum likelihood estimates, as reported in studies of smooth bias correction and robust estimation [2, 5].

Right-censored observations must be retained in the likelihood because each one states that the lifetime exceeded the recorded time. Parametric analyses of incomplete field-failure data discuss this issue in detail [3]. For the log-logistic model, recent work has considered progressive censoring, interval censoring, and accelerated life tests [4, 7, 8]. Related extensions have also been proposed for data that require different tail or hazard patterns [9]. Recent work has also considered frequentist and Bayesian predictive inference for the log-logistic model under progressive

*Correspondence to: Rusul Ismail Khalil (Email: basicmathte7@uodiyala.edu.iq). Department of Mathematics, College of Basic Education, Diyala University, Diyala, Iraq. ORCID: 0009-0006-2032-9497.

Type-II censoring [18]. Recent SOIC contributions demonstrate continuing interest in reliability estimation under censoring, including progressive first-failure censoring, Type-II censoring, and E-Bayesian reliability analysis [19, 20, 21].

1.1. Regularization and shrinkage for small-sample reliability estimation

Small samples and heavy censoring motivate several stabilization strategies. Penalized-likelihood approaches modify the likelihood itself; for example, Firth-type penalization has been studied for Weibull, log-normal, and log-logistic accelerated-failure-time models and can reduce small-sample bias and convergence failures [17]. Bayesian methods instead combine a full likelihood with prior distributions and propagate prior uncertainty through a posterior distribution; recent log-logistic work under progressive Type-II censoring illustrates that route [18, 8]. Shrinkage procedures occupy a different position: they retain a classical sample estimator but borrow toward a target when external information is credible. In reliability problems, pre-test and adaptive shrinkage have been studied for Burr XII, Pareto, and exponentiated-Weibull models under censoring and precautionary risk [10, 11, 12, 13].

The present method does not penalize the two log-logistic parameters and does not require a prior distribution for them. Instead, it targets the scalar mission-time reliability, performs shrinkage on the logit scale, and chooses the amount of borrowing by an estimated-risk criterion. This distinction is important when the available prior information is naturally expressed as an engineering reliability target R_0 rather than as a joint prior distribution for (β, λ) .

1.2. Motivation and main contributions

The central practical problem is not merely how to shrink, but how much to trust a target whose accuracy is uncertain. Fixed rules can have very low risk when the target is exactly correct, yet they can fail sharply after target misspecification. The proposed rule addresses this by allowing $k \in [0, 1]$ to be selected from the data, with $k = 0$ corresponding to full borrowing and $k = 1$ to the classical MLE.

The contributions are: (i) a probability-respecting logit-shrinkage estimator for right-censored log-logistic reliability; (ii) continuous tuning of k by a stratified Gaussian risk approximation based on the observed information matrix; (iii) theoretical existence, quadrature consistency, and an explicit large-sample MLE-reversion result for a fixed incorrect target; (iv) simulations that separate exact-target, misspecified-target, administrative-censoring, random-censoring, boundary-aware, and large-sample behavior; (v) explicit break-even and B -sensitivity analyses plus a computational comparison with random-Gaussian and nonparametric pair bootstraps; and (vi) a repeated-split heart-transplant illustration with pilot-size and mission-time sensitivity. The resulting claim is deliberately conditional: the method can improve reliability estimation when the target is informative, but it is not uniformly superior to the MLE.

2. Log-Logistic Model and Right-Censored Likelihood

Let X follow a two-parameter log-logistic distribution with shape $\beta > 0$ and scale $\lambda > 0$. Its density, distribution, reliability, and hazard functions are [1]

$$f(x; \beta, \lambda) = \frac{\beta}{\lambda} \left(\frac{x}{\lambda} \right)^{\beta-1} \left[1 + \left(\frac{x}{\lambda} \right)^{\beta} \right]^{-2}, \quad x > 0, \quad (1)$$

$$F(x; \beta, \lambda) = \frac{(x/\lambda)^{\beta}}{1 + (x/\lambda)^{\beta}}, \quad R(t; \beta, \lambda) = \frac{1}{1 + (t/\lambda)^{\beta}}, \quad (2)$$

$$h(t; \beta, \lambda) = \frac{\beta t^{\beta-1}}{\lambda^{\beta} + t^{\beta}}. \quad (3)$$

The inverse-transform representation used in simulation is

$$X = \lambda \left(\frac{U}{1-U} \right)^{1/\beta}, \quad U \sim \text{Uniform}(0, 1). \quad (4)$$

Let C_i be independent censoring times and observe

$$Y_i = \min(X_i, C_i), \quad \delta_i = I(X_i \leq C_i), \quad i = 1, \dots, n. \quad (5)$$

The likelihood and log-likelihood are

$$L(\beta, \lambda) = \prod_{i=1}^n f(y_i; \beta, \lambda)^{\delta_i} R(y_i; \beta, \lambda)^{1-\delta_i}, \quad (6)$$

$$\ell(\beta, \lambda) = \sum_{i=1}^n \delta_i \left[\log \beta - \log \lambda + (\beta - 1) \log \left(\frac{y_i}{\lambda} \right) \right] \quad (7)$$

$$- \sum_{i=1}^n (1 + \delta_i) \log \left[1 + \left(\frac{y_i}{\lambda} \right)^\beta \right]. \quad (8)$$

Numerical maximization uses $a = \log \beta$ and $e = \log \lambda$. With $\theta = (a, e)^\top$, the observed information is

$$\mathcal{I}(\hat{\theta}) = -\nabla_{\hat{\theta}}^2 \ell(\hat{\theta}), \quad (9)$$

and the classical plug-in reliability estimator is

$$\hat{R}_c(t) = \frac{1}{1 + (t/\hat{\lambda})^\beta}. \quad (10)$$

3. Shrinkage Estimators

Suppose an independently obtained target $R_0 \in (0, 1)$ is available. The fixed linear rule is

$$\tilde{R}_k^{(lin)}(t) = k\hat{R}_c(t) + (1-k)R_0, \quad 0 \leq k \leq 1. \quad (11)$$

The reliability logit is

$$\psi(t) = \text{logit}\{R(t)\} = \log \frac{R(t)}{1-R(t)} = -\beta \log \left(\frac{t}{\lambda} \right). \quad (12)$$

Let $\hat{\psi} = \text{logit}(\hat{R}_c)$ and $\psi_0 = \text{logit}(R_0)$. The logit-shrinkage family is

$$\tilde{R}_k^{(logit)}(t) = \text{expit}\{k\hat{\psi}(t) + (1-k)\psi_0\}. \quad (13)$$

Proposition 1

If $R_0, \hat{R}_c(t) \in (0, 1)$ and $k \in [0, 1]$, then $0 < \tilde{R}_k^{(logit)}(t) < 1$, with $\tilde{R}_0^{(logit)} = R_0$ and $\tilde{R}_1^{(logit)} = \hat{R}_c$.

Proof

Both logits are finite. Their convex combination is therefore finite, and $\text{expit}(x) \in (0, 1)$ for every finite real x . Substitution of $k = 0$ and $k = 1$ gives the endpoints. $\square$

For comparison, with $\hat{V}_R$ obtained by the delta method,

$$A_R = \max\{(\hat{R}_c - R_0)^2 - \hat{V}_R, 0\}, \quad \hat{k}_A = \frac{A_R}{A_R + \hat{V}_R}, \quad (14)$$

with $\tilde{R}_A = \text{expit}\{\hat{k}_A\hat{\psi} + (1-\hat{k}_A)\psi_0\}$.

3.1. Target availability and MLE fallback

The method is intended for a defensible external target, such as reliability from a comparable historical study or an engineering specification. If no such target is available, an arbitrary generic target should not be manufactured. Operationally one sets $k = 1$, in which case the logit-shrinkage estimator is exactly $\hat{R}_c$. Thus the no-target version of the procedure is the classical MLE rather than an unrelated generic regularizer. A large discrepancy between R_0 and $\hat{R}_c$ is also a warning that borrowing may be risky; Section 5.4 develops a scenario-specific break-even analysis and a practical sensitivity diagnostic.

4. Precautionary Loss and Gaussian-Bootstrap Tuning

4.1. Loss function and decision interpretation

For an estimator $T > 0$ of $R > 0$, the precautionary loss of Norström [14] is

$$L(R, T) = \left(\sqrt{T/R} - \sqrt{R/T} \right)^2 = \frac{T}{R} + \frac{R}{T} - 2. \quad (15)$$

Its derivatives are

$$\frac{\partial L}{\partial T} = \frac{1}{R} - \frac{R}{T^2}, \quad \frac{\partial^2 L}{\partial T^2} = \frac{2R}{T^3} > 0, \quad (16)$$

so the loss is strictly convex in T and uniquely minimized at $T = R$.

The loss is asymmetric for equal additive deviations $T = R \pm d$: a sufficiently severe downward error receives a larger penalty because of the R/T term. It is, however, symmetric for reciprocal multiplicative errors because $L(R, T)$ depends on $T/R + (T/R)^{-1}$. Therefore this loss should not be interpreted as universally appropriate for every engineering decision. Underestimating reliability can drive premature replacement, excessive inspection, or unnecessary redundancy, whereas overestimating reliability can be the more serious safety error in other settings. If a particular application assigns a larger directional cost to overestimation, the same tuning framework can be used with a different continuous loss function; the numerical conclusions in this paper are specific to precautionary loss.

4.2. Gaussian-bootstrap approximation

Write $\theta = (a, e)^\top$ and

$$\psi(\theta; t) = -\exp(a)(\log t - e), \quad g_\psi(\theta; t) = \begin{pmatrix} -\exp(a)(\log t - e) \\ \exp(a) \end{pmatrix}. \quad (17)$$

Under the usual regularity conditions,

$$\hat{V}_\psi(t) = g_\psi(\hat{\theta}; t)^\top \mathcal{I}(\hat{\theta})^{-1} g_\psi(\hat{\theta}; t). \quad (18)$$

The local Gaussian bootstrap approximates the sampling distribution of the reliability logit by

$$\psi_b^G(t) = \hat{\psi}(t) + \sqrt{\hat{V}_\psi(t)} z_b, \quad z_b = \Phi^{-1} \left(\frac{b - 1/2}{B} \right), \quad b = 1, \dots, B. \quad (19)$$

For candidate k ,

$$\tilde{R}_{k,b}^G(t) = \text{expit}\{k\psi_b^G(t) + (1-k)\psi_0\}, \quad (20)$$

and the estimated precautionary risk is

$$\hat{\mathcal{R}}_G(k) = \frac{1}{B} \sum_{b=1}^B \left[\frac{\tilde{R}_{k,b}^G}{\hat{R}_c} + \frac{\hat{R}_c}{\tilde{R}_{k,b}^G} - 2 \right]. \quad (21)$$

The tuned coefficient and estimator are

$$\hat{k}_G = \arg \min_{0 \leq k \leq 1} \hat{\mathcal{R}}_G(k), \quad \tilde{R}_G(t) = \text{expit}\{\hat{k}_G \hat{\psi}(t) + (1 - \hat{k}_G) \psi_0\}. \quad (22)$$

The minimization is continuous and bounded, not a discrete search over candidate coefficients.

The term ‘‘Gaussian bootstrap’’ here refers to a local parametric approximation on the scalar logit scale, not to resampling the observed censored pairs; the distinction follows the usual separation between parametric/normal approximations and resampling-based bootstrap procedures [15]. Its computational purpose is to avoid re-estimating (β, λ) hundreds of times inside every candidate- k evaluation. The midpoint normal quantiles act as deterministic quadrature nodes: they remove the Monte Carlo integration noise that would remain if B independent normal variates were drawn at each fit. The approximation is most defensible when the observed-information normal approximation is locally adequate; it is not claimed to dominate a full nonparametric bootstrap in all finite samples.

Proposition 2 (Existence)

At least one minimizer $\hat{k}_G \in [0, 1]$ exists.

Proof

For fixed data, $\hat{R}_c, R_0 \in (0, 1)$ and every z_b is finite. Hence $k \mapsto \tilde{R}_{k,b}^G$ is continuous and remains in $(0, 1)$. The precautionary loss is continuous on this image, so every summand and their finite average $\hat{\mathcal{R}}_G(k)$ are continuous on the compact interval $[0, 1]$. The extreme-value theorem therefore guarantees that the minimum is attained. $\square$

Proposition 3 (Conditional quadrature consistency)

For a fixed fitted data set, let $0 < \hat{R}_c, R_0 < 1$ and $0 \leq \hat{V}_\psi < \infty$. Define

$$Q_G(k) = E_Z L\left(\hat{R}_c, \text{expit}\{\psi_0 + k(\hat{\psi} + \sqrt{\hat{V}_\psi} Z - \psi_0)\}\right), \quad Z \sim N(0, 1).$$

Let $Q_B(k)$ be its midpoint-quantile approximation obtained with $z_b = \Phi^{-1}((b - 1/2)/B)$. Then $\sup_{k \in [0, 1]} |Q_B(k) - Q_G(k)| \rightarrow 0$ as $B \rightarrow \infty$. If Q_G has a unique minimizer k_G^* , every sequence of minimizers $\hat{k}_{G,B}$ satisfies $\hat{k}_{G,B} \rightarrow k_G^*$.

Proof

Put $h(k, z) = L\{\hat{R}_c, \text{expit}[\psi_0 + k(\hat{\psi} + \sqrt{\hat{V}_\psi} z - \psi_0)]\}$. The map $(k, z) \mapsto h(k, z)$ is continuous. Uniformly in $k \in [0, 1]$, $h(k, z)$ is bounded by a constant multiple of $1 + \exp(c|z|)$ for finite $c = \sqrt{\hat{V}_\psi}$, which is integrable under the standard normal density. Fix $M < \infty$. On $[0, 1] \times [-M, M]$, uniform continuity makes the midpoint-quantile Riemann sum converge uniformly in k to the truncated Gaussian integral. The normal-tail contribution outside $[-M, M]$ is uniformly controlled by the integrable envelope. Letting first $B \rightarrow \infty$ and then $M \rightarrow \infty$ gives uniform convergence. The final statement follows from the standard argmin theorem on the compact set $[0, 1]$. $\square$

Proposition 4 (Reversion under a fixed incorrect target)

Assume $\hat{\psi} \xrightarrow{P} \psi$, $\hat{V}_\psi \xrightarrow{P} 0$, and a fixed target $R_0 \neq R(t)$ is used. If the Gaussian-risk minimizer is unique with probability tending to one, then $\hat{k}_G \xrightarrow{P} 1$ and $\tilde{R}_G(t) - \hat{R}_c(t) \xrightarrow{P} 0$.

Proof

For $k = 1$, the Gaussian draws collapse to ψ as $\hat{V}_\psi \rightarrow 0$, so the limiting loss is zero. For any $k < 1$, the limiting shrunken logit is $\psi_0 + k(\psi - \psi_0) \neq \psi$ because $\psi_0 \neq \psi$. Strict convexity of the loss then gives a strictly positive limiting loss. On every compact interval $[0, 1 - \varepsilon]$, this separation is uniform. Hence the minimizer must enter $(1 - \varepsilon, 1]$ with probability tending to one for each $\varepsilon > 0$, proving $\hat{k}_G \rightarrow 1$; continuity of the inverse logit gives the estimator result. If $R_0 = R(t)$ exactly, this argument does not apply because the limiting criterion is not separated away from zero for $k < 1$, so persistent borrowing is possible even at large n . $\square$

4.3. Quadrature size and comparison with standard bootstrap

The B -sensitivity experiment used 3000 common fitted samples at $n = 50$, 20% administrative censoring, and $R(t) = 0.5$. Table 1 shows that moving from $B = 500$ to $B = 1000$ changes mean k_G by only about 2×10^{-4} in both the exact-target and +0.10 target-shift settings. Thus $B = 500$ provides a stable numerical integration grid for the reported design.

Table 1. Sensitivity of Gaussian-bootstrap tuning to B ($N = 3000$).

Shift	B	Risk	RR	Mean k
+0.0	100	0.00798	2.0935	0.3483
+0.0	250	0.00795	2.1020	0.3470
+0.0	500	0.00794	2.1049	0.3466
+0.0	1000	0.00793	2.1064	0.3464
+0.1	100	0.01545	1.0811	0.6239
+0.1	250	0.01545	1.0815	0.6227
+0.1	500	0.01545	1.0816	0.6222
+0.1	1000	0.01545	1.0817	0.6220

A second benchmark used 30 representative censored data sets, $B = 500$ draws or nodes per tuning, and a +0.10 target shift. The stratified Gaussian rule and random-Gaussian rule produced almost identical mean coefficients (0.5384 versus 0.5376); the mean absolute coefficient difference was 0.0022. The nonparametric pair bootstrap gave mean $k = 0.5274$, with mean absolute difference 0.0147 from the stratified Gaussian rule. Table 2 reports representative tuning-overhead timing after the initial model fit. On the analysis machine, the two Gaussian calculations were both sub-millisecond, whereas pair resampling was hundreds of times slower because each bootstrap sample requires a new censored MLE fit. These times are machine-dependent and are included only to quantify the computational trade-off, not as universal speed constants.

Table 2. Representative tuning-overhead benchmark after the initial model fit ($B = 500$, 30 data sets).

Method	Mean k	Mean conditional SD	Time per tuning
Stratified Gaussian	0.5384	0 (deterministic)	0.161 ms
Random Gaussian	0.5376	0.0139	0.135 ms
Nonparametric pair bootstrap	0.5274	0.0118	145.0 ms

5. Monte Carlo Study

5.1. Design and reproducibility

Administrative censoring is imposed at $C_i = c$, where for $0 < q < 1$,

$$c = \lambda \left(\frac{1-q}{q} \right)^{1/\beta}, \quad (23)$$

and $c = \infty$ at $q = 0$. The core reported conditions use $N = 10000$ independent Monte Carlo replications and $B = 500$ stratified Gaussian quantiles. Computationally heavier sensitivity analyses use smaller replication counts that are stated explicitly. Numerical optimization is conducted in $(\log \beta, \log \lambda)$ coordinates.

The primary interpretive comparison is between the two genuinely adaptive rules—the positive-part rule and the Gaussian-bootstrap rule. Fixed linear and fixed-logit rules are retained as secondary reference/oracle comparators because they illustrate the price of borrowing a fixed amount from a misspecified target.

Table 3. Monte Carlo design and sensitivity settings.

Component	Values
Monte Carlo replications	$N = 10000$ for core reported conditions
Gaussian grid	$B = 500$ stratified standard-normal quantiles
Sample sizes	$n = 20, 50, 100$; large-sample sensitivity: 500, 1000
Shape and scale	Baseline $\beta = 2, \lambda = 1$; shape sensitivity $\beta = 0.75, 1.5, 3$
Censoring	Administrative $q = 0, 0.20, 0.40$; independent exponential random censoring
Mission reliabilities	$R(t) = 0.8, 0.5, 0.2$
Target shifts	$R_0 - R = -0.20, -0.10, 0, 0.10, 0.20$ plus local break-even grid
Core seeds	20260721 (main), 271828 (shape), 888 (application)
Sensitivity-analysis seeds	20260819–20260824

For estimator T ,

$$\text{Bias}(T) = N^{-1} \sum_{s=1}^N (T_s - R), \quad \text{MSE}(T) = N^{-1} \sum_{s=1}^N (T_s - R)^2, \quad (24)$$

$$\widehat{\mathcal{R}}(T) = N^{-1} \sum_{s=1}^N L(R, T_s), \quad RR(T) = \frac{\widehat{\mathcal{R}}(\widehat{R}_c)}{\widehat{\mathcal{R}}(T)}. \quad (25)$$

Thus $RR > 1$ favors the competing estimator and $RR < 1$ favors the MLE.

5.2. Sample size

Table 4 reports the exact-target case at 20% administrative censoring and $R(t) = 0.5$. Fixed rules are an oracle benchmark because 70% of their weight is placed on the true target. Among adaptive rules, the positive-part rule is strongest in this special setting, while Gaussian-bootstrap RR changes only from 2.094 to 2.150 over $n = 20$ to 100.

Table 4. Exact-target results at 20% administrative censoring and $R(t) = 0.5$.

n	Method	Bias	MSE	Risk (MCSE)	RR	Mean k
20	Classical MLE	-0.0003	0.0099	0.0453 (0.0008)	1.0000	—
20	Linear $k = 0.3$	-0.0001	0.0009	0.0036 (0.0001)	12.5406	—
20	Logit $k = 0.3$	-0.0001	0.0010	0.0039 (0.0001)	11.5443	—
20	Adaptive logit	0.0000	0.0040	0.0183 (0.0006)	2.4684	0.1705
20	Gaussian-bootstrap	0.0015	0.0051	0.0216 (0.0005)	2.0943	0.3520
50	Classical MLE	0.0011	0.0038	0.0160 (0.0002)	1.0000	—
50	Linear $k = 0.3$	0.0003	0.0003	0.0014 (0.0000)	11.5065	—
50	Logit $k = 0.3$	0.0003	0.0004	0.0014 (0.0000)	11.1810	—
50	Adaptive logit	0.0009	0.0014	0.0056 (0.0002)	2.8415	0.1589
50	Gaussian-bootstrap	0.0016	0.0019	0.0075 (0.0002)	2.1301	0.3498
100	Classical MLE	-0.0011	0.0019	0.0079 (0.0001)	1.0000	—
100	Linear $k = 0.3$	-0.0003	0.0002	0.0007 (0.0000)	11.4159	—
100	Logit $k = 0.3$	-0.0003	0.0002	0.0007 (0.0000)	11.2573	—
100	Adaptive logit	-0.0005	0.0006	0.0027 (0.0001)	2.9332	0.1553
100	Gaussian-bootstrap	-0.0005	0.0009	0.0037 (0.0001)	2.1497	0.3480

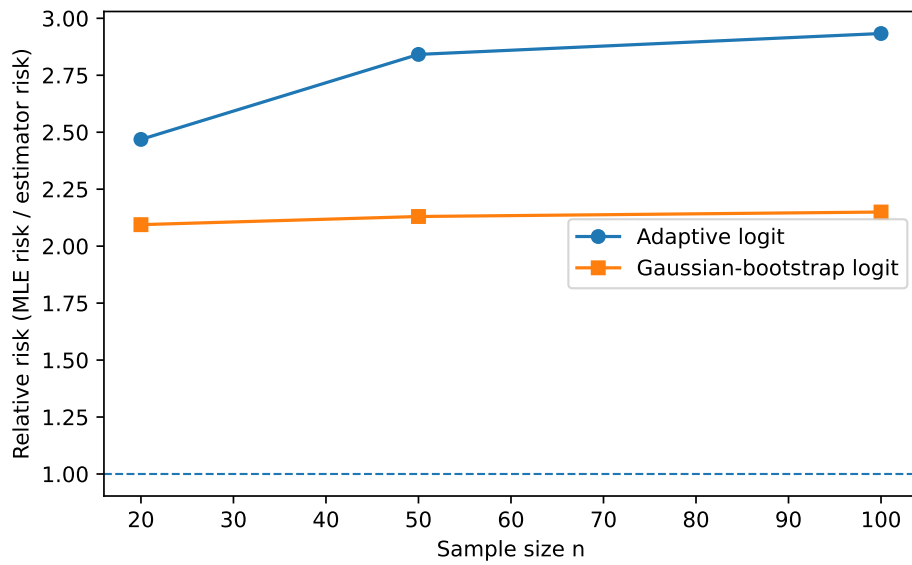


Figure 1. Relative-risk behavior across sample sizes; fixed oracle rules are omitted to emphasize the two adaptive methods.

5.3. Administrative and random censoring

Under administrative censoring, the realized censoring proportions were 0.000, 0.201, and 0.401. The Gaussian-bootstrap RR remained between 2.112 and 2.130, while mean k_G changed only from 0.347 to 0.352.

Table 5. Adaptive-method comparison under administrative censoring ($n = 50$, exact target).

Target q	Observed	Method	Risk (MCSE)	RR	Mean k
0.00	0.000	Classical MLE	0.0158 (0.0003)	1.0000	—
0.00	0.000	Adaptive logit	0.0057 (0.0002)	2.7902	0.1570
0.00	0.000	Gaussian-bootstrap	0.0075 (0.0002)	2.1115	0.3474
0.20	0.201	Classical MLE	0.0160 (0.0002)	1.0000	—
0.20	0.201	Adaptive logit	0.0056 (0.0002)	2.8415	0.1589
0.20	0.201	Gaussian-bootstrap	0.0075 (0.0002)	2.1301	0.3498
0.40	0.401	Classical MLE	0.0180 (0.0003)	1.0000	—
0.40	0.401	Adaptive logit	0.0065 (0.0002)	2.7743	0.1597
0.40	0.401	Gaussian-bootstrap	0.0085 (0.0002)	2.1220	0.3518

To test a distinct censoring mechanism, independent censoring times were also drawn from an exponential distribution. Its rate was numerically calibrated to obtain target censoring proportions 0.20 and 0.40. Across $N = 3000$ replications, the realized proportions were 0.199 and 0.399. The Gaussian-bootstrap RR values were 2.101 and 2.114, respectively, which closely match the administrative-censoring results and support applicability beyond a fixed censoring time.

5.4. Target misspecification and practical break-even behavior

Table 7 shows that target quality is the main limitation of shrinkage. The Gaussian-bootstrap rule has $RR = 2.130$ at the exact target and remains slightly above one at a $+0.10$ shift, but it is inferior to the MLE for several larger or directionally unfavorable shifts. The important adaptive response is that k_G moves toward one as disagreement increases: its mean rises from 0.350 at the exact target to 0.937 at -0.20 and 0.861 at $+0.20$. By contrast, the fixed rules cannot reduce borrowing and their risks deteriorate much more sharply.

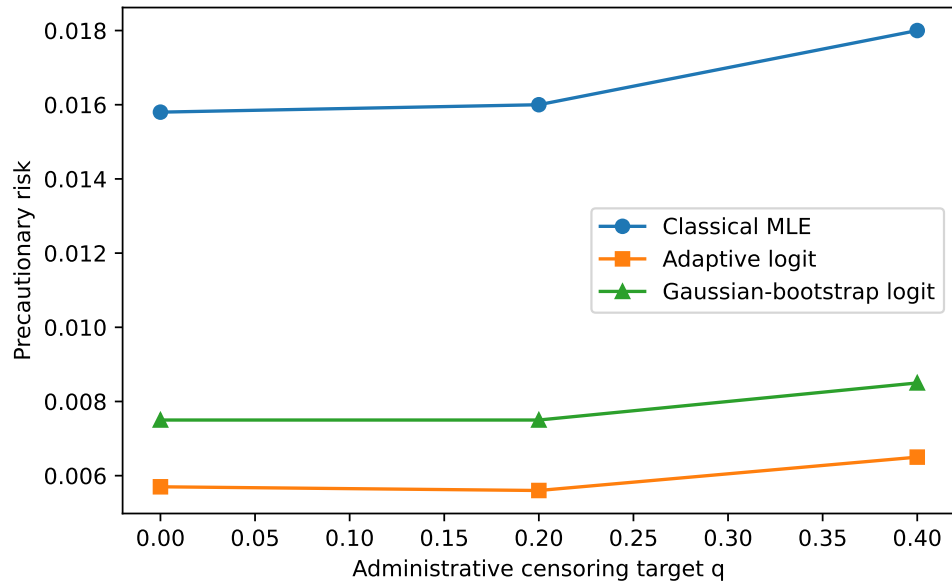


Figure 2. Precautionary risk under increasing administrative censoring.

Table 6. Exact-target results under independent exponential random censoring ($n = 50$, $R(t) = 0.5$, $N = 3000$).

Target q	Observed	Method	Risk (MCSE)	RR	Mean k
0.2	0.199	Classical MLE	0.0185 (0.0006)	1.000	—
0.2	0.199	Adaptive logit	0.0068 (0.0004)	2.720	0.166
0.2	0.199	Gaussian-bootstrap logit	0.0088 (0.0004)	2.101	0.355
0.4	0.399	Classical MLE	0.0226 (0.0007)	1.000	—
0.4	0.399	Adaptive logit	0.0086 (0.0005)	2.637	0.161
0.4	0.399	Gaussian-bootstrap logit	0.0107 (0.0005)	2.114	0.348

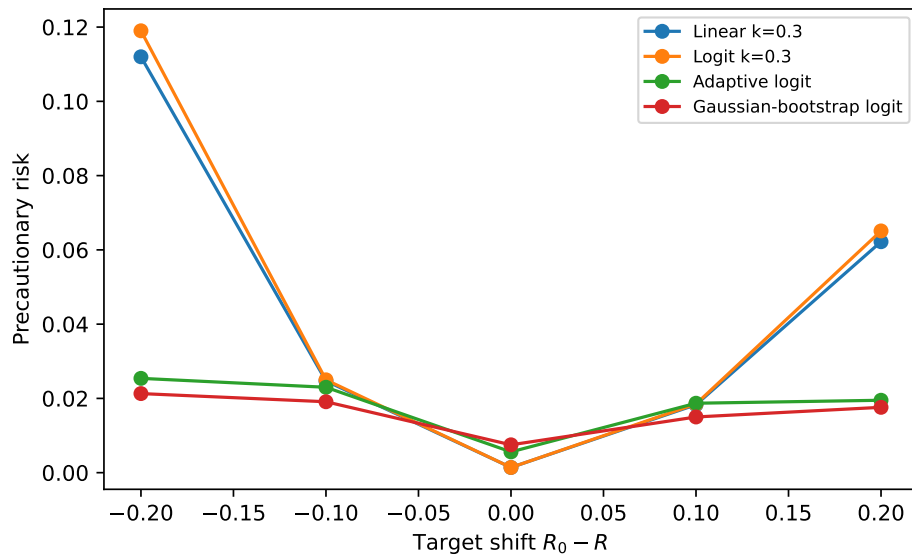
Figure 3. Precautionary risk under target misspecification. The horizontal variable is explicitly $R_0 - R$.

Table 7. Target sensitivity for $n = 50$, 20% administrative censoring, and $R(t) = 0.5$.

$R_0 - R$	Method	Bias	Risk (MCSE)	RR	Mean k
-0.20	Classical MLE	0.0011	0.0160 (0.0002)	1.0000	—
-0.20	Linear $k = 0.3$	-0.1397	0.1120 (0.0004)	0.1425	—
-0.20	Logit $k = 0.3$	-0.1436	0.1190 (0.0003)	0.1341	—
-0.20	Adaptive logit	-0.0205	0.0254 (0.0004)	0.6274	0.8642
-0.20	Gaussian-bootstrap	-0.0076	0.0213 (0.0003)	0.7484	0.9366
-0.10	Classical MLE	0.0011	0.0160 (0.0002)	1.0000	—
-0.10	Linear $k = 0.3$	-0.0697	0.0247 (0.0001)	0.6450	—
-0.10	Logit $k = 0.3$	-0.0700	0.0250 (0.0001)	0.6385	—
-0.10	Adaptive logit	-0.0299	0.0230 (0.0002)	0.6939	0.4983
-0.10	Gaussian-bootstrap	-0.0171	0.0191 (0.0002)	0.8360	0.6623
+0.00	Classical MLE	0.0011	0.0160 (0.0002)	1.0000	—
+0.00	Linear $k = 0.3$	0.0003	0.0014 (0.0000)	11.5065	—
+0.00	Logit $k = 0.3$	0.0003	0.0014 (0.0000)	11.1810	—
+0.00	Adaptive logit	0.0009	0.0056 (0.0002)	2.8415	0.1589
+0.00	Gaussian-bootstrap	0.0016	0.0075 (0.0002)	2.1301	0.3498
+0.10	Classical MLE	0.0011	0.0160 (0.0002)	1.0000	—
+0.10	Linear $k = 0.3$	0.0703	0.0183 (0.0001)	0.8726	—
+0.10	Logit $k = 0.3$	0.0707	0.0185 (0.0001)	0.8647	—
+0.10	Adaptive logit	0.0319	0.0187 (0.0002)	0.8509	0.4896
+0.10	Gaussian-bootstrap	0.0256	0.0150 (0.0002)	1.0647	0.6127
+0.20	Classical MLE	0.0011	0.0160 (0.0002)	1.0000	—
+0.20	Linear $k = 0.3$	0.1403	0.0622 (0.0001)	0.2566	—
+0.20	Logit $k = 0.3$	0.1442	0.0651 (0.0001)	0.2450	—
+0.20	Adaptive logit	0.0230	0.0195 (0.0003)	0.8195	0.8593
+0.20	Gaussian-bootstrap	0.0252	0.0176 (0.0002)	0.9090	0.8613

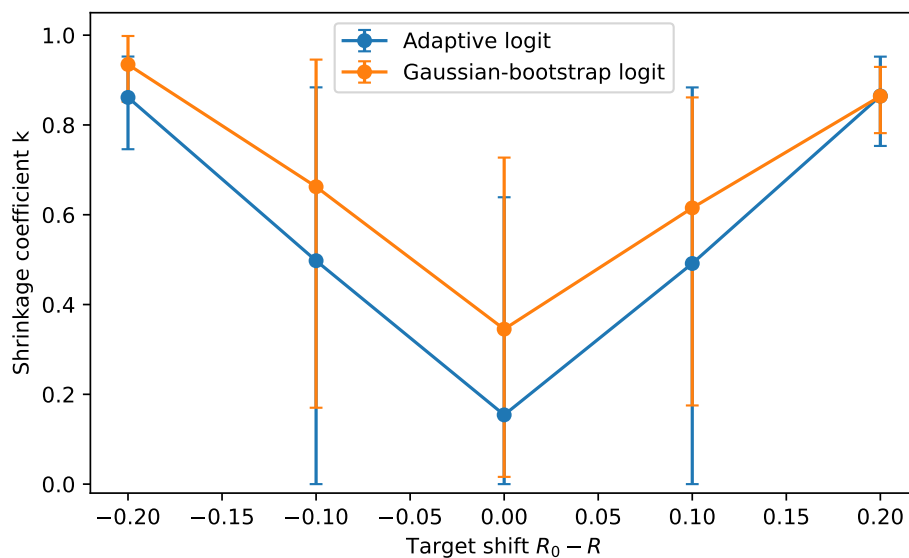


Figure 4. Response of adaptive coefficients to target misspecification. Bars show the 10th–90th percentile range across $N = 10000$ Monte Carlo fits. Their width reflects genuine data-to-data variability in the selected coefficient, especially near transition regions where target discrepancy and sampling variance compete.

A denser local grid was used to quantify the break-even behavior explicitly. Using $N = 3000$ common simulated fits and linear interpolation between neighboring grid points, $RR = 1$ occurred at approximately

$$R_0 - R = -0.081 \quad \text{and} \quad R_0 - R = +0.122.$$

Hence, for this specific baseline design, the proposed estimator was risk-favorable roughly over $-0.081 < R_0 - R < 0.122$. These values are not universal thresholds: they depend on n , censoring, mission reliability, model shape, and the loss function. They should therefore be used as an illustration of how to construct a target-validity envelope, not as a fixed rule for other data sets.

Table 8. Local target-shift analysis around the break-even region ($N = 3000$, $B = 500$).

$R_0 - R$	Risk	RR	Mean k	Bias
-0.10	0.0194	0.861	0.655	-0.0201
-0.08	0.0166	1.007	0.570	-0.0195
-0.06	0.0135	1.238	0.486	-0.0172
+0.00	0.0079	2.105	0.347	-0.0011
+0.10	0.0154	1.082	0.622	+0.0229
+0.12	0.0166	1.005	0.693	+0.0238
+0.13	0.0170	0.982	0.724	+0.0240
+0.20	0.0177	0.943	0.866	+0.0219

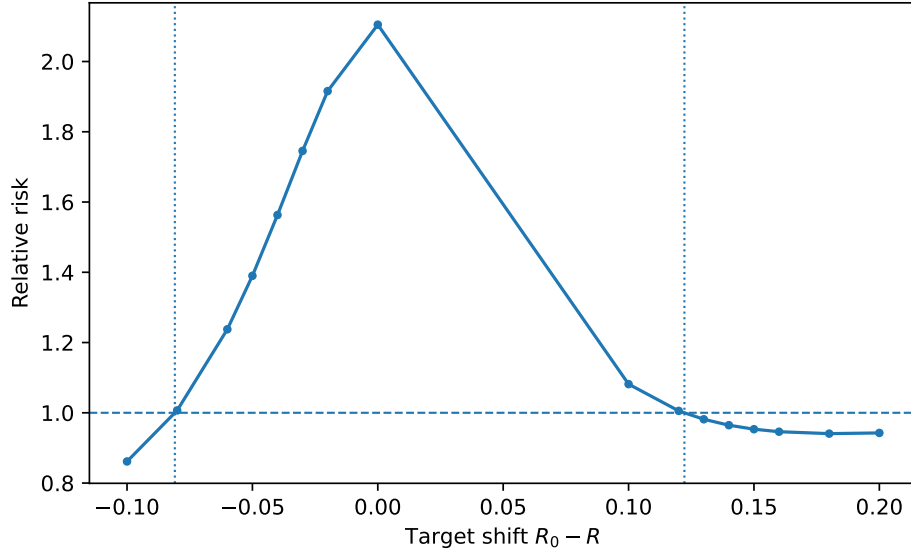


Figure 5. Local break-even analysis. Vertical dotted lines indicate the interpolated points where $RR = 1$.

In practice the true $R(t)$ is unknown, so $R_0 - R$ cannot be observed directly. A more useful diagnostic is to report the target provenance together with the standardized discrepancy

$$D = \frac{|\text{logit}(R_0) - \hat{\psi}|}{\sqrt{\hat{V}_{\psi}}}.$$

Large D indicates that the target and current data are difficult to reconcile. Rather than imposing a universal cutoff, we recommend a local sensitivity analysis around the proposed R_0 and reporting how $\hat{k}_G$ and the resulting

reliability estimate change. The $RR = 1$ crossings above are simulation-calibration quantities because the true $R(t)$ is known there; they are not directly observable in an application. If the target is weakly documented, D is large, or modest plausible perturbations of R_0 materially change the fitted conclusion, the conservative action is to use $k = 1$ and report the MLE.

A boundary-aware design was also examined for $R(t) = 0.2$. A -0.20 shift would put the target at the boundary $R_0 = 0$ and is therefore inadmissible for the logit formulation. We instead considered $-0.10, 0, +0.10, +0.20$. This preserves $R_0 \in (0, 1)$ and makes the target-shift design consistent with the estimator's mathematical domain.

Table 9. Boundary-aware target shifts for $R(t) = 0.2$ ($n = 50$, 20% censoring, $N = 3000$).

$R_0 - R$	Method	Risk (MCSE)	RR	Mean k
-0.1	Classical MLE	0.0661 (0.0020)	1.000	—
-0.1	Adaptive logit	0.1316 (0.0030)	0.502	0.629
-0.1	Gaussian-bootstrap logit	0.0919 (0.0023)	0.719	0.814
+0.0	Classical MLE	0.0661 (0.0020)	1.000	—
+0.0	Adaptive logit	0.0254 (0.0016)	2.598	0.153
+0.0	Gaussian-bootstrap logit	0.0303 (0.0014)	2.185	0.339
+0.1	Classical MLE	0.0661 (0.0020)	1.000	—
+0.1	Adaptive logit	0.0830 (0.0019)	0.797	0.608
+0.1	Gaussian-bootstrap logit	0.0654 (0.0016)	1.011	0.661
+0.2	Classical MLE	0.0661 (0.0020)	1.000	—
+0.2	Adaptive logit	0.0766 (0.0021)	0.863	0.907
+0.2	Gaussian-bootstrap logit	0.0702 (0.0018)	0.942	0.870

5.5. Mission reliability and shape

At the exact target, the classical risk increased as mission reliability decreased, whereas Gaussian-bootstrap RR stayed near 2.1 across $R(t) = 0.8, 0.5, 0.2$. Table 10 emphasizes the MLE and adaptive rules. For shape sensitivity,

Table 10. Exact-target results over mission reliabilities ($n = 50$, 20% censoring), emphasizing the adaptive comparators.

$R(t)$	Method	Risk (MCSE)	RR	Mean k
0.8	Classical MLE	0.0036 (0.0001)	1.0000	—
0.8	Adaptive logit	0.0012 (0.0000)	2.9692	0.1643
0.8	Gaussian-bootstrap	0.0017 (0.0000)	2.1496	0.3480
0.5	Classical MLE	0.0160 (0.0002)	1.0000	—
0.5	Adaptive logit	0.0056 (0.0002)	2.8415	0.1589
0.5	Gaussian-bootstrap	0.0075 (0.0002)	2.1301	0.3498
0.2	Classical MLE	0.0708 (0.0011)	1.0000	—
0.2	Adaptive logit	0.0284 (0.0009)	2.4933	0.1700
0.2	Gaussian-bootstrap	0.0335 (0.0008)	2.1119	0.3502

mission times were chosen so that $R(t) = 0.2$ for each β . The Gaussian-bootstrap RR was 2.120 for $\beta = 0.75$, 2.079 for $\beta = 1.5$, and 2.098 for $\beta = 3$, covering decreasing and unimodal hazard shapes without changing mission reliability.

5.6. Large-sample behavior

The large-sample behavior depends on target correctness. Proposition 4 makes this distinction explicit. Table 12 confirms this distinction. With an exact target, mean k_G remains near 0.35–0.36 at $n = 500$ and $n = 1000$ because the target continues to coincide with the truth. With a fixed $+0.10$ misspecification, mean k_G increases to 0.950 and 0.976 and RR moves toward one, so the proposed estimator approaches the MLE as sampling information accumulates.

Table 11. Shape sensitivity at $n = 50$, 20% censoring, and $R(t) = 0.2$, emphasizing the adaptive comparators.

β	t	Method	Risk (MCSE)	RR	Mean k
0.75	6.350	Classical MLE	0.0701 (0.0012)	1.0000	–
0.75	6.350	Adaptive logit	0.0282 (0.0009)	2.4894	0.1628
0.75	6.350	Gaussian-bootstrap	0.0331 (0.0008)	2.1195	0.3477
1.50	2.520	Classical MLE	0.0730 (0.0012)	1.0000	–
1.50	2.520	Adaptive logit	0.0307 (0.0010)	2.3782	0.1696
1.50	2.520	Gaussian-bootstrap	0.0351 (0.0009)	2.0794	0.3495
3.00	1.587	Classical MLE	0.0717 (0.0012)	1.0000	–
3.00	1.587	Adaptive logit	0.0295 (0.0010)	2.4317	0.1669
3.00	1.587	Gaussian-bootstrap	0.0342 (0.0009)	2.0980	0.3486

Table 12. Large-sample sensitivity at 20% administrative censoring and $R(t) = 0.5$.

n	N	Shift	Risk	RR	Mean k
500	1000	+0.0	0.000724	2.131	0.350
500	1000	+0.1	0.001716	0.899	0.950
1000	800	+0.0	0.000436	1.976	0.359
1000	800	+0.1	0.000901	0.956	0.976

6. Heart-Transplant Survival Application

6.1. Model choice and interpretation

The heart-transplant data contain 69 survival times and censoring indicators and are available in the `statsmodels` data library; the original source is Miller [16]. The full sample was used to compare four parametric models.

Table 13. Parametric model comparison for the heart-transplant data.

Model	log-likelihood	AIC	BIC
Log-normal	-313.72	631.44	635.91
Log-logistic	-314.36	632.71	637.18
Weibull	-316.08	636.15	640.62
Exponential	-331.16	664.32	666.56

The log-normal model has the smallest AIC, 631.44, while the log-logistic AIC is 632.71, a difference of only 1.27. The log-logistic model is therefore not asserted to be uniquely selected by AIC; it is retained because the proposed estimator is derived for log-logistic reliability and the application is intended to illustrate that method. Model uncertainty remains relevant. As a direct sensitivity check, fitted survival probabilities from the log-normal and log-logistic models are 0.5511 versus 0.5495 at 180 days, 0.4294 versus 0.4203 at 365 days, and 0.3161 versus 0.3034 at 730 days. Thus the two best-AIC models give closely similar mission-time survival probabilities over the range used below, although a fully model-averaged shrinkage procedure is beyond the scope of this paper. Similar difficulty in discriminating log-normal from log-logistic models under censoring has been discussed in [6].

6.2. Repeated pilot/main splits

The data were split 500 times into a pilot sample of 21 and a main sample of 48. For each split, the pilot fit supplied R_0 at 365 days, while the main sample supplied $\hat{R}_c$, the observed information, and adaptive coefficients. Because the true reliability for these real data is unknown, “risk” in Table 14 is the model-based Gaussian risk criterion averaged over splits rather than an empirical loss against a known truth.

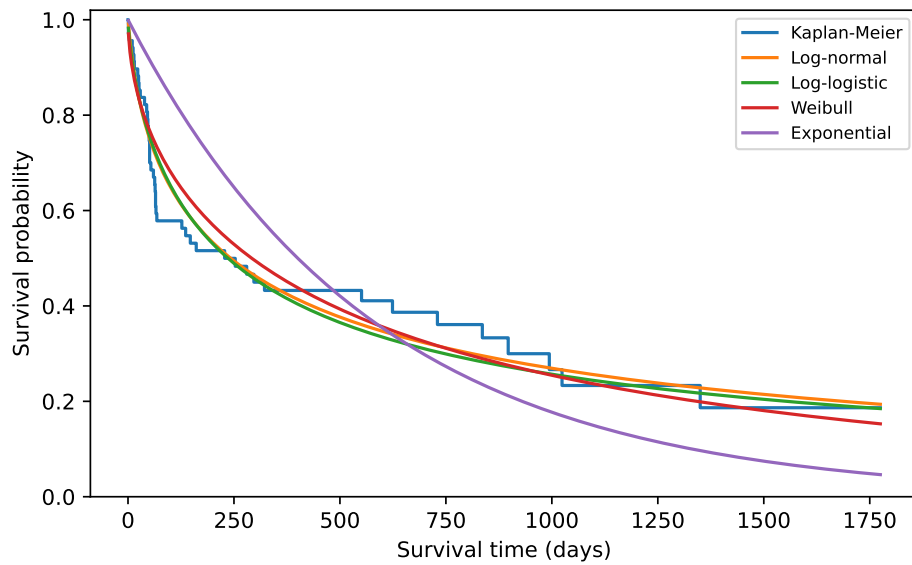


Figure 6. Kaplan–Meier estimate and fitted parametric survival functions.

These splits are an internal cross-validation-style proof of concept, not a substitute for a genuinely external historical target. Splitting the same data deliberately reduces the main-sample size and adds target sampling variability. In the intended external-information use case, R_0 would come from a previous comparable study, a validated engineering specification, or other independent evidence; the present repeated splitting is used only to study the tuning response without making the result depend on one arbitrary partition. The Gaussian-bootstrap

Table 14. Results across 500 repeated pilot/main splits at 365 days, including the positive-part adaptive comparator.

Method	Median estimate	Mean est. risk	Aggregate RR	Median split RR	Median k (IQR)
Classical MLE	0.4227	0.0265	1.0000	1.0000	—
Linear $k = 0.3$	0.4165	0.0505	0.5257	1.0890	—
Logit $k = 0.3$	0.4165	0.0528	0.5023	1.0911	—
Adaptive logit	0.4201	0.0173	1.5322	1.3579	0.4254 (0.0000–0.7876)
Gaussian-bootstrap	0.4226	0.0147	1.8069	1.5985	0.6306 (0.3202–0.8316)

aggregate RR is 1.8069, corresponding to a 44.7% reduction in the averaged internally estimated Gaussian-risk criterion relative to the classical criterion. Its median coefficient is 0.6306. The positive-part adaptive rule also improves this internal criterion, but less strongly in this application (aggregate RR 1.5322). Because $\hat{k}_G$ is selected by minimizing the same model-based Gaussian criterion, this reduction is a proof-of-concept measure of the tuning mechanism and should not be interpreted as an externally validated 44.7% reduction in the unknown true prediction risk.

6.3. Pilot-size and mission-time sensitivity

Pilot size creates a genuine trade-off: a larger pilot stabilizes R_0 but leaves fewer observations in the main fit. Table 15 shows that target standard deviation falls from 0.155 at pilot size 10 to 0.062 at pilot size 35. Correspondingly, median k_G decreases from 0.815 to 0.526, meaning more borrowing when the target is more stable. Aggregate RR increases from 1.479 to 2.075, although part of that increase reflects the higher MLE risk caused by the smaller main sample.

Table 15. Sensitivity to pilot-sample size at 365 days (500 repeated splits).

Pilot	Main	Target SD	Mean MLE risk	Mean GB risk	Aggregate RR	Median k_G
10	59	0.155	0.0214	0.0145	1.479	0.815
21	48	0.089	0.0265	0.0147	1.807	0.631
35	34	0.062	0.0389	0.0188	2.075	0.526

The application was also repeated at 180 and 730 days to test dependence on the originally chosen 365-day mission time. Aggregate RR values were 1.864, 1.807, and 1.791 at 180, 365, and 730 days, respectively, while median k_G ranged from 0.597 to 0.640. The qualitative conclusion is therefore stable across these three mission times.

Table 16. Mission-time sensitivity for pilot size 21 (500 repeated splits).

Mission time	Median MLE	Median GB	Aggregate RR	Median k_G	IQR k_G
180	0.5521	0.5528	1.864	0.597	0.297–0.817
365	0.4227	0.4226	1.807	0.631	0.320–0.832
730	0.3053	0.3044	1.791	0.640	0.331–0.840

7. Discussion and Practical Guidance

The simulations identify three distinct regimes. First, when a target is accurate, both adaptive rules can substantially reduce risk, while fixed rules provide an oracle upper benchmark. Second, when the target is moderately wrong, the benefit may disappear before k_G reaches one; adaptive tuning mitigates but does not eliminate misspecification risk. Third, when the target is strongly discrepant, k_G moves toward one and protects substantially better than fixed borrowing, and Proposition 4 explains why this reversion becomes complete asymptotically for a fixed incorrect target.

For applied use, the target should therefore be documented before fitting. A historical target is most credible when population, component design, censoring definition, mission time, and operating conditions are comparable to the current study. Users should report the target source, its uncertainty if available, the standardized discrepancy D , and a local target-sensitivity analysis. The baseline break-even interval found here is an example of that analysis, not a universal acceptance band. If no defensible external target exists, or if plausible target perturbations make conclusions unstable, the recommended fallback is $k = 1$ and the classical MLE.

The Gaussian approximation is a computational device rather than a claim that censoring uncertainty is exactly normal in every small sample. The benchmark and B -sensitivity results support its numerical stability in the scenarios studied, while a full pair bootstrap remains a useful diagnostic when computational cost is acceptable. Likewise, precautionary loss encodes one error geometry. Applications with explicitly directional safety costs should replace it with a loss that matches those costs and retune k under that loss.

8. Conclusion

This paper develops a logit-scale shrinkage estimator for the mission-time reliability of a two-parameter log-logistic model under independent right censoring. The shrinkage coefficient is selected over $[0, 1]$ by a deterministic stratified Gaussian approximation to precautionary risk based on the observed information matrix. The estimator remains inside the probability range, and the theoretical results establish existence of a minimizer, convergence of

the stratified Gaussian quadrature criterion, and reversion toward the MLE for a fixed incorrect target as information increases.

The numerical evidence supports a conditional rather than universal improvement claim. With an exact target, Gaussian-bootstrap RR was about 2.09–2.15 in the core simulation, and performance was similar under independent exponential random censoring. In the baseline misspecification study, the approximate break-even region was $-0.081 < R_0 - R < 0.122$; outside that region the MLE could have lower risk, although adaptive tuning remained much safer than fixed shrinkage at large target errors. $B = 500$ was numerically stable relative to $B = 1000$, and the stratified Gaussian calculation closely matched random-Gaussian tuning while avoiding its integration noise and the repeated refitting cost of a pair bootstrap.

In the heart-transplant proof-of-concept, 500 repeated pilot/main splits produced an internally estimated Gaussian-risk ratio of 1.807 at 365 days, with similar values at 180 and 730 days; this is an internal tuning criterion rather than an external estimate of true predictive risk. Pilot-size sensitivity showed that the method appropriately borrowed less when the pilot target was more variable. When no credible external target is available, the procedure should use $k = 1$ and reduce exactly to the classical plug-in MLE. Future work should propagate uncertainty from truly external targets, study additional censoring schemes and model averaging, and extend tuning from a single mission time to the full reliability curve.

Data Availability

The heart-transplant data are included in the `statsmodels` data library and originate from Miller [16]. Simulation data are generated from the models and settings reported in this paper. The accompanying reproducibility folder contains portable scripts and the numerical outputs used for the sensitivity analyses.

Conflict of Interest

The author declares no conflict of interest.

Author Contribution

Rusul Ismail Khalil was responsible for the conception, mathematical development, computational design, interpretation of results, and preparation of the manuscript.

REFERENCES

1. S. Bennett, *Log-logistic regression models for survival data*, Journal of the Royal Statistical Society: Series C (Applied Statistics), vol. 32, no. 2, pp. 165–171, 1983. doi: 10.2307/2347295.
2. X. Zheng, J.-Y. Chiang, T.-R. Tsai, and S. Wang, *Estimating the failure rate of the log-logistic distribution by smooth adaptive and bias-correction methods*, Computers & Industrial Engineering, vol. 156, Art. 107188, 2021. doi: 10.1016/j.cie.2021.107188.
3. T. Emura and H. Michimae, *Left-truncated and right-censored field failure data: Review of parametric analysis for reliability*, Quality and Reliability Engineering International, vol. 38, no. 7, pp. 3919–3934, 2022. doi: 10.1002/qre.3161.
4. K. Maiti and S. Kayal, *Estimating reliability characteristics of the log-logistic distribution under progressive censoring with two applications*, Annals of Data Science, vol. 10, no. 1, pp. 89–128, 2023. doi: 10.1007/s40745-020-00292-y.
5. Z. Ma, M. Wang, and C. Park, *Robust explicit estimation of the log-logistic distribution with applications*, Journal of Statistical Theory and Practice, vol. 17, no. 2, Art. 21, 2023. doi: 10.1007/s42519-023-00322-x.
6. B. Diyali, D. Kumar, and S. Singh, *Discriminating between log-normal and log-logistic distributions in the presence of Type-II censoring*, Computational Statistics, vol. 39, pp. 1459–1483, 2024. doi: 10.1007/s00180-023-01351-7.
7. V. P. Vidas, E. M. M. Ortega, A. K. Suzuki, G. M. Cordeiro, and P. C. dos Santos Junior, *The generalized odd log-logistic-G regression with interval-censored survival data*, Journal of Applied Statistics, vol. 51, no. 9, pp. 1642–1663, 2024. doi: 10.1080/02664763.2023.2230533.
8. A. K. Mahto et al., *Bayesian estimation and prediction under progressive-stress accelerated life test for a log-logistic model*, Alexandria Engineering Journal, vol. 101, pp. 330–342, 2024. doi: 10.1016/j.aej.2024.05.045.

9. V. Kariuki et al., *Properties, estimation, and applications of the extended log-logistic distribution*, Scientific Reports, vol. 14, Art. 20967, 2024. doi: 10.1038/s41598-024-68843-4.
10. M. R. Sabr and A. K. Jiheel, *Pre-test shrinkage estimation for reliability function of Burr XII distribution using progressive Type II censored sample under precautionary loss function (PLF)*, Journal of Education for Pure Science, vol. 14, no. 3, 2024. doi: 10.32792/jeps.v14i3.549.
11. Z. A. Al-Hemyari, A. K. Jiheel, and I. J. Atewi, *A class of efficient shrinkage estimators for modelling the reliability of Burr XII distribution*, Mathematics and Statistics, vol. 12, no. 2, pp. 184–198, 2024. doi: 10.13189/ms.2024.120208.
12. O. J. A. Al-Qaragholi and A. K. Jiheel, *Optimization-oriented double pre-test shrinkage estimators for Pareto reliability under progressive Type-II censoring and precautionary loss*, Control and Optimization in Applied Mathematics, vol. 11, no. 2, pp. 83–103, 2026. doi: 10.30473/coam.2026.77066.1388.
13. O. J. A. Al-Qaragholi and A. K. Jiheel, *Adaptive reliability estimation under progressive censoring for industrial decision support and predictive maintenance applications*, Journal of Innovations in Business and Industry, vol. 4, no. 4, pp. 397–406, 2026. doi: 10.61552/JIBI.2026.04.004.
14. J. G. Norström, *The use of precautionary loss functions in risk analysis*, IEEE Transactions on Reliability, vol. 45, no. 3, pp. 400–403, 1996. doi: 10.1109/24.536992.
15. A. C. Davison and D. V. Hinkley, *Bootstrap Methods and Their Application*. Cambridge University Press, 1997. doi: 10.1017/CBO9780511802843.
16. R. G. Miller, *Least squares regression with censored data*, Biometrika, vol. 63, no. 3, pp. 449–464, 1976. doi: 10.1093/biomet/63.3.449.
17. T. F. Alam, M. S. Rahman, and W. Bari, *On estimation for accelerated failure time models with small or rare event survival data*, BMC Medical Research Methodology, vol. 22, Art. 169, 2022. doi: 10.1186/s12874-022-01638-1.
18. Z. Zhang and W. Gui, *Frequentist and Bayesian predictive inference for the log-logistic distribution under progressive Type-II censoring*, Entropy, vol. 28, no. 4, Art. 466, 2026. doi: 10.3390/e28040466.
19. O. Abuelamayem, *Reliability estimation of a multicomponent stress-strength model based on copula function under progressive first failure censoring*, Statistics, Optimization & Information Computing, vol. 12, no. 6, pp. 1601–1611, 2024. doi: 10.19139/soic-2310-5070-1894.
20. O. Shojaee, H. Zarei, and F. Naruei, *E-Bayesian estimation and the corresponding E-MSE under progressive type-II censored data for some characteristics of Weibull distribution*, Statistics, Optimization & Information Computing, vol. 12, no. 4, pp. 962–981, 2023. doi: 10.19139/soic-2310-5070-1709.
21. K. A. Mohammed, B. G. Al-Ani, and Z. Al-Khaledi, *Different methods to estimate reliability for Exponentiated Mukherjee–Islam distribution under type II censoring*, Statistics, Optimization & Information Computing, vol. 15, no. 5, pp. 3555–3571, 2026. doi: 10.19139/soic-2310-5070-3183.