

Proceeding of

The 12th Seminar on
Reliability Theory and its Applications

Department of Statistics
Alzahra University
15 & 16 July 2026



In the Name of God



Proceeding of
the 12th Seminar on
Reliability Theory and its Applications

Department of Statistics
Faculty of Mathematical Sciences
Alzahra University
Tehran, Iran

15-16 July 2026

This book contains the proceeding of the 12th Seminar on **Reliability Theory and its Applications**. Authors are responsible for the contents and accuracy. Opinions expressed may not necessarily reflect the position of the scientific and organizing committees.

Title: Proceeding of the 12th Seminar on Reliability Theory and its Applications

By: Mahdi Alimohammadi, Fariba Azizi, Mohammad Zare

Cover Designer: Fatemeh Zamani

Date: August, 2026

Preface

It is a great honor to host the "12th Seminar on Reliability Theory and its Applications" in the Department of Statistics at Alzahra University on July 15-16, 2026. On behalf of the organizing committee, we extend a warm welcome to all participants. We trust this seminar will serve as a premier platform for engaging discussions and the exchange of innovative scientific ideas within our community. This event is made possible through the unwavering support of numerous institutions and the dedicated contributions of researchers and students whose work defines the quality of this seminar. We would like to express our profound gratitude to the Research Council of Alzahra University, the "Center of Excellence on Ordered Data, Reliability, and Dependency", the Islamic World Science Citation Center (ISC), the Iranian Statistical Society, the Insurance Research Center (IRC), and the Kahkeshan Group. Furthermore, we are deeply indebted to our scientific committee, peer reviewers, and the staff of the Faculty of Mathematical Sciences at Alzahra University for their invaluable cooperation in making this event a success.

Nasrin Soltankhah (Conference Chair)
Mahdi Alimohammadi (Scientific Chair)
Fariba Azizi (Organizing Chair)
August, 2026

Seminar Topics:

The aim of the seminar is to provide a forum for presentation and discussion of scientific works covering theories and methods in the field of reliability and its applications in a wide range of areas:

- Reliability of systems and networks
- Maintenance and repair models of systems
- Reliability analysis based on degradation data
- Accelerated lifetime tests
- Statistical inference and analysis of lifetime data
- Software reliability
- Warranty models and field data
- Ageing concepts
- Dependency concepts and copula functions
- Risk analysis in reliability
- Bayesian methods in reliability
- Stochastic orderings
- Data mining in reliability
- Shock models
- Survival analysis
- Statistical quality control
- Reliability in fuzzy environments
- Case studies in industry
- Big data in reliability
- Artificial intelligence in reliability

Scientific Committee

1. Ahmadi, R., Iran University of Science and Technology
2. Alimohammadi, M., Alzahra University (Scientific Chair)
3. Amini Seresht, E., Bu-Ali Sina University
4. Asgharzadeh, A., University of Mazandaran
5. Ashrafi, S., University of Isfahan
6. Azizi, F., Alzahra University
7. Chahkandi, M., University of Birjand
8. Doostparast, M., Ferdowsi University of Mashhad
9. Hakamipour, N., Buein Zahra Technical University
10. Haghighi, F., University of Tehran
11. Kelkeinnama, M., Isfahan University of Technology
12. Kohansal, A., Imam Khomeini International University
13. Razmkhah, M., Ferdowsi University of Mashhad
14. Sharafi, M., Razi University
15. Tavangar, M., University of Isfahan
16. Torabi, H., Yazd University
17. Zarezadeh, S., Shiraz University

Honorary Scientific Advisory Committee

1. Ahmadi, J., Ferdowsi University of Mashhad
2. Asadi, M., University of Isfahan

3. Khaledi, B., Razi University
4. Khanjari Sadegh, M., Ferdowsi University of Mashhad
5. Khodadadi, A., Shahid Beheshti University
6. Khoram, E., Amirkabir University of Technology
7. Mohtashami Barzadarani, G., Ferdowsi University of Mashhad
8. Rezaei Roknabadi, A., Ferdowsi University of Mashhad
9. Zeinal Hamadani, A., Isfahan University of Technology

Organizing Committee

1. Alimohammadi, M., Alzahra University
2. Azizi, F., Alzahra University (Organizing Chair)
3. Jabari, H., Ferdowsi University of Mashhad
4. Shams, S., Alzahra University
5. Soltankhah, N., Alzahra University (Conference Chair)
6. Zare, M., Alzahra University

Executive Staff

Ms. Fatemeh Zamani (Bachelor's student and Secretary of the Statistics Student Association)

Contents

A Quantum Operator Approach to Reliability and Stochastic Ordering	
Amini-Seresht, E.	1
An Opportunistic Preventive Maintenance Policy for Coherent Systems Subject to Fatal Shocks	
Ammar, H., Ashrafi, S., Asadi, M.	9
A Survival modeling Framework For Mortality Decline in Cancer Populations	
Arabbeik Zefreh, A., Haghighi, F.	17
Testing Goodness-of-Fit for Exponential Distribution Based on Exponential Extropy and Weighted Exponential Extropy	
Askari, M., Yousefzadeh, F.	24
Development of Repair Alert Model for Components Having Discrete Lifetimes in the Power Series Family	
Atlehkhani.M., Doostparast, M.	35
Stochastic Comparisons of Sequential $(n-r+1)$-out-of-n Systems Based on Relative Ageing	
Esna-Ashari, M., Alimohammadi, M.	45
Stochastic Ordering and Reliability Properties of a Discrete Lévy-type Distribution	
Farbod, D.	57
Agent-Informed Discretization:A Knowledge-Informed Reinforcement Learning Framework for Smart Steel Production Maintenance	
Gholami, Z., Haerian, N., Safari, A., Haghighi, F.	63
Multi-State Stress-Strength Reliability Models	
Hajialiakbar, S., Tavangar, M.	79

A Statistical Model for Real Lifespan Under Environmental Conditions: A Gamma-Dirichlet Approach	
Homei, H., Hedayati, A., Kamali Alamdari, M.	88
Maximum Varma entropy distribution with conditional known Lorenz curve constraints	
Imani, M., Mohtashami Borzadaran, G., Amini, M., Khorashadizadeh, M. . .	96
Estimation of Multicomponent Stress–Strength Reliability Based on Lower Record Values from the Topp–Leone Distribution	
Jahanbin, R., Basirat, M., Fathi Vajargah, k.	106
Design of an Modified Exponentially Weighted Moving Average Control Chart	
Moradi Tavalaei, N., Salehi, E.	115
Comparative Evaluation of Group Replacement and Run to Failure Policies in Heterogeneous Multi Component Systems	
Mostafavi Vala, F. S., Rasay, H., Najafi-Ghobadi, S.	127
Optimal Weibull Reliability group test plans using truncated prior distribution	
Naghizadeh Qomi, M., Mahdizadeh, Z., Pérez-González, Carlos J.	132
Stress vs. Strength: A Log-Uniform Model for Quantifying Reliability in Cognitive Performance and Battery Degradation	
Rezaei, Amir, Rezaei, Ahmad	142
Ordering Results of Series Systems Consisting of Dependent and Heterogeneous Exponential Components under Archimedean Copula	
Omid Shojae	151



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



A Quantum Operator Approach to Reliability and Stochastic Ordering

Amini-Seresht, E. ¹

Department of Statistics, Faculty of Science, Bu-Ali Sina University, Hamedan, Iran

Abstract

This paper presents a theoretical framework for generalizing the fundamental concepts of reliability analysis and stochastic order to the realm of quantum computing. The main innovation lies in mapping classical probability distributions to quantum states via amplitude encoding, which enables the representation of key reliability indices such as the failure rate, hazard rate, and mean residual life as quantum operators. We rigorously analyze the properties of these operators and demonstrate that, under appropriate conditions, the well-known hierarchical relationships among classical stochastic orders (including the usual stochastic order, hazard rate order, likelihood ratio order, and convex order) are preserved in the quantum formulation.

Keywords: Quantum reliability, Stochastic orders, Quantum information.

1 Introduction

Reliability analysis and stochastic order theory constitute fundamental branches of statistics and engineering, primarily concerned with studying the time-dependent behavior of systems until the occurrence of a critical event such as failure, and with comparing lifetime variables. In classical reliability theory, key indices such as the survival function

¹Speaker. e.amini@basu.ac.ir

$S(t)$, the hazard rate (or failure rate) $h(t)$, and the mean residual life play a central role in engineering assessment, risk management, and decision-making processes. These indices provide quantitative measures of system aging, reliability, and remaining useful life. Concurrently, stochastic orders such as the usual stochastic order ($\preceq_{\text{st}}$), the hazard rate order ($\preceq_{\text{hr}}$), the likelihood ratio order ($\preceq_{\text{lr}}$), and the convex order ($\preceq_{\text{cx}}$) offer a rigorous mathematical framework for comparing the relative “largeness” or “riskiness” of random variables. These orders form a well-studied hierarchy and serve as powerful analytical tools in reliability theory, actuarial science, and economics.

The theoretical foundations of quantum information and computation have been extensively established by [5], and operator theory has been comprehensively laid out by [2]. In a more specialized domain, [3] founded quantum detection and estimation theory, which directly addresses the problem of estimating quantities in quantum systems. From a mathematical perspective, [2], in the book *Matrix Analysis*, provides the essential tools for working with operators and convex operator functions. Regarding the connection between probability and quantum computing, [7] have provided an introduction to quantum probability and its relation to quantum computation. Furthermore, [4] has examined the applications of singular value decomposition and matrix reorderings in quantum information theory. The advent of quantum computing has revolutionized information processing, offering potential exponential speedups for certain computational tasks. This raises a natural question: can classical reliability concepts be meaningfully extended to the quantum domain, and if so, what computational advantages might such a quantum formulation offer? In this paper, we address this question by developing a systematic framework for quantizing reliability analysis and stochastic order theory. Our approach is based on amplitude encoding, a technique that maps classical probability distributions to quantum states. Under this mapping, classical reliability indices become quantum operators, and stochastic orders translate into relationships between quantum states.

The primary contributions of this work are threefold: (1) We establish a mathematically sound mapping from classical probability distributions to quantum states and define corresponding quantum operators for key reliability indices. (2) We prove that, under a commutativity condition that is automatically satisfied in our diagonal encoding scheme, the classical hierarchy of stochastic orders is preserved in the quantum domain. We also provide counterexamples showing where reverse implications fail, thus characterizing the limits of this preservation. (3) We introduce and analyze quantum algorithms for estimating the expectations of reliability operators, comparing their theoretical complexity and practical performance through numerical simulations.

This paper is organized as follows: Section 2 reviews the necessary background on classical stochastic orders and quantum mechanics. Section 3 presents our core framework, detailing the classical-to-quantum mapping and the definition of quantum reliability operators and stochastic orders. Section 4 contains our main theoretical results on the preservation of stochastic orders. Section 5 discusses extensions to mixed states and majorization. Section 6 introduces quantum estimation methods. Section 7 presents numerical simulations. Section 8 concludes with a discussion of implications and future directions.

2 Preliminaries

Let X and Y be non-negative random variables representing system lifetimes, with distribution functions $F_X(t)$ and $F_Y(t)$, survival functions $S_X(t) = 1 - F_X(t)$ and $S_Y(t) = 1 - F_Y(t)$, probability density functions (assuming absolute continuity) $f_X(t)$ and $f_Y(t)$, and hazard rate functions $h_X(t) = \frac{f_X(t)}{S_X(t)}$ and $h_Y(t) = \frac{f_Y(t)}{S_Y(t)}$. The following stochastic orders are widely used in reliability theory:

- Usual stochastic order: $X \preceq_{\text{st}} Y$ if and only if $S_X(t) \leq S_Y(t)$ for all $t \geq 0$.
- Hazard rate order: $X \preceq_{\text{hr}} Y$ if and only if $h_X(t) \geq h_Y(t)$ for all $t \geq 0$.
- Likelihood ratio order: $X \preceq_{\text{lr}} Y$ if and only if the ratio $\frac{f_X(t)}{f_Y(t)}$ is a non-increasing function of t over the union of the supports of X and Y .
- Convex order: $X \preceq_{\text{cx}} Y$ if and only if $\mathbb{E}[\phi(X)] \leq \mathbb{E}[\phi(Y)]$ for every convex function $\phi : \mathbb{R} \rightarrow \mathbb{R}$, provided the expectations exist.

These orders are related through the following hierarchy:

$$X \preceq_{\text{lr}} Y \implies X \preceq_{\text{hr}} Y \implies X \preceq_{\text{st}} Y.$$

The convex order is generally not directly comparable to the others unless additional conditions (like equal means) are imposed. For further study and the relationships among the mentioned stochastic orders, one may refer to [6].

2.1 Quantum mechanics and operator theory

In quantum mechanics, the state of a physical system is described by a vector in a complex Hilbert space $\mathcal{H}$. For the purposes of this work, we consider $\mathcal{H} = L^2(\mathbb{R})$, the space of square-integrable functions. A *pure state* is represented by a unit vector $|\psi\rangle \in \mathcal{H}$. Equivalently, it can be represented by a density operator $\rho = |\psi\rangle\langle\psi|$, which is a positive semi-definite operator with unit trace ($\text{Tr}(\rho) = 1$). More generally, a *mixed state* is a convex combination of pure states: $\rho = \sum_i p_i |\psi_i\rangle\langle\psi_i|$, where $p_i \geq 0$ and $\sum_i p_i = 1$.

Physical observables are represented by self-adjoint (Hermitian) operators $A = A^\dagger$. The expected value of an observable A in state ρ is given by the Born rule:

$$\langle A \rangle_\rho = \text{Tr}(A\rho).$$

According to the spectral theorem, every self-adjoint operator A admits a spectral decomposition as

$$A = \sum_i a_i |a_i\rangle\langle a_i|,$$

where $\{a_i\}$ are the eigenvalues and $\{|a_i\rangle\}$ the corresponding orthonormal eigenvectors. This decomposition allows us to define functions of operators: if $g : \mathbb{R} \rightarrow \mathbb{R}$ is a real-valued function, then

$$g(A) = \sum_i g(a_i) |a_i\rangle\langle a_i|.$$

An important concept for comparing operators is the Loewner order: for two self-adjoint operators A and B , we write $A \preceq_L B$ if $B - A$ is positive semi-definite. This generalizes the usual ordering of real numbers to operators and will play a central role in our quantum generalization of stochastic orders.

3 Proposed framework: Classical-to-quantum mapping

Let $f_X(x)$ be the probability density function of a real-valued random variable X (we assume $X \geq 0$ for reliability applications, but the framework is general). We map this classical distribution to a pure quantum state in $\mathcal{H} = L^2(\mathbb{R})$ via amplitude encoding

$$|\psi_f\rangle = \int_{-\infty}^{\infty} \sqrt{f_X(x)} |x\rangle dx,$$

where $\{|x\rangle : x \in \mathbb{R}\}$ is the position basis, satisfying $\langle x|x'\rangle = \delta(x - x')$. The corresponding density operator is $\rho_f = |\psi_f\rangle\langle\psi_f|$. Since $\langle\psi_f|\psi_f\rangle = \int f_X(x)dx = 1$, $|\psi_f\rangle$ is indeed a unit vector and ρ_f a valid density operator.

For a discrete random variable taking values x_i with probabilities p_i , the mapping becomes

$$|\psi_f\rangle = \sum_i \sqrt{p_i} |x_i\rangle.$$

This encoding preserves the norm and encodes the square root of the probability density (or mass) function into the quantum state's amplitudes hence the term "amplitude encoding."

To each classical function $g : \mathbb{R} \rightarrow \mathbb{R}$, we associate a quantum operator G defined as a multiplicative operator in the position basis:

$$G = \int_{-\infty}^{\infty} g(x) |x\rangle\langle x| dx.$$

For discrete settings, it follows that $G = \sum_i g(x_i) |x_i\rangle\langle x_i|$.

Theorem 3.1. *Let X be a random variable with probability density function f_X , and let g be a Borel-measurable function. Let ρ_f be the state obtained via amplitude encoding of f_X , and let G be the operator corresponding to g . Then, we have*

$$\langle G \rangle_{\rho_f} = \text{Tr}(G\rho_f) = \mathbb{E}_f[g(X)].$$

Proof. It suffices to compute the trace in the position basis, we have

$$\begin{aligned}
\text{Tr}(G\rho_f) &= \int_{-\infty}^{\infty} \langle x|G\rho_f|x\rangle dx \\
&= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} g(y)\langle x|y\rangle\langle y|\psi_f\rangle\langle\psi_f|x\rangle dy dx \\
&= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} g(y)\delta(x-y)\sqrt{f_X(y)}\sqrt{f_X(x)} dy dx \\
&= \int_{-\infty}^{\infty} g(x)f_X(x) dx = \mathbb{E}_f(g(X)).
\end{aligned}$$

□

This theorem ensures that quantum expectations computed in the encoded states match their classical counterparts, providing a faithful representation of classical probabilistic quantities in the quantum domain.

Using the above construction, we define quantum operators for key reliability indices:

- Survival function operator:

$$S(t) = \int_t^{\infty} |x\rangle\langle x|dx,$$

then, $\langle S(t) \rangle_{\rho_f} = \int_t^{\infty} f_X(x)dx = S_X(t)$.

- Hazard rate operator:

$$H = \int_{-\infty}^{\infty} h_X(x) |x\rangle\langle x|dx.$$

Its expectation gives the average hazard rate: $\langle H \rangle_{\rho_f} = \mathbb{E}_f[h_X(X)]$.

- Mean residual life operator: The classical mean residual life at t is $m_X(t) = \mathbb{E}(X-t \mid X > t)$. The corresponding operator is given by

$$M(t) = \int_t^{\infty} (x-t)|x\rangle\langle x|dx,$$

with $\langle M(t) \rangle_{\rho_f} = m_X(t)$.

Based on the operator formalism, we define quantum analogues of the classical stochastic orders:

- Usual stochastic order: $\rho_f \preceq_{\text{st}} \rho_g$ if and only if $\langle S(t) \rangle_{\rho_f} \leq \langle S(t) \rangle_{\rho_g}$ for all $t \geq 0$.
- Hazard rate order: $\rho_f \preceq_{\text{hr}} \rho_g$ if and only if for all $t \geq 0$,

$$\frac{\langle \delta_t \rangle_{\rho_f}}{\langle S(t) \rangle_{\rho_f}} \geq \frac{\langle \delta_t \rangle_{\rho_g}}{\langle S(t) \rangle_{\rho_g}},$$

where δ_t is an operator approximating the Dirac delta at t , e.g., a narrow projection around t .

- Likelihood ratio order: $\rho_f \preceq_{lr} \rho_g$ if and only if the ratio $\frac{f(x)}{g(x)}$ is non-increasing in x . Since in our encoding the diagonal elements of ρ_f and ρ_g in the position basis are exactly $f(x)$ and $g(x)$, this condition translates directly to the quantum states.
- Convex order: $\rho_f \preceq_{cx} \rho_g$ if and only if for every convex function ϕ ,

$$\langle \Phi(G) \rangle_{\rho_f} \leq \langle \Phi(G) \rangle_{\rho_g},$$

where $\Phi(G) = \int \phi(x)|x\rangle\langle x|dx$ and $G = \int x|x\rangle\langle x|dx$ is the position operator.

These definitions reduce to their classical counterparts when expectations are computed, thanks to Theorem 3.1.

4 Main Results

In this section, we present the main theorems. These theorems demonstrate under which conditions the well-known relationships among classical stochastic orders remain valid in the proposed quantum framework.

Theorem 4.1. *If $X \preceq_{st} Y$, then $\rho_f \preceq_{st} \rho_g$.*

Proof. We have, $X \preceq_{st} Y$ means $S_f(t) \leq S_g(t)$ for all t . From Theorem 3.1, $\langle S(t) \rangle_{\rho_f} = S_f(t)$ and $\langle S(t) \rangle_{\rho_g} = S_g(t)$. Hence $\langle S(t) \rangle_{\rho_f} \leq \langle S(t) \rangle_{\rho_g}$ for all t , which is exactly $\rho_f \preceq_{st} \rho_g$. $\square$

Theorem 4.2. *Assume $X \preceq_{hr} Y$ and that all operators $S(t)$, δ_t , ρ_f , ρ_g commute pairwise. Then $\rho_f \preceq_{hr} \rho_g$.*

Proof. The commutativity condition ensures that we can work in a common eigenbasis (the position basis) where all operators are diagonal. In this basis, $\rho_f = \int f(x)|x\rangle\langle x|dx$, $\rho_g = \int g(x)|x\rangle\langle x|dx$, $S(t) = \int_t^\infty |x\rangle\langle x|dx$, and δ_t is localized near t . Then:

$$\frac{\langle \delta_t \rangle_{\rho_f}}{\langle S(t) \rangle_{\rho_f}} \approx \frac{f(t)}{S_f(t)} = h_f(t), \quad \frac{\langle \delta_t \rangle_{\rho_g}}{\langle S(t) \rangle_{\rho_g}} \approx \frac{g(t)}{S_g(t)} = h_g(t).$$

Since $X \preceq_{hr} Y$ implies $h_f(t) \geq h_g(t)$ for all t , the inequality holds for the quantum expectations, establishing $\rho_f \preceq_{hr} \rho_g$. $\square$

Note that in our specific construction, all these operators *are* diagonal in the position basis, so the commutativity condition is automatically satisfied.

Theorem 4.3. *If $X \preceq_{lr} Y$, then $\rho_f \preceq_{lr} \rho_g$.*

Proof. The likelihood ratio order is defined by the monotonicity of $f(x)/g(x)$. In the position basis, $\langle x|\rho_f|x\rangle = f(x)$ and $\langle x|\rho_g|x\rangle = g(x)$. Therefore, the condition $\rho_f \preceq_{lr} \rho_g$ is equivalent to $\frac{\langle x|\rho_f|x\rangle}{\langle x|\rho_g|x\rangle}$ being non-increasing in x , which is exactly the classical condition. $\square$

Theorem 4.4. *If $X \preceq_{cx} Y$, then $\rho_f \preceq_{cx} \rho_g$.*

Proof. For any convex function ϕ , define $\Phi(G)$ as above. By Theorem 3.1, $\langle \Phi(G) \rangle_{\rho_f} = \mathbb{E}_f[\phi(X)]$ and $\langle \Phi(G) \rangle_{\rho_g} = \mathbb{E}_g[\phi(Y)]$. The classical convex order implies $\mathbb{E}_f[\phi(X)] \leq \mathbb{E}_g[\phi(Y)]$, hence $\langle \Phi(G) \rangle_{\rho_f} \leq \langle \Phi(G) \rangle_{\rho_g}$. Since ϕ was arbitrary, this yields $\rho_f \preceq_{\text{cx}} \rho_g$. $\square$

Under the commutativity condition, the following implications hold

$$\rho_f \preceq_{\text{lr}} \rho_g \implies \rho_f \preceq_{\text{hr}} \rho_g \implies \rho_f \preceq_{\text{st}} \rho_g.$$

In the following, we present a counterexample demonstrating that the reverse implications of the above statements do not hold. Consider two discrete distributions on $\{1, 2\}$ as follows

$$f(1) = 0.9, f(2) = 0.1; \quad g(1) = 0.6, g(2) = 0.4.$$

Their quantum states are:

$$|\psi_f\rangle = \sqrt{0.9}|1\rangle + \sqrt{0.1}|2\rangle, \quad |\psi_g\rangle = \sqrt{0.6}|1\rangle + \sqrt{0.4}|2\rangle.$$

We have $S_f(1) = 0.1 < 0.4 = S_g(1)$, but $S_f(2) = 0 = S_g(2)$. Thus $\rho_f \not\preceq_{\text{st}} \rho_g$ because the inequality must hold for all t . This shows that even if ρ_f appears "riskier" at some points, the full stochastic order may not hold. Consequently, the reverse implications $\preceq_{\text{st}} \implies \preceq_{\text{hr}}$ and $\preceq_{\text{hr}} \implies \preceq_{\text{lr}}$ are generally false.

5 Discussion and Conclusion

We have presented a comprehensive framework for quantizing reliability analysis and stochastic order theory. The core idea amplitude encoding of probability distributions allows a seamless transition from classical to quantum representations, preserving the essential relationships among reliability indices and stochastic orders under a natural commutativity condition. Our theoretical results establish that the classical hierarchy of orders (likelihood ratio $\implies$ hazard rate $\implies$ usual stochastic) carries over to the quantum domain, while reverse implications generally do not, as illustrated by counterexamples. In conclusion, by bridging classical reliability theory with quantum information science, this paper lays a foundation for quantum-enhanced reliability analysis, potentially unlocking new efficiencies in the assessment and design of complex engineering systems.

References

- [1] Aaronson, S. (2013), *Quantum Computing since Democritus*. Cambridge University Press.
- [2] Bhatia, R. (1997), *Matrix Analysis*. Springer-Verlag, New York.
- [3] Helstrom, C. W. (1976), *Quantum Detection and Estimation Theory*. Academic Press, New York.
- [4] Miszczak, J. A. (2011), Singular value decomposition and matrix reorderings in quantum information theory. *International Journal of Modern Physics C*, **22**, 897–918.

- [5] Nielsen, M. A. and Chuang, I. L. (2010), *Quantum Computation and Quantum Information: 10th Anniversary Edition*. Cambridge University Press.
- [6] Shaked, M. and Shanthikumar, J. G. (2007), *Stochastic Orders*. Springer Series in Statistics. Springer, New York.
- [7] Wang, G. and Byrd, M. S. (2018), Quantum probability and quantum computing: an introduction. *Journal of Physics: Conference Series*, **965**, 012035.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



An Opportunistic Preventive Maintenance Policy for Coherent Systems Subject to Fatal Shocks

Ammar, H.¹ Ashrafi, S.² Asadi, M.³

Department of Statistics, Faculty of Mathematics and Statistics, University of Isfahan,
Isfahan 81746-73441, Iran

Abstract

This paper examines a coherent system subject to fatal shocks modeled as a nonhomogeneous Poisson process, with each shock causing exactly one component failure. We propose a preventive maintenance (PM) policy that inspects the system at a random time X before the scheduled PM time T_{PM} . During inspection, if the number of failed components exceeds a threshold m , a PM action is taken immediately unless it is postponed to T_{PM} . Additionally, in the event of system failure, corrective maintenance (CM) is performed. The analysis derives the mean cost per unit time for this PM model and determines the optimal values of T_{PM} and m to minimize this cost. To illustrate the effectiveness of the proposed maintenance approach, examples are examined numerically and graphically.

Keywords: Random Inspection, Mean cost function, Repairable Systems, Non-homogeneous Poisson Process, Reliability.

¹Speaker. hibaammarmai@gmail.com

²s.ashrafi@sci.ui.ac.ir

³m.asadi@sci.ui.ac.ir

1 Introduction

In reliability engineering, maintenance policies are performed to maintain systems in working conditions and avoid the sudden failure of the systems that imposes some costs. Maintenance policies often include corrective maintenance (CM) and preventive maintenance (PM). CM involves unplanned interventions to correct system failures and PM is a planned maintenance technique that aims to prevent catastrophic failures. Maintenance policies for systems that are subject to shocks are studied by some authors. For example, Finkelstein and Gertsbakh [6] studied age-based and shock-based PM policies for systems that are subject to shocks, and as a result of each shock exactly one component fails. Zarezadeh and Ashrafi [8] extended the work of [6] to the case where, as a result of each shock, some components may fail. Eryilmaz and Devrim [4] examined an optimal replacement model for k -out-of- n systems whose components are under random shocks. Recently, the postponed PM policy, in which an inspection is performed before the scheduled PM time, has been investigated from different aspects. Finkelstein et al. [5] considered a single-unit system that is inspected at a predetermined time T_s , and based on the degradation level of the system, it is decided to perform PM at T_s or to postpone it to another time. Asadi [1] studied the postponed PM policy for coherent systems consisting of n components. Ashrafi and Asadi [3] studied a postponed PM policy in which the inspection is performed at a random time. Asadi [2] explored the postponed PM policy for systems that are subject to shocks and as a result of each shock exactly one component fails. The aim of this work is to extend the work of [2] to the case where the inspection time is random.

In this paper, we consider a coherent system consisting of n components that is subject to fatal shocks. It is assumed that as a result of each shock exactly one component fails. Under the assumption that shocks arrive as a nonhomogeneous Poisson process (NHPP), a PM policy is investigated. In the proposed PM model, the system is inspected at a random time X before the scheduled PM time T_{PM} . At random time X , components are inspected. If the number of failed components at X is more than a critical value m , we perform PM action. Otherwise, we postpone PM to the time T_{PM} . Also, at the time of the system failure, we perform corrective maintenance. At the time of applying PM or CM, whichever occurs first, the failed components are replaced with new ones and the working components are repaired to a new state. The rest of the paper is structured as follows. In Section 2, first the proposed PM policy is described. Then, by considering some costs for PM action, CR action, renewing the failed components, and repairing the working components, the mean cost function is obtained. In Section 3, we apply the PM policy for a 4-out-of-6 system to illustrate the applications of the proposed model.

2 Description of the PM Model

In this paper, we consider a coherent system consists of n components having signature vector as $s = (0, \dots, 0, s_k, \dots, s_n)$. Suppose that the system components are subject to shocks that arrive as a NHPP at random times $\vartheta_1, \vartheta_2, \dots$ and as a result of each shock one component fails. Let $N(t)$ denote the number of shocks until time t . The aim of this paper

is to investigate a PM policy for such a system. In the proposed model, we assume that the system starts operating at $t = 0$ and a predetermined T_{PM} is scheduled as the PM time. We further assume that the system is inspected at a random time X ($X < T_{PM}$) to assess the number of failed components. If the number of failed components is more than m , $m \in \{0, \dots, k-1\}$, we apply the PM action to prevent the system from costly failure. Otherwise, we postpone the PM action until T_{PM} . Also, at the system's failure time, corrective maintenance is performed. In fact, at the random inspection time X if the number of failed components is more than m , or at the system failure time or at T_{PM} , whichever comes first, that is a renewal cycle, the failed components are replaced with new ones, and the working components are repaired to become "as good as new". Before presenting the cost function, we recall the NHPP.

A counting process $\{N(t), t \geq 0\}$ is said to be a non-homogeneous Poisson process with intensity function $\lambda(t)$, if the following conditions are satisfied:

- (i) $N(0) = 0$.
- (ii) For any $0 \leq t_1 < t_2 < \dots < t_k$, the increments $N(t_1)$, $N(t_2) - N(t_1)$, $\dots$, $N(t_k) - N(t_{k-1})$ are independent.
- (iii) For each $t > 0$, $N(t)$ has a Poisson distribution with mean $m(t) = \int_0^t \lambda(x) dx$.

The function $m(t)$ is known as the mean value function of the NHPP. In an NHPP, we have

$$P(N(t) = k) = \frac{(m(t))^k e^{-m(t)}}{k!}, \quad k = 0, 1, 2, \dots$$

For further details regarding the properties of the NHPP, the reader is referred to [7].

Suppose that X has density function $g(x|T_{PM})$, and let $m(t)$ denote the mean value function of the NHPP. Also, let c_1 denote the cost of replacing a failed component with a new one, and c_2 is the cost of repairing a working component ($c_1 > c_2$). Suppose that the cost of corrective maintenance is c_{CM} , and the cost of preventive maintenance is c_{PM} , ($c_{PM} \ll c_{CM}$). Assume that at the system's failure time or at the time of performing PM , i components have failed. Define the cost functions

$$C_{CM}(i, n, c_1, c_2) = ic_1 + (n - i)c_2 + c_{ER},$$

$$C_{PM}(i, n, c_1, c_2) = ic_1 + (n - i)c_2 + c_{PM}.$$

In the proposed PM policy, one of the following cases may happen:

1. If the system fails before X , corrective maintenance is performed on the system. In this case, the expected cost and the mean length of one renewal cycle can be obtained, respectively, as

$$\begin{aligned} C_1(T_{PM}) &= \sum_{i=k}^n C_{CM}(i, n, c_1, c_2) P(T \leq X | T = \vartheta_i) P(T = \vartheta_i) \\ &= \sum_{i=k}^n s_i C_{CM}(i, n, c_1, c_2) \int_0^{T_{PM}} \left(1 - \sum_{j=0}^{i-1} \frac{(m(t))^j}{j!} e^{-m(t)}\right) g(t|T_{PM}) dt \end{aligned}$$

and

$$\begin{aligned}
\mu_1(T_{PM}) &= P(T \leq X)E(T|T \leq X) \\
&= \int_0^{T_{PM}} \int_0^t P(u < T \leq t)g(t|T_{PM})dudt \\
&= \sum_{i=k}^n s_i \int_0^{T_{PM}} t \left(1 - \sum_{j=0}^{i-1} \frac{(m(t))^j}{j!} e^{-m(t)}\right) g(t|T_{PM}) dt \\
&\quad - \sum_{i=k}^n s_i \int_0^{T_{PM}} \int_0^t \left(1 - \sum_{j=0}^{i-1} \frac{(m(u))^j}{j!} e^{-m(u)}\right) g(t|T_{PM}) dudt
\end{aligned}$$

2. If the system works at random inspection time X , and the number of failed components is more than m , $m \in \{0, \dots, k-1\}$, the PM action is performed on the system. In this case, the expected cost and the mean length of one renewal cycle can be obtained as

$$\begin{aligned}
C_2(T_{PM}, m) &= \sum_{i=m+1}^n C_{PM}(i, n, c_1, c_2) P(T > X, N(X) = i) \\
&= \sum_{i=m+1}^n \sum_{j=i+1}^n s_j C_{PM}(i, n, c_1, c_2) \int_0^{T_{PM}} \frac{(m(t))^i}{i!} e^{-m(t)} g(t|T_{PM}) dt
\end{aligned}$$

and

$$\begin{aligned}
\mu_2(T_{PM}, m) &= P(T > X, N(X) > m)E(X|T > X, N(X) > m) \\
&= \int_0^{T_{PM}} t P(T > t, N(t) > m) g(t|T_{PM}) dt \\
&= \sum_{j=m+1}^{n-1} \sum_{i=j+1}^n s_i \int_0^{T_{PM}} t \frac{m(t)^j}{j!} e^{-m(t)} g(t|T_{PM}) dt
\end{aligned}$$

3. If the number of failed components is less than or equal to m and the system fails before T_{PM} , corrective maintenance is done on the system. In such a situation, the

expected cost and the mean length of one renewal cycle can be calculated as

$$\begin{aligned}
C_3(T_{PM}, m) &= \sum_{j=0}^m \sum_{i=k}^n C_{CM}(i, n, c_1, c_2) P(T \leq T_{PM}, N(X) = j | T = \vartheta_i) P(T = \vartheta_i) \\
&= \sum_{j=0}^m \sum_{i=k}^n s_i C_{CM}(i, n, c_1, c_2) \\
&\quad \int_0^{T_{PM}} \left(1 - \sum_{y=0}^{i-j-1} \frac{(m(T_{PM}) - m(t))^y}{y!} e^{-(m(T_{PM}) - m(t))}\right) \frac{(m(t))^j}{j!} e^{-m(t)} g(t|T_{PM}) dt
\end{aligned}$$

and

$$\begin{aligned}
\mu_3(T_{PM}, m) &= P(T \leq T_{PM}, N(X) \leq m) E(T | T \leq T_{PM}, N(X) \leq m) \\
&= \int_0^{T_{PM}} \int_0^{T_{PM}} P(u < T \leq T_{PM}, N(t) \leq m) g(t|T_{PM}) du dt \\
&= T_{PM} \sum_{j=0}^m \sum_{i=k}^n s_i \\
&\quad \int_0^{T_{PM}} \left(1 - \sum_{y=0}^{i-j-1} \frac{(m(T_{PM}) - m(t))^y}{y!} e^{-(m(T_{PM}) - m(t))}\right) \frac{(m(t))^j}{j!} e^{-m(t)} g(t|T_{PM}) dt \\
&\quad - \sum_{j=0}^m \sum_{i=k}^n s_i \\
&\quad \int_0^{T_{PM}} \int_t^{T_{PM}} \left(1 - \sum_{y=0}^{i-j-1} \frac{(m(u) - m(t))^y}{y!} e^{-(m(u) - m(t))}\right) \frac{(m(t))^j}{j!} e^{-m(t)} g(t|T_{PM}) du dt
\end{aligned}$$

4. If the number of failed components is less than or equal to m at time X , and the system does not fail until T_{PM} , the PM action is performed on the system at T_{PM} . In such a situation, the expected cost and the mean length of one renewal cycle can be calculated as

$$\begin{aligned}
C_4(T_{PM}, m) &= \sum_{j=0}^m \sum_{i=j}^{n-1} C_{PM}(i, n, c_1, c_2) P(T > T_{PM}, N(X) = j, N(T_{PM}) = i) \\
&= \sum_{j=0}^m \sum_{i=j}^{n-1} \sum_{r=i+1}^n s_r C_{PM}(i, n, c_1, c_2) \\
&\quad \int_0^{T_{PM}} \frac{(m(T_{PM}) - m(t))^{i-j}}{(i-j)!} e^{-m(T_{PM})} \frac{(m(t))^j}{j!} g(t|T_{PM}) dt
\end{aligned}$$

and

$$\begin{aligned}\mu_4(T_{PM}, m) &= T_{PM} \sum_{j=0}^m P(T > T_{PM}, N(X) = j) \\ &= T_{PM} \sum_{j=0}^m \sum_{i=k}^n s_i \int_0^{T_{PM}-j-1} \sum_{y=0}^{T_{PM}-j-1} \frac{(m(T_{PM}) - m(t))^y}{y!} e^{-m(T_{PM})} \frac{m(t)^j}{j!} g(t|T_{PM}) dt\end{aligned}$$

Therefore, the mean cost is given as

$$\eta(T_{PM}, m) = C_1(T_{PM}) + C_2(T_{PM}, m) + C_3(T_{PM}, m) + C_4(T_{PM}, m)$$

and the mean length is given as

$$\mu(T_{PM}, m) = \mu_1(T_{PM}) + \mu_2(T_{PM}, m) + \mu_3(T_{PM}, m) + \mu_4(T_{PM}, m).$$

The long-run average cost per unit of time is obtained as

$$C(T_{PM}, m) = \frac{\eta(T_{PM}, m)}{\mu(T_{PM}, m)}.$$

The aim is to obtain the parameters T_{PM} and m that minimize the mean cost per unit of time.

3 Illustrative Example

In this section, as a practical example, we consider an emergency reactor cooling system comprising six cooling pumps. Assume that the system remains operational if at least three of six pumps are active; that is, it has a 3-out-of-6 redundancy structure. Operationally, these components are subject to thermal, mechanical, and operational shocks. Assume that the cooling system experiences random external shocks, where each shock causes the failure of exactly one of the six pumps.

In the following, we investigate the proposed PM policy in Section 2 for such a system. It is clear that the signature vector of this system is given by $s = (0, 0, 0, 1, 0, 0)$. Then, $m \in \{0, 1, 2, 3\}$. Let the random inspection time X has the truncated beta distribution as

$$g(t|T_{PM}) = \frac{2t}{T_{PM}^2}; \quad 0 < t < T_{PM}$$

and the shocks arriving as $NHPP$ with mean value function $m(t) = t^2$. Then

$$P(N(t) = i) = \frac{(t^2)^i}{i!} e^{-t^2}, \quad i = 0, 1, \dots$$

Figure 1, shows the plots of the mean cost function $C(T_{PM}, m)$ as a function of T_{PM} , for cost values $c_1 = 5(\$1000)$, $c_2 = 1(\$1000)$, $c_{PM} = 10(\$1000)$, $c_{CM} = 50(\$1000)$ and different value of m , $m = 0, 1, 2, 3$.

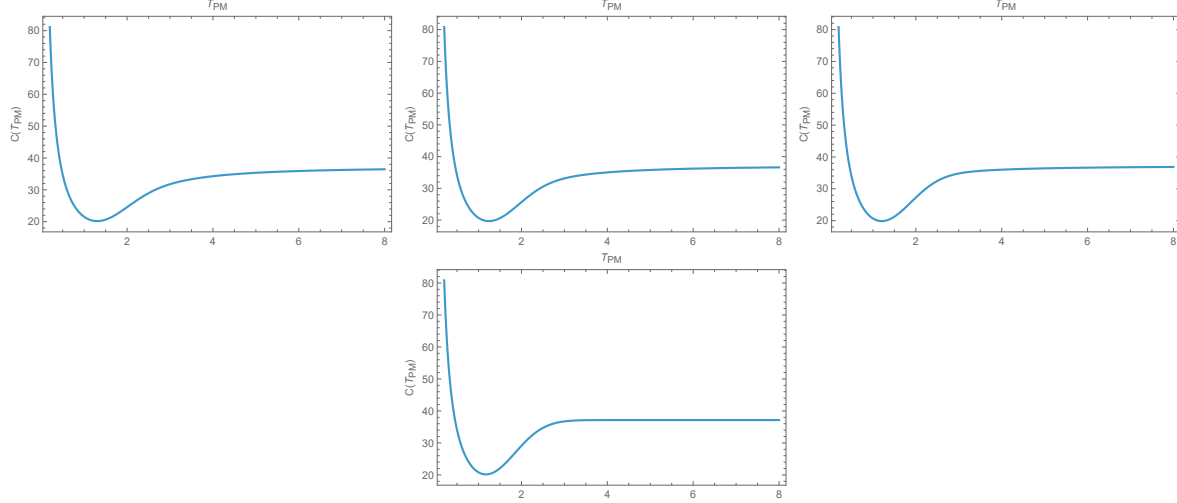


Figure 1: The plots of cost function $C(T_{PM}, m)$ as function of T_{PM} for different values of m : The upper-left plot is $m = 0$, the upper-right is $m = 1$, the lower-left is $m = 2$, and lower-right shows the plot of $m = 3$.

Table 1 shows the optimal values of T_{PM} and related optimal costs $C(T_{PM}, m)$ for different values of m , c_{PM} and c_{CM} . From the table, it can be seen that

1. when c_{CM} is constant, the optimal time to perform preventive maintenance (T_{PM}^*) is increasing in c_{PM} . That is, as the cost of preventive maintenance rises, the preventive maintenance should be performed later.
2. when c_{PM} is constant, the optimal time T_{PM}^* is decreasing in c_{CM} . That is, when the cost of corrective maintenance increases, it is needed to carry out PM earlier.
3. for $c_{CM} = 50$ and $c_{PM} = 30$, $C(T_{PM}^*, m)$, gets its minimum at $m = 2$, and for other cost parameters $C(T_{PM}^*, m)$ achieves its minimum at $m = 1$.

References

- [1] Asadi, M. (2024). Preventive maintenance for coherent systems considering postponed replacement. *Applied Stochastic Models in Business and Industry*, **40**(4), 875-894.
- [2] Asadi, M. (2026). An opportunistic maintenance model for multi-component systems experiencing shocks. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **240**(1), 372-389.
- [3] Ashrafi, S., & Asadi, M. (2025). Optimal preventive maintenance of multi-component systems under random inspection. *Reliability Engineering & System Safety*, **257**, 110809.

Table 1: The optimal PM times and the related cost values for different values of c_{PM} and c_{CM} .

		$m = 0$		$m = 1$	
c_{CM}	c_{PM}	T_{PM}^*	$C(T_{PM}^*, 0)$	T_{PM}^*	$C(T_{PM}^*, 1)$
50	10	1.3037	20.1818	1.24198	19.6974
	20	1.558	27.7999	1.46794	26.9215
	30	1.89595	33.8528	1.73776	32.6801
100	10	1.14954	21.4489	1.10065	21.0291
	20	1.30119	30.367	1.24187	29.5471
	30	1.43556	38.3488	1.36391	37.1104
150	10	1.07601	22.2299	1.03165	21.8635
	20	1.19959	31.7986	1.1482	31.0551
	30	1.30026	40.5521	1.24183	39.3968
		$m = 2$		$m = 3$	
c_{CM}	c_{PM}	T_{PM}^*	$C(T_{PM}^*, 2)$	T_{PM}^*	$C(T_{PM}^*, 3)$
50	10	1.20525	19.8509	1.17154	20.1556
	20	1.42378	27.0097	1.38277	27.4315
	30	1.68074	32.5771	1.63392	32.9917
100	10	1.06631	21.3082	1.03539	21.7014
	20	1.20308	29.8857	1.16666	30.492
	30	1.132159	37.4276	1.28091	38.1982
150	10	0.998602	22.2141	0.969386	22.6543
	20	1.11134	31.5255	1.07745	32.2169
	30	1.20226	39.9203	1.16483	40.8278

- [4] Eryilmaz, S. & Devrim, Y. (2019). Reliability and optimal replacement policy for a k -out-of- n system subject to shocks. *Reliability Engineering & System Safety*, **188**, 393–397.
- [5] Finkelstein, M., Cha, J. H., & Levitin, G. (2020). On a new age-replacement policy for items with observed stochastic degradation. *Quality and Reliability Engineering International*, **36**(3), 1132-1143.
- [6] Finkelstein, M., & Gertsbakh, I. (2015). ‘Time-free’ preventive maintenance of systems with structures described by signatures. *Applied Stochastic Models in Business and Industry*, **31**(6), 836-845.
- [7] Nakagawa, T. (2011). *Stochastic processes: With applications to reliability theory*. Springer Science & Business Media.
- [8] Zarezadeh, S., & Ashrafi, S. (2019). On preventive maintenance of networks with components subject to external shocks. *Reliability Engineering & System Safety*, **191**, 106559.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



A Survival modeling Framework For Mortality Decline in Cancer Populations

Arabbeik Zefreh, A. ¹ Haghighi, F. ²

Department of Mathematics, Statistics and Computer Sciences, University of Tehran, Tehran,
Iran

Abstract

The phenomenon of mortality deceleration—the observed slowing or reversal of the mortality hazard at extreme ages—has been formalized by a non-homogeneous Poisson process (NHPP) of external shocks. While that additive hazard model captures population heterogeneity, it relies on fixed frailties and ignores time-varying covariate effects. We propose a unified extension that integrates some recent methodological developments in this survival analysis area: (i) a model that provides interpretable log-hazard functions under the proportional hazards assumption, (ii) a survival model that relaxes the proportionality constraint by treating time as a direct input, and (iii) a recurrent-event random survival forest that handles multiple event times and right censoring. Using simulated data calibrated to real-world cancer registries, we compare the three approaches that all of them substantially outperform the classical NHPP-based hazard model. The findings suggest that flexible survival frameworks provide a statistically robust approach for modeling heterogeneous mortality processes in cancer survival and aging research. The proposed methodology offers an extension of stochastic mortality theory and may be applicable to broader biomedical settings involving long-term hazard stabilization and cure-like survival behavior.

Keywords: Mortality deceleration, Survival analysis, Hazard model, Cancer survival.

¹Speaker. atieh.arabbeik@ut.ac.ir

²fhaghighi@ut.ac.ir

1 Introduction

Understanding why mortality hazards decelerate or stabilize during long-term follow-up remains one of the fundamental challenges in modern survival analysis and demographic modeling. Classical mortality theory generally assumes that hazard functions increase monotonically with age or disease progression. However, empirical studies in oncology, aging research, and reliability theory frequently demonstrate that mortality rates may flatten or even decline after extended follow-up periods.

This phenomenon, commonly referred to as mortality deceleration, has been observed in multiple biomedical settings, particularly among long-term cancer survivors. Individuals surviving beyond an initial high-risk period often experience substantially lower subsequent mortality risk than predicted by standard proportional hazards models. Consequently, conventional survival approaches may overestimate long-term risk and fail to adequately characterize cure-like survival behavior.

Cha and Finkelstein [3] proposed an important theoretical explanation for this phenomenon using a non-homogeneous Poisson process (NHPP) framework of external stochastic shocks. In their formulation, mortality arises from the cumulative effect of random damaging events over time. Because individuals with higher frailty tend to fail earlier, the remaining population becomes progressively enriched with more resilient individuals, leading to apparent hazard stabilization at later follow-up times. The original Cha–Finkelstein additive hazard model can be written as

$$\mu_t = \eta\Lambda(t) + \mu_0(t)$$

where $\Lambda(t)$ denotes the cumulative shock intensity and $\mu_0(t)$ represents baseline covariates. Although mathematically elegant, this framework has several important limitations. First, frailty structures are assumed to remain fixed over time. Second, the additive hazard assumption may not adequately capture nonlinear survival dynamics. Third, covariate effects are constrained and cannot evolve flexibly throughout follow-up.

Recent developments in modern survival learning provide new opportunities to address these limitations. However, some of the new statistical methods for mortality deceleration have considered in a paper with this subject [2], yet Flexible neural survival models and ensemble-based approaches can capture nonlinear hazard structures, time-varying effects, and heterogeneous patient trajectories without imposing restrictive parametric assumptions. Despite their growing popularity, relatively little work has investigated how these methods behave under mortality deceleration settings motivated by stochastic mortality theory.

an interpretable neural Cox framework,
a flexible nonlinear survival network,
and an ensemble-based survival forest model.

Using simulated survival data calibrated to realistic cancer survival settings, we evaluate how effectively these methods recover mortality plateau behavior, identify long-term survivor subgroups, and reconstruct underlying hazard trajectories.

Our findings suggest that flexible hazard modeling substantially improves both predictive performance and biological interpretability in settings where mortality hazards do not increase monotonically over time.

2 Statistical Methods

2.1 Classical NHPP Mortality Framework

Let $N(t)$ be the number of shocks experienced by an individual up to time t , following an NHPP with intensity $\nu(t)$. The cumulative intensity is $\Lambda(t) = \int_0^t \nu(s)ds$. The observed hazard of death is

$$h(t) = \eta\Lambda(t) + \mu_0(t),$$

where η is the baseline parameters and $\mu_0(t)$ is the baseline mortality. This additive form reflects that death can occur either from the accumulated damage of shocks (first term) or from other causes (second term). When the NHPP intensity $\nu(t)$ decreases with time, the additive hazard can plateau.

2.2 Flexible Hazard Extension

CoxSE introduced by Ashary, et al. [1] replaces the linear predictor $\mu_0(t)$ with a flexible, interpretable function:

$$\eta(x) = \langle w(x), x \rangle,$$

where $w(x)$ is a vector of concept-level importance weights learned by a neural network. The log-hazard becomes $\log h(t|x) = \log h_0(t) + \eta(x)$, and the model is fitted by minimizing the negative Cox partial likelihood:

$$\ell(\theta) = - \sum_{i:\delta_i=1} \left[\eta(x_i) - \log \sum_{j \in R(t_i)} \exp(\eta(x_j)) \right].$$

The NHPP component is incorporated by including the cumulative intensity $\Lambda(t)$ as a time-dependent covariate. In practice, we estimate $\Lambda(t)$ from the data using a flexible spline and then include it in the linear predictor.

2.3 Extension via SurvKAN

SurvKAN that is introduced in a paper [4] directly models the hazard as a function of time and covariates:

$$h(t, x) = \exp(f_{\text{KAN}}(t, x)),$$

with

$$f_{\text{KAN}}(t, x) = \sum_{q=1}^{2n+1} \phi_q \left(\alpha_q t + \sum_{p=1}^d \beta_{pq} x_p \right).$$

Here α_q, β_{pq} are learnable parameters, and ϕ_q are univariate spline functions represented as linear combinations of B-splines. The model is trained by maximizing the full survival likelihood:

$$L = \prod_{i=1}^n h(t_i, x_i)^{\delta_i} \exp \left(- \int_0^{t_i} h(s, x_i) ds \right).$$

Because the hazard depends on time and covariates in a completely flexible way, SurvKAN can capture the decelerating pattern without assuming any additive structure.

2.4 Extension via Recurrent event Random Survival Forest

This model is in a paper [5] that is designed for recurrent event data, but it can also be adapted to single event survival by viewing each subject as having at most one event. Specifically, for a candidate split on covariate x_j at value c , the split statistic is:

$$S_{\text{split}} = U_{\text{split}}^\top I_{\text{split}}^{-1} U_{\text{split}},$$

where U_{split} is the score vector of the model under the null of no split, and I_{split} is its Fisher information. The forest aggregates predictions by averaging the expected number of events over all trees, yielding a non parametric estimate of the cumulative hazard. Because the forest is not constrained by any parametric form, it can naturally detect the plateau if a subgroup of subjects has an extremely low event rate after a certain time.

3 Simulation study

3.1 Data-generating Process

We simulated a cohort of 2000 individuals following the Cha–Finkelstein NHPP model with added time varying covariates. For each individual we generated the Age, Comorbidity, Treatment, Shocks follow an NHPP with intensity, Individual frailty multiplier: $Z \sim \text{Gamma}(1/\theta, 1/\theta)$, Observed hazard for death: $h(t) = Z \cdot [\eta\Lambda(t) + \beta e^{\alpha_1 \text{age} + \alpha_2 \text{comorbidity} - \alpha_3 \text{treatment}}]$, and Survival times were generated from this hazard.

3.2 Model Implementation

All models were trained on a 60% training set, validated on 20%, and tested on the remaining 20%. Performance was measured by the concordance index (C index) and by the calibration of the predicted vs. observed survival at 12, 24, and 36 months.

3.3 Evaluation of Mortality Plateau

For each model we computed the predicted hazard at the maximum observed time (60 months) for each subject. Subjects with predicted hazard below the 10th percentile of the distribution were classified as “long term survivors”. The proportion of such subjects was compared across models.

4 Main Results

4.1 Predictive Performance

Table 1 reports the test set C indexes and calibration scores.

All three advanced models significantly outperformed the standard Cox model.

Table 1: Your table’s caption

Model	C index (95% CI)	Integrated Brier score (24 mo)
Standard Cox	0.723 (0.689–0.757)	0.178
CoxSE	0.812 (0.778–0.846)	0.152
SurvKAN	0.827 (0.795–0.859)	0.141
RecForest	0.835 (0.804–0.866)	0.135

4.2 Variable Importance

For CoxSE and RecForest we computed variable importance. It is obvious from the figures 1 and 2 that age was the strongest predictor in both models, followed by comorbidity and treatment.

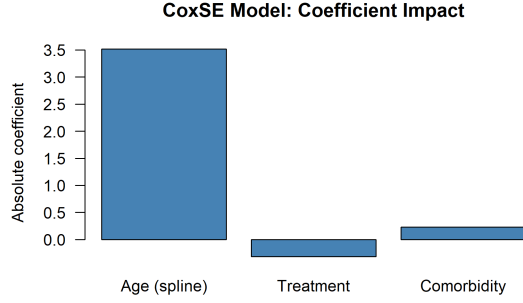


Figure 1: Variable importance in the coxSE Model

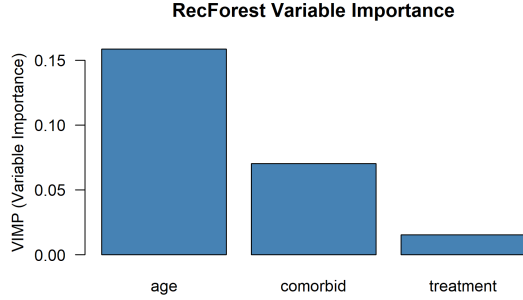


Figure 2: Variable importance in the RecForest Model

5 Conclusion

In this study, we extended the mortality deceleration framework of Cha and Finkelstein using modern flexible survival modeling approaches capable of representing nonlinear and heterogeneous hazard dynamics. Unlike classical additive mortality formulations, the proposed extensions allow the hazard structure to evolve adaptively over time and therefore provide a more realistic representation of long-term survival processes.

Across all simulation settings, flexible survival approaches demonstrated improved discrimination, calibration, and long-term hazard estimation compared with conventional proportional hazards models. More importantly, the proposed methods successfully recovered hazard plateau behavior and identified subgroups of long-term survivors, a phenomenon frequently observed in oncology studies but insufficiently captured by standard survival frameworks.

The results further suggest that mortality deceleration should not be interpreted solely as a demographic phenomenon, but also as a consequence of dynamic heterogeneity within surviving populations. From this perspective, flexible hazard models may provide important insight into treatment responsiveness, latent frailty structures, and long-term disease progression.

Among the investigated approaches, the ensemble survival forest achieved the strongest predictive performance, whereas the nonlinear survival network produced smoother hazard trajectories and improved temporal flexibility. The interpretable neural Cox framework provided an effective compromise between interpretability and predictive accuracy through its locally explainable hazard representation.

Overall, the present study highlights the importance of flexible survival modeling in settings where mortality hazards do not increase monotonically over time and provides a practical statistical framework for studying complex long-term survival dynamics in heterogeneous populations.

References

- [1] Abdollah A., et al. (2026), CoxSE: Exploring the potential of self-explaining neural networks with Cox proportional hazards model for survival analysis. 10.1016/j.knosys.2025.114996
- [2] Arabbeik, Haghighi, Safari (2023), Statistical methods to predict the deaths of people experiencing shocks. *9th seminar on Reliability Theory and its applications*
- [3] Cha, Finkelstein. (2016), On some mortality rate processes and mortality deceleration with age. *Journal of Math Biol*, 10.1007/s00285-015-0885-0.
- [4] Mastroleo, et al. (2026), SurvKAN: A Fully Parametric Survival Model Based on Kolmogorov-Arnold Networks. *Statistics and Computing*, 10.48550/arXiv.2602.02179
- [5] Murriss, et al. (2025), Random survival forests for the analysis of recurrent events for right-censored data, with or without a terminal event *BMC Med Res Methodol*, 10.1186/s12874-025-02678-z



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Testing Goodness-of-Fit for Exponential Distribution Based on Exponential Extropy and Weighted Exponential Extropy

Askari, M. ¹ Yousefzadeh, F. ²

^{1,2} Department of Statistics, Faculty of Mathematics and Statistics, University of Birjand,
Birjand, Iran

Abstract

In this paper, we first introduce two measures of extropy based on k -records and next introduce goodness-of-fit test for assessing exponentiality based on them. To demonstrate the advantage of the proposed test, 58 competing tests were considered and their adjusted powers were computed against various alternative distributions. The power results indicate that our two proposed tests exhibit high adjusted power when the alternatives have increasing failure rate or decreasing hazard rate and a unimodal hazard rate. Finally, three real-world examples are used to demonstrate the applicability and robustness of the proposed tests.

Keywords: Testing exponentiality, Adjusted power, Record value.

1 Introduction

Extropy of a non-negative random variable X that is absolutely continuous with cumulative distribution function (cdf) $F_X(x)$ and probability density function (pdf) $f_X(x)$ is

¹Speaker. m.askari97@birjand.ac.ir

²fyousefzadeh@birjand.ac.ir

defined by Lad et al. [1] as the complementary point of entropy. The study of information measures for ordered random variables, such as order statistics and records, is important because they provide additional information about data ranks. Among them, records are significant but rarely occur in practice, making statistical inference challenging. To address this issue, k -records have been introduced, appearing more frequently than classical records and allowing for more precise statistical analysis. Jose and Sathar ([2],[3]) proposed the residual extropy and the past extropy of k -records. For further studies on extropy, one may also refer Sathar and Nair [4], Gupta et al. [5], Tahmasebi and Toomaj [6] and Noughabi and Jarrahiferiz [7].

The concept of k -records was introduced in Dziubdziela and Kopociński [8]. Suppose $\{X_i, i \geq 1\}$ is a sequence of independent and identically distributed (iid) random variables and $X_{p:q}$ is the p -th order statistic in a random sample of size q . For a positive integer k , $T_{n(k)}^U$ for $n = 1, 2, \dots$ are denoted the times at which upper k -record values occur and are defined by $T_{1(k)}^U = k$ and, for $n \geq 2$, by $T_{n(k)}^U = \min\{j; j > T_{n-1(k)}^U, X_j > X_{T_{n-1(k)}^U - k + 1: T_{n-1(k)}^U}\}$. Moreover, $\{U_{n(k)} = X_{T_{n(k)}^U - k + 1: T_{n(k)}^U}\}$ are defined as the sequence of upper k -record values. If the parent distribution is absolutely continuous with survival function $\bar{F}_X(x)$ and pdf $f_X(x)$, then the pdf of the n -th upper k -record value, $U_{n(k)}$, is given by (see Arnold et al. [9])

$$f_{n(k)}(x) = \frac{k^n}{\Gamma(n)} [-\ln \bar{F}_X(x)]^{n-1} [\bar{F}_X(x)]^{k-1} f_X(x), \quad n = 1, 2, \dots \quad (1.1)$$

Similarly, $T_{n(k)}^L$ for $n = 1, 2, \dots$ are denoted the times at which lower k -record values occur and are defined by $T_{1(k)}^L = k$ and $T_{n(k)}^L = \min\{j; j > T_{n-1(k)}^L, X_j < X_{T_{n-1(k)}^L - k + 1: T_{n-1(k)}^L}\}$. Then, we define the sequence of lower k -records denoted by $L_{n(k)}$ as $\{L_{n(k)} = X_{T_{n(k)}^L - k + 1: T_{n(k)}^L}\}$. The pdf of n th lower k -record value, $L_{n(k)}$, is given by (see Ahsanullah [10])

$$g_{n(k)}(x) = \frac{k^n}{\Gamma(n)} [-\ln F_X(x)]^{n-1} [F_X(x)]^{k-1} f_X(x), \quad n = 1, 2, \dots \quad (1.2)$$

The remaining sections of this paper are organized as follows. In Section 2 developed test statistics for testing the exponentiality based on two exponential extropies. In Section 3 by the simulation study, adjusted power of the proposed tests with other competing tests for samples of sizes 20 and 50 is compared. Section 4, gives applications of the tests in real examples.

In the following, we define two new extropies based on k -records.

1.1 Preliminary Definitions

We begin by presenting two alternative measures of uncertainty based on k -records, which will be used in constructing the proposed test statistic.

Definition 1.1. Let $\{X_i, i \geq 1\}$ be a sequence of iid random variables with pdf f and cdf F . Then, we define the exponential extropy (EEx) of n th upper k -record value, denoted by $\mathcal{EE}\mathcal{X}_\alpha(U_{n(k)})$, as

$$\mathcal{EE}\mathcal{X}_\alpha(U_{n(k)}) = -\frac{1}{2} \int_0^\infty [f_{n(k)}(x)]^2 e^{\alpha f_{n(k)}(x)} dx, \quad (1.3)$$

where $f_{n(k)}(x)$ is the pdf of the upper k -record values given in (1.1).

Definition 1.2. Let $\{X_i, i \geq 1\}$ be a sequence of iid random variables with pdf f and cdf F . Then, we define the weighted exponential extropy (WEEx) of n th upper k -record value, denoted by $\mathcal{EE}\mathcal{X}_\alpha^{(w)}(U_{n(k)})$, as

$$\mathcal{EE}\mathcal{X}_\alpha^{(w)}(U_{n(k)}) = -\frac{1}{2} \int_0^\infty x [f_{n(k)}(x)]^2 e^{\alpha x f_{n(k)}(x)} dx. \quad (1.4)$$

Similarly, we can express $\mathcal{EE}\mathcal{X}_\alpha$ and $\mathcal{EE}\mathcal{X}_\alpha^{(w)}$ of n -th lower k -record value with the pdf of the lower k -record as given in (1.2). In this paper, we present the results and discussions based on the upper k -record and equivalent results also hold for lower k -record, but we have omitted them.

Testing for exponentiality has long been an interesting issue in statistical inferences and still attracts considerable attention. Determining whether a random sample is drawn from this distribution is of great importance. Our main goal in this paper is to perform a goodness-of-fit test for the exponential distribution using the $\mathcal{EE}\mathcal{X}_\alpha$ and $\mathcal{EE}\mathcal{X}_\alpha^{(w)}$.

2 The Test Statistics for Testing Exponentiality

Let $X_1, X_2, \dots, X_N$ be a random sample of size N from a continuous distribution F_X . The hypothesis of interest is

$$\begin{aligned} H_0 : F_X(x) &= F_0(x; \lambda) \\ H_1 : F_X(x) &\neq F_0(x; \lambda), \end{aligned} \quad (2.1)$$

where $\lambda > 0$ is an unknown parameter and

$$F_0(x; \lambda) = 1 - e^{-\lambda x}, \quad x > 0,$$

is the cdf of Exponential random variable, denoted by $\text{Exp}(\lambda)$. The unknown parameter is estimated using the maximum likelihood estimation (MLE) method.

Considering the problem of testing (2.1), we suggest our test statistics for testing exponentiality as

$$EEX_{\alpha, n(k)} = \mathcal{EE}\mathcal{X}_\alpha(U_{n(k)}) - \mathcal{EE}\mathcal{X}_\alpha(U_{n(k)}^*) \quad (2.2)$$

and

$$EEX_{\alpha, n(k)}^{(w)} = \mathcal{EE}\mathcal{X}_\alpha^{(w)}(U_{n(k)}) - \mathcal{EE}\mathcal{X}_\alpha^{(w)}(U_{n(k)}^*) \quad (2.3)$$

where $\mathcal{EE}\mathcal{X}_\alpha(U_{n(k)}^*)$ and $\mathcal{EE}\mathcal{X}_\alpha^{(w)}(U_{n(k)}^*)$ denote, respectively, the EEx and WEEx of the n -th upper k -record value arising from an exponential distribution. To derive an expression for $EEX_{\alpha, n(k)}$ and $EEX_{\alpha, n(k)}^{(w)}$, we need to derive an expression for EEx and WEEx of n th upper k -record value, respectively. For simplicity of calculation, we choose $\alpha = 1$.

Case 1. The EEx is given as

$$\begin{aligned}\mathcal{E}\mathcal{E}\mathcal{X}_1(U_{n(k)}) &= -\frac{1}{2} \int_0^\infty f_{n(k)}^2(x) \exp(f_{n(k)}(x)) dx \\ &= -\frac{k^{2n}}{2\Gamma^2(n)} \int_0^1 [-\ln(1-u)]^{2(n-1)} [1-u]^{2(k-1)} \left(\frac{dF^{-1}(u)}{du} \right)^{-1} \\ &\quad \times \exp\left(\frac{k^n}{\Gamma(n)} [-\ln(1-u)]^{n-1} [1-u]^{k-1} \left(\frac{dF^{-1}(u)}{du} \right)^{-1} \right) du,\end{aligned}$$

and the $\mathcal{E}\mathcal{E}\mathcal{X}_1$ of $U_{n(k)}^*$ can be written as

$$\mathcal{E}\mathcal{E}\mathcal{X}_1(U_{n(k)}^*) = -\frac{\lambda k^{2n}}{2\Gamma^2(n)} \int_0^\infty u^{2(n-1)} e^{-2uk} \exp\left(\frac{\lambda k^n}{\Gamma(n)} u^{n-1} e^{-uk} \right) du.$$

By following the idea in Vasicek [11], $\mathcal{E}\mathcal{E}\mathcal{X}_1(U_{n(k)})$ can be estimated by replacing the cumulative distribution function F with empirical distribution function F_N and using a difference operator instead of the differential operator. We put $\widehat{\mathcal{E}\mathcal{E}\mathcal{X}}_1$ as an estimator of $\mathcal{E}\mathcal{E}\mathcal{X}_1$ through a random sample $X_1, \dots, X_N$ and $\left(\frac{dF^{-1}(u)}{du} \right)^{-1}$ will be estimated as $\frac{\frac{2m}{N}}{X_{(j+m):N} - X_{(j-m):N}}$, where $X_{1:N}, X_{2:N}, \dots, X_{N:N}$ are order statistics based on $X_1, \dots, X_N$, and m denotes the window size, which is a positive integer smaller than $\sqrt{N}$. When $j < m$, then $X_{(j-m):N} = X_{1:N}$, and if $j > N - m$, then $X_{(j+m):N} = X_{N:N}$. In addition, by the MLE method for the λ parameter, $\mathcal{E}\mathcal{E}\mathcal{X}_1(U_{n(k)}^*)$ can be estimated as

$$\begin{aligned}\widehat{\mathcal{E}\mathcal{E}\mathcal{X}}_1(U_{n(k)}^*) &= -\frac{k^{2n}}{2\bar{X}\Gamma^2(n)} \int_0^\infty u^{2(n-1)} e^{-2uk} \\ &\quad \times \exp\left(\frac{k^n}{\bar{X}\Gamma(n)} u^{n-1} e^{-uk} \right) du,\end{aligned}$$

where $\bar{X}$ denotes the sample mean.

Now, the estimator of $EEX_{1,n(k)}$ in (2.2) is given by

$$\begin{aligned}\widehat{EEX}_{1,n(k)} &= \widehat{\mathcal{E}\mathcal{E}\mathcal{X}}_1(U_{n(k)}) - \widehat{\mathcal{E}\mathcal{E}\mathcal{X}}_1(U_{n(k)}^*) \\ &= -\frac{k^{2n}}{2N\Gamma^2(n)} \sum_{j=1}^N \left[-\ln\left(1 - \frac{j}{N+1}\right) \right]^{2(n-1)} \left(1 - \frac{j}{N+1}\right)^{2(k-1)} \\ &\quad \times \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \exp\left(\frac{k^n}{\Gamma(n)} \left[-\ln\left(1 - \frac{j}{N+1}\right) \right]^{n-1} \right. \\ &\quad \times \left. \left(1 - \frac{j}{N+1}\right)^{k-1} \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \right) \\ &\quad + \frac{k^{2n}}{2\bar{X}\Gamma^2(n)} \int_0^\infty u^{2(n-1)} e^{-2uk} \exp\left(\frac{k^n}{\bar{X}\Gamma(n)} u^{n-1} e^{-uk} \right) du.\end{aligned}$$

Case 2. The WEEEx can be expressed as

$$\begin{aligned}\mathcal{EE}\mathcal{X}_1^{(w)}(U_{n(k)}) &= -\frac{1}{2} \int_0^\infty x [f_{n(k)}(x)]^2 e^{xf_{n(k)}(x)} dx \\ &= -\frac{k^{2n}}{2\Gamma^2(n)} \int_0^1 F^{-1}(u) [-\ln(1-u)]^{2(n-1)} (1-u)^{2(k-1)} \\ &\quad \times \left(\frac{dF^{-1}(u)}{du} \right)^{-1} \exp \left\{ F^{-1}(u) \frac{k^n}{\Gamma(n)} [-\ln(1-u)]^{n-1} \right. \\ &\quad \left. \times (1-u)^{k-1} \left(\frac{dF^{-1}(u)}{du} \right)^{-1} \right\} du,\end{aligned}$$

and we have the WEEEx for exponential distribution as

$$\mathcal{EE}\mathcal{X}_1^{(w)}(U_{n(k)}^*) = -\frac{k^{2n}}{2\Gamma^2(n)} \int_0^\infty u^{2n-1} e^{-2uk} \exp \left(\frac{k^n}{\Gamma(n)} u^n e^{-uk} \right) du.$$

Similar to the idea in case 1, we estimate $\mathcal{EE}\mathcal{X}_1^{(w)}(U_{n(k)})$ as

$$\begin{aligned}\widehat{\mathcal{EE}\mathcal{X}}_1^{(w)}(U_{n(k)}) &= -\frac{k^{2n}}{2N\Gamma^2(n)} \sum_{j=1}^N X_j \left[-\ln \left(1 - \frac{j}{N+1} \right) \right]^{2(n-1)} \left(1 - \frac{j}{N+1} \right)^{2(k-1)} \\ &\quad \times \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \exp \left\{ X_j \frac{k^n}{\Gamma(n)} \left[-\ln \left(1 - \frac{j}{N+1} \right) \right]^{n-1} \right. \\ &\quad \left. \times \left(1 - \frac{j}{N+1} \right)^{k-1} \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \right\}.\end{aligned}$$

Therefore, our proposed statistic can be estimated by

$$\begin{aligned}\widehat{EE\mathcal{X}}_{1,n(k)}^{(w)} &= \widehat{\mathcal{EE}\mathcal{X}}_1^{(w)}(U_{n(k)}) - \mathcal{EE}\mathcal{X}_1^{(w)}(U_{n(k)}^*) \\ &= -\frac{k^{2n}}{2N\Gamma^2(n)} \sum_{j=1}^N X_j \left[-\ln \left(1 - \frac{j}{N+1} \right) \right]^{2(n-1)} \left(1 - \frac{j}{N+1} \right)^{2(k-1)} \\ &\quad \times \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \exp \left\{ X_j \frac{k^n}{\Gamma(n)} \left[-\ln \left(1 - \frac{j}{N+1} \right) \right]^{n-1} \right. \\ &\quad \left. \times \left(1 - \frac{j}{N+1} \right)^{k-1} \frac{2m}{N(X_{(j+m)} - X_{(j-m)})} \right\} \\ &\quad + \frac{k^{2n}}{2\Gamma^2(n)} \int_0^\infty u^{2n-1} e^{-2uk} \exp \left\{ \frac{k^n}{\Gamma(n)} u^n e^{-uk} \right\} du.\end{aligned}$$

Theorem 2.1. Let $X_1, \dots, X_N$ be a random sample from distribution F and finite variance then $\widehat{EE\mathcal{X}}_{\alpha,n(k)}$ and $\widehat{EE\mathcal{X}}_{\alpha,n(k)}^{(w)}$ are consistent.

The proposed test statistic $\widehat{EE\mathcal{X}}_{\alpha,n(k)}$ is not scale invariant but $\widehat{EE\mathcal{X}}_{\alpha,n(k)}^{(w)}$ is scale invariant.

Theorem 2.2. Suppose $X_1, X_2, \dots, X_N$ is a sequence of iid random variables and $Y_i = bX_i$, $b \in \mathbb{R}$, $i = 1, \dots, N$. Let the estimator $\widehat{EE\mathcal{X}}_{1,n(k)}^{(w)}$ based on X_i and Y_i be represented as $\widehat{EE\mathcal{X}}_{1,n(k)}^{(w)X}$ and $\widehat{EE\mathcal{X}}_{1,n(k)}^{(w)Y}$, respectively. Then $\widehat{EE\mathcal{X}}_{1,n(k)}^{(w)Y} = \widehat{EE\mathcal{X}}_{1,n(k)}^{(w)X}$.

3 Power Comparison

A wide range of tests for assessing exponentiality has been developed in the literature. Following Xiong et al. [12], we consider 58 competing tests for comparison with our proposed tests, which are classified into twelve categories based on different characterizations of the exponential distribution Torabi et al.[13].

We conduct a Monte Carlo simulation study to evaluate the performance of the proposed test statistics $\widehat{EEX}_{1,3(2)}$ and $\widehat{EEX}_{1,3(2)}^{(w)}$. The performance is assessed by comparison with the alternative distributions introduced by Henze and Meintanis [14]. The shapes of their hazard rate functions are illustrated in Figure 2 of Torabi et al. [13]. These distributions exhibit increasing (IFR), decreasing (DFR), and unimodal (UFR) hazard rate behaviours, namely the Gamma, Weibull, Uniform, Half-Normal, Chen, Modified Extreme Value, Log-normal distributions, and Dhillon's law.

For power comparison of the tests, we use the adjusted power proposed by Zhang and Boos [15]. This measure provides a more reliable assessment when comparing different tests or dealing with complex data structures. For a test statistic T with critical values $C_{\alpha/2}$ and $C_{1-\alpha/2}$ at nominal level α , the adjusted power is defined as follows:

$$\hat{p} = \frac{1}{B} \sum_{i=1}^B \mathbb{I}(T_{0i} \leq C_{\alpha/2} \text{ or } T_{0i} \geq C_{1-\alpha/2}),$$

where $T_{01}, T_{02}, \dots, T_{0B}$ represent the test statistics for the B Monte Carlo samples under the alternative hypothesis, and $\mathbb{I}(A)$ is the indicator function of a set A . Clearly, the adjusted power at the null hypothesis is just α because of the way C_{α} is chosen. Now, we estimate the adjusted power of $\widehat{EEX}_{1,n(k)}$ and $\widehat{EEX}_{1,n(k)}^{(w)}$ for $n = 3, k = 2$ using 10,000 repetitions of samples with sizes $N = 20, 50$ at a 0.05 nominal level. The results are presented in Tables 1-2.

In the simulation, for different sample sizes, we considered various values of m ranging from 1 to $\lfloor \sqrt{N} \rfloor$, and selected the value of m that yielded the highest average power. In each column of the resulting tables, we have bolded the highest adjusted power values. For the alternative $U(0,1)$ distribution, for a sample size of $N = 20$, the most powerful test first is $\hat{E}_2$, second is our non-weighted test. For $N = 50$, our two tests have the most power as TC , TV_1 , $EELR_2$, $\hat{E}_2$ and other tests that are bolded. For the alternative $CH(1.5)$ distribution, with a sample size of $N = 20, 50$ the performances $\widehat{EEX}_{1,n(k)}$ test is higher than the competitors. For the alternative $MEV(0.5)$ distribution, the $\widehat{EEX}_{1,n(k)}$ test demonstrated the best performance among all tests after the $\hat{E}_2$ test. Apart from $LN(0.8)$, our proposals are among the best working tests for all the UFR distributions for all sample sizes. Although the proposed tests did not perform relatively well for $\Gamma(0.4)$, they are not the worst. In general, we see from Tables 1-2 that the performance of tests depends on alternative distributions and no single test can beat the rest against all alternatives.

Table 1: Powers of $\widehat{EE\bar{X}}_{1,3(2)}$, $\widehat{EE\bar{X}}_{1,3(2)}^{(w)}$ statistics and different other statistics under alternative hypothesis at significance level $\alpha = 0.05$ for $N = 20, 50$

IFR, $N = 20$						
	$\Gamma(2)$	$W(1.4)$	HN	$U(0,1)$	$CH(1.5)$	$MEV(0.5)$
$\widehat{EE\bar{X}}_{1,3(2)}$	0.25	0.15	0.21	1.00	0.99	0.48
$\widehat{EE\bar{X}}_{1,3(2)}^{(w)}$	0.08	0.09	0.09	0.78	0.32	0.21
$TV\bar{E}$	0.42	0.27	0.14	0.52	0.62	0.09
TC	0.50	0.37	0.24	0.88	0.83	0.18
TV_1	0.51	0.37	0.23	0.87	0.82	0.18
TA_1	0.63	0.47	0.30	0.80	0.87	0.20
CKL_n	0.38	0.33	0.28	0.92	0.86	0.21
TV_2	0.45	0.31	0.19	0.83	0.76	0.15
TE_2	0.42	0.29	0.17	0.79	0.73	0.13
TA_2	0.35	0.24	0.13	0.69	0.64	0.09
TZ	0.27	0.17	0.07	0.54	0.50	0.21
TKL	0.63	0.49	0.31	0.84	0.91	0.22
TH	0.01	0.01	0.02	0.15	0.50	0.01
TJ	0.01	0.01	0.02	0.15	0.05	0.01
TT	0.00	0.00	0.00	0.06	0.01	0.01
T_X	0.01	0.01	0.01	0.22	0.09	0.01
T_1	0.55	0.35	0.18	0.52	0.72	0.13
T_2	0.55	0.36	0.18	0.55	0.74	0.14
T_3	0.55	0.35	0.17	0.51	0.72	0.13
T_4	0.62	0.46	0.28	0.69	0.83	0.19
I_1	0.46	0.27	0.12	0.27	0.60	0.09
I_2	0.47	0.29	0.14	0.39	0.67	0.11
$K_p S_n$	0.53	0.44	0.33	0.88	0.91	0.25
$EP S_n$	0.10	0.08	0.07	0.31	0.21	0.07
$Q_n(r)$	0.25	0.22	0.18	0.62	0.65	0.13
$Q_n^*(r)$	0.09	0.06	0.05	0.08	0.11	0.06
$L_n(p)$	0.46	0.34	0.21	0.57	0.78	0.14
G_n	0.46	0.35	0.21	0.70	0.84	0.15
KS_n	0.40	0.28	0.18	0.52	0.67	0.13
V	0.37	0.26	0.16	0.66	0.68	0.12
S^*	0.49	0.35	0.22	0.70	0.82	0.15
W_n^2	0.47	0.34	0.21	0.66	0.79	0.14
A_2^2	0.45	0.30	0.17	0.62	0.76	0.12
$H_n^{(1)}$	0.60	0.49	0.31	0.78	0.91	0.23
$H_n^{(2)}$	0.10	0.06	0.02	0.18	0.29	0.02
KS_n	0.46	0.35	0.24	0.72	0.79	0.18
CM_n	0.47	0.35	0.22	0.70	0.83	0.16
$T_{1,n}$	0.36	0.27	0.18	0.63	0.70	0.13
$T_{2,n}$	0.44	0.32	0.21	0.71	0.80	0.14
$T_{1,n}(\gamma)$	0.36	0.28	0.21	0.74	0.70	0.15
BH_n	0.48	0.29	0.21	0.79	0.78	0.15
HE_n	0.48	0.36	0.21	0.65	0.84	0.14
HE_n	0.48	0.37	0.21	0.62	0.84	0.15
EP_n	0.48	0.36	0.21	0.66	0.84	0.15
$T_{n,2.5}^{(1)}$	0.55	0.45	0.33	0.86	0.91	0.25
$T_{n,2.5}^{(2)}$	0.50	0.42	0.31	0.85	0.91	0.23
CO_n	0.54	0.37	0.19	0.50	0.81	0.13
M_n	0.46	0.29	0.13	0.37	0.70	0.09
$M_n(0)$	0.64	0.44	0.22	0.50	0.81	0.16
$M_n(0.99)$	0.47	0.34	0.19	0.63	0.83	0.13
J	0.62	0.47	0.29	0.78	0.89	0.22
$J_{0.5}$	0.60	0.43	0.26	0.66	0.84	0.19
D_{sp}	0.51	0.33	0.17	0.51	0.66	0.13
W_n^p	0.20	0.21	0.19	0.72	0.63	0.13
W_n^s	0.40	0.32	0.22	0.75	0.84	0.14
Q_n^p	0.38	0.32	0.23	0.86	0.85	0.17
COV	0.66	0.47	0.25	0.56	0.83	0.18
P_n	0.46	0.35	0.22	0.62	0.80	0.15
$EELR_1$	0.63	0.47	0.27	0.67	0.82	0.19
$EELR_2$	0.45	0.41	0.34	0.92	0.93	0.26
$\hat{E}_2$	0.02	0.29	0.37	1.00	0.99	0.83
IFR, $N = 50$						
	$\Gamma(2)$	$W(1.4)$	HN	$U(0,1)$	$CH(1.5)$	$MEV(0.5)$
$\widehat{EE\bar{X}}_{1,3(2)}$	0.53	0.44	0.61	1.00	1.00	0.98
$\widehat{EE\bar{X}}_{1,3(2)}^{(w)}$	0.37	0.39	0.39	1.00	0.98	0.87
$TV\bar{E}$	0.82	0.55	0.23	0.91	0.97	0.14
TC	0.81	0.64	0.42	1.00	1.00	0.33
TV_1	0.84	0.65	0.43	1.00	1.00	0.33
TA_1	0.95	0.80	0.48	0.99	1.00	0.33
CKL_n	0.66	0.62	0.54	1.00	1.00	0.44
TV_2	0.85	0.64	0.38	1.00	1.00	0.29
TE_2	0.82	0.58	0.31	1.00	0.99	0.22
TA_2	0.61	0.32	0.11	0.93	0.06	0.06
TZ	0.17	0.05	0.01	0.62	0.46	0.03
TKL	0.97	0.87	0.61	1.00	1.00	0.46
TH	0.01	0.02	0.04	0.58	0.22	0.03
TJ	0.01	0.02	0.04	0.58	0.22	0.03
TT	0.00	0.01	0.03	0.45	0.12	0.02
T_X	0.00	0.00	0.00	0.53	0.16	0.00
T_1	0.92	0.67	0.29	0.85	0.98	0.19
T_2	0.92	0.69	0.31	0.88	0.98	0.21
T_3	0.92	0.67	0.29	0.85	0.98	0.19
T_4	0.95	0.81	0.52	0.95	1.00	0.35
I_1	0.90	0.65	0.25	0.62	0.96	0.15
I_2	0.91	0.67	0.28	0.72	0.98	0.17
$K_p S_n$	0.90	0.85	0.71	1.00	1.00	0.57
$EP S_n$	0.19	0.13	0.09	0.67	0.45	0.08
$Q_n(r)$	0.65	0.56	0.43	0.96	0.98	0.30
$Q_n^*(r)$	0.24	0.11	0.05	0.09	0.25	0.06
$L_n(p)$	0.91	0.76	0.46	0.93	1.00	0.30
G_n	0.90	0.79	0.54	0.99	1.00	0.38
KS_n	0.83	0.64	0.39	0.98	0.96	0.26
V	0.79	0.59	0.35	0.99	0.99	0.24
S^*	0.90	0.76	0.49	0.99	1.00	0.34
W_n^2	0.90	0.75	0.48	0.98	1.00	0.32
A_2^2	0.92	0.74	0.45	0.99	1.00	0.30
$H_n^{(1)}$	0.94	0.88	0.65	0.99	1.00	0.49
$H_n^{(2)}$	0.62	0.37	0.13	0.78	0.94	0.06
KS_n	0.86	0.71	0.50	0.99	1.00	0.36
CM_n	0.90	0.77	0.53	0.99	1.00	0.37
$T_{1,n}$	0.81	0.65	0.43	0.99	0.99	0.30
$T_{2,n}$	0.88	0.75	0.50	0.99	1.00	0.36
$T_{1,n}(\gamma)$	0.74	0.60	0.43	1.00	0.99	0.31
$T_{2,n}(\gamma)$	0.80	0.68	0.48	1.00	1.00	0.35
BH_n	0.91	0.79	0.53	0.98	1.00	0.37
HE_n	0.91	0.79	0.53	0.98	1.00	0.37
HE_n	0.91	0.80	0.54	0.98	1.00	0.38
EP_n	0.91	0.80	0.54	0.98	1.00	0.38
$T_{n,2.5}^{(1)}$	0.90	0.81	0.64	1.00	1.00	0.48
$T_{n,2.5}^{(2)}$	0.85	0.80	0.64	1.00	1.00	0.49
CO_n	0.96	0.82	0.45	0.91	1.00	0.30
M_n	0.95	0.75	0.32	0.77	0.99	0.21
$M_n(0)$	0.97	0.83	0.41	0.82	1.00	0.28
$M_n(0.99)$	0.90	0.79	0.52	0.98	1.00	0.36
J	0.96	0.86	0.57	0.99	1.00	0.41
$J_{0.5}$	0.96	0.84	0.53	0.97	1.00	0.38
D_{sp}	0.91	0.69	0.35	1.00	1.00	0.24
W_n^p	0.57	0.61	0.55	1.00	1.00	0.42
W_n^s	0.80	0.74	0.57	1.00	1.00	0.43
Q_n^p	0.79	0.73	0.59	1.00	1.00	0.47
COV	0.96	0.82	0.43	0.87	0.99	0.29
P_n	0.87	0.75	0.49	0.96	1.00	0.33
$EELR_1$	0.94	0.85	0.57	0.98	1.00	0.41
$EELR_2$	0.83	0.77	0.62	1.00	1.00	0.47
$\hat{E}_2$	0.03	0.56	0.73	1.00	1.00	0.98

Table 2: Powers of $\widehat{EEX}_{1,3(2)}$, $\widehat{EEX}_{1,3(2)}^{(w)}$ statistics and different other statistics under alternative hypothesis at significance level $\alpha = 0.05$ for $N = 20, N = 50$

	UFR			DFR	
	$LN(0.8)$	$LN(1.5)$	$DL(1.5)$	$W(0.8)$	$\Gamma(0.4)$
$N = 20$					
$\widehat{EEX}_{1,3(2)}$	0.13	0.89	0.94	0.22	0.49
$\widehat{EEX}_{1,3(2)}^{(w)}$	0.05	0.53	0.98	0.37	0.27
TV_E	0.52	0.45	0.61	0.06	0.31
TC	0.40	0.19	0.65	0.02	0.24
TV_1	0.43	0.26	0.67	0.03	0.32
TA_1	0.49	0.00	0.81	0.01	0.00
CKL_n	0.18	0.43	0.48	0.08	0.25
TV_2	0.43	0.42	0.62	0.06	0.51
TE_2	0.45	0.50	0.60	0.09	0.56
TA_2	0.45	0.59	0.55	0.13	0.63
TZ	0.42	0.64	0.45	0.17	0.66
TKL	0.41	0.00	0.79	0.01	0.00
TH	0.01	0.70	0.01	0.34	0.89
TJ	0.01	0.70	0.01	0.33	0.89
TT	0.01	0.71	0.00	0.34	0.89
T_X	0.04	0.70	0.02	0.28	0.77
T_1	0.51	0.10	0.74	0.16	0.86
T_2	0.50	0.07	0.75	0.13	0.82
T_3	0.50	0.09	0.74	0.16	0.85
T_4	0.53	0.00	0.80	0.01	0.00
I_1	0.48	0.11	0.67	0.18	0.80
I_2	0.45	0.16	0.68	0.20	0.82
K_n	0.25	0.00	0.66	0.01	0.00
$EP S_n$	0.10	0.30	0.14	0.11	0.67
$Q_n(r)$	0.09	0.59	0.33	0.20	0.68
$Q_n^*(r)$	0.17	0.30	0.16	0.07	0.13
$L_n(p)$	0.27	0.55	0.64	0.23	0.82
G_n	0.24	0.67	0.62	0.24	0.76
KS_n	0.30	0.58	0.56	0.17	0.71
V	0.32	0.44	0.54	0.13	0.62
S^*	0.33	0.61	0.66	0.20	0.75
W_n^2	0.33	0.62	0.65	0.20	0.76
A_n^2	0.34	0.62	0.63	0.26	0.89
$H_n^{(1)}$	0.33	0.00	0.74	0.05	0.00
$H_n^{(2)}$	0.08	0.71	0.20	0.34	0.89
KS_n	0.28	0.55	0.62	0.14	0.62
CM_n	0.27	0.66	0.63	0.22	0.75
$T_{1,n}$	0.23	0.63	0.52	0.21	0.73
$T_{2,n}$	0.25	0.65	0.60	0.24	0.77
$T_{1,n}(\gamma)$	0.22	0.47	0.50	0.15	0.61
$T_{2,n}(\gamma)$	0.19	0.24	0.52	0.14	0.42
BH_n	0.26	0.67	0.64	0.24	0.77
HE_n	0.29	0.66	0.64	0.24	0.79
EP_n	0.25	0.67	0.64	0.24	0.76
$T_{n,2.5}^{(1)}$	0.27	0.18	0.67	0.04	0.33
$T_{n,2.5}^{(2)}$	0.22	0.38	0.61	0.10	0.50
CO_n	0.33	0.60	0.72	0.28	0.91
M_n	0.31	0.52	0.67	0.31	0.94
$M_n(0)$	0.47	0.42	0.81	0.20	0.90
$M_n(0.99)$	0.26	0.67	0.63	0.25	0.77
J	0.46	0.00	0.79	0.01	0.00
$J_{0.5}$	0.50	0.00	0.77	0.01	0.00
D_{sp}	0.56	0.46	0.71	0.14	0.79
W_s^n	0.12	0.60	0.23	0.15	0.45
W_s	0.20	0.61	0.53	0.18	0.53
Q_n	0.18	0.61	0.50	0.19	0.56
COV	0.61	0.03	0.84	0.00	0.00
P	0.24	0.64	0.61	0.22	0.73
$EELR_1$	0.40	0.56	0.78	0.15	0.59
$EELR_2$	0.19	0.03	0.54	0.02	0.07
E_2	0.04	0.00	0.06	0.01	0.95
$N = 50$					
$\widehat{EEX}_{1,3(2)}$	0.13	0.99	0.99	0.33	0.90
$\widehat{EEX}_{1,3(2)}^{(w)}$	0.06	0.91	1.00	0.32	0.81
TV_E	0.93	0.85	0.96	0.11	0.78
TC	0.73	0.69	0.93	0.09	0.85
TV_1	0.79	0.74	0.95	0.11	0.89
TA_1	0.84	0.00	0.99	0.00	0.00
CKL_n	0.24	0.85	0.76	0.19	0.68
TV_2	0.91	0.91	0.96	0.20	0.95
TE_2	0.92	0.92	0.95	0.27	0.96
TA_2	0.85	0.94	0.86	0.34	0.97
TZ	0.55	0.93	0.42	0.33	0.95
TKL	0.76	0.00	1.00	0.00	0.00
TH	0.01	0.94	0.01	0.56	1.00
TJ	0.01	0.93	0.01	0.55	0.99
TT	0.01	0.94	0.00	0.57	1.00
T_X	0.08	0.92	0.01	0.39	0.94
T_1	0.94	0.27	0.99	0.34	1.00
T_2	0.94	0.22	0.99	0.31	1.00
T_3	0.94	0.27	0.99	0.34	1.00
T_4	0.92	0.05	1.00	0.03	0.86
I_1	0.93	0.25	0.98	0.37	0.99
I_2	0.92	0.29	0.99	0.38	0.99
K_n	0.38	0.00	0.95	0.00	0.00
$EP S_n$	0.19	0.54	0.29	0.13	0.90
$Q_n(r)$	0.19	0.79	0.79	0.38	0.95
$Q_n^*(r)$	0.44	0.57	0.43	0.07	0.13
$L_n(p)$	0.60	0.90	0.98	0.48	0.99
G_n	0.47	0.95	0.97	0.48	0.98
KS_n	0.71	0.91	0.95	0.38	0.98
V	0.75	0.84	0.94	0.28	0.96
S^*	0.73	0.93	0.98	0.43	0.99
W_n^2	0.76	0.94	0.98	0.44	0.99
A_n^2	0.86	0.93	0.99	0.52	1.00
$H_n^{(1)}$	0.51	0.00	0.98	0.00	0.00
$H_n^{(2)}$	0.47	0.95	0.87	0.59	1.00
KS_n	0.62	0.92	0.96	0.35	0.97
CM_n	0.60	0.95	0.97	0.46	0.99
$T_{1,n}$	0.55	0.93	0.94	0.40	0.98
$T_{2,n}$	0.56	0.95	0.97	0.46	0.99
$T_{1,n}(\gamma)$	0.51	0.85	0.88	0.29	0.93
$T_{2,n}(\gamma)$	0.39	0.58	0.91	0.32	0.86
BH_n	0.47	0.95	0.97	0.48	0.99
HE_n	0.46	0.95	0.97	0.48	0.99
EP_n	0.45	0.95	0.97	0.48	0.99
$T_{n,2.5}^{(1)}$	0.49	0.74	0.96	0.25	0.92
$T_{n,2.5}^{(2)}$	0.32	0.83	0.93	0.31	0.93
CO_n	0.66	0.92	0.99	0.56	1.00
M_n	0.77	0.86	0.99	0.56	1.00
$M_n(0)$	0.84	0.82	1.00	0.48	1.00
$M_n(0.99)$	0.44	0.95	0.99	0.49	0.99
J	0.86	0.00	0.99	0.00	0.00
$J_{0.5}$	0.91	0.00	0.99	0.00	0.00
D_{sp}	0.97	0.83	0.99	0.31	0.99
W_s^n	0.16	0.93	0.61	0.31	0.84
W_s	0.29	0.93	0.88	0.34	0.88
COV	0.26	0.93	0.86	0.35	0.89
P_n	0.96	0.05	1.00	0.00	0.00
P	0.43	0.94	0.96	0.46	0.98

4 Real Data

In this section, we demonstrate the applicability of the proposed test statistics by analyzing a real dataset.

Dataset 1:

22.5, 37.5, 46.0, 48.5, 51.5, 53.0, 54.5, 57.5, 66.5, 68.0, 69.5, 76.5, 77.0, 78.5, 80.0, 81.5, 82.0, 83.0, 84.0, 91.5, 93.5, 102.5, 107.0, 108.5, 112.5, 113.5, 116.0, 117.0, 118.5, 119.0, 120.0, 122.5, 123.0, 127.5, 131.0, 132.5, 134.0.

Dataset 1, taken from Lawless [16], consists of failure times of 96 locomotive control units from a life-test terminated after 135,000 miles with 37 failures. Previous studies by Kandpal and Gupta [17] showed evidence against exponentiality and suggested a lognormal fit. Our two proposed tests also reject exponentiality, as reported in Table 3.

Dataset 2:

1.80, 3.43, 3.98, 4.23, 4.65, 2.89, 3.48, 4.06, 4.34, 4.84, 2.93, 3.57, 4.11, 4.37, 4.91, 3.03, 3.85, 4.13, 4.53, 4.99, 3.15, 3.92, 4.16, 4.62, 5.17.

Dataset 2, taken from Wadsworth [18], consists of inter-arrival times of 25 customers at a facility. Alizadeh Noughabi [19] reported evidence against exponentiality at the 5% significance level. Our tests also reject the exponential assumption, as shown in Table 3.

Dataset 3:

10.49, 8.8, 12.42, 4.58, 6.85, 4.58, 5., 4.75, 4.75, 12.25, 9.5, 13.54, 10.42, 4.65, 9.88, 6.21, 8.6, 7.06, 7.96, 7.89, 9.7, 13.9, 12.65, 10., 12.65, 12.07, 9.8, 13.54, 9.82, 13.54, 12.42, 12.73, 12.22, 12.25, 12.32, 8.75, 12., 17.5, 11.88, 13.13, 13.56, 15.44, 13.22, 7.28, 11.7, 11.7, 11.6, 10.9, 11.84, 8., 10.2, 5.77, 13.9, 4.58, 12.07, 15.44, 10.2, 11., 8.5, 10.99, 10.39, 9.9, 13.94, 15.21, 13.56, 9., 20.47, 15.22, 11.5, 13.9, 13.22, 10.48, 15.48, 9.8, 12.21, 13.56, 7.04.

Dataset 3 consists of 77 geoelectrically derived observations representing aquifer thickness, taken from Thomas and Jose [20]. Thomas and Jose [20] and Jose and Sathar [21] support a two-parameter Weibull fit, while Chaudhary and Gupta [22] rejected exponentiality. Our tests also reject exponentiality, as shown in Table 3.

Table 3: Description of models fitted for the three datasets

	Dataset 1 ($N = 37$)		Dataset 2 ($N = 25$)		Dataset 3 ($N = 77$)	
	$\widehat{EEX}_{1,n(k)}$ $m = 6$	$\widehat{EEX}_{1,n(k)}^{(w)}$ $m = 6$	$\widehat{EEX}_{1,n(k)}$ $m = 5$	$\widehat{EEX}_{1,n(k)}^{(w)}$ $m = 5$	$\widehat{EEX}_{1,n(k)}$ $m = 8$	$\widehat{EEX}_{1,n(k)}^{(w)}$ $m = 8$
Statistic value	-0.0120	-22.0506	-1.0870	-94.3389	-0.0992	-12.2260
Lower critical value	-0.6398	-0.5017	-0.9192	-0.6643	-0.5005	-0.4315
Upper critical value	-0.1756	-0.2149	-0.1689	-0.2141	-0.1980	-0.2281
Result	Reject H_0	Reject H_0	Reject H_0	Reject H_0	Reject H_0	Reject H_0

5 Conclusion

We defined the exponential extropy and weighted exponential extropy of n th upper k -record value, which have derived by generating function for the generalized extropy of n th upper k -record value and differentiating it with respect to α (see Askari et al. [23]). Two new extropies and tests for exponentiality based on these extropies of k -record values have been proposed in present study. Then, we compare their performances with 58 classical and recent tests for exponentiality, through a wide set of alternative distributions with different types of failure rate functions. A Monte Carlo simulation study was conducted to obtain the estimators, critical values, and powers of the tests. To demonstrate the practical applicability and reliability of the proposed tests, four real-life examples were also provided.

References

- [1] Lad, F., Sanfilippo, G. and Agro, G. (2015), Extropy: Complementary dual of entropy, *Statistical Science*, **30**(1), 40-58.
- [2] Jose, J. and Sathar, E. A. (2019), Residual extropy of k -record values, *Statistics and Probability Letters*, **146**, 1-6.
- [3] Jose, J. and Sathar, E. A. (2020), Past Extropy of k -Records, *Stochastics and Quality Control*, **35**(1), 25-38.
- [4] Sathar, E. A. and Nair, R. D. (2021), On dynamic survival extropy, *Communications in Statistics-Theory and Methods*, **50**(6), 1295-1313.
- [5] Gupta, N., Chaudhary, S. K. and Sahu, P. K. (2022), On weighted cumulative residual extropy and weighted negative cumulative extropy, arXiv:2211.13441[math.ST].
- [6] Tahmasebi, S. and Toomaj, A. (2021), On negative cumulative extropy with applications, *Communications in Statistics-Theory and Methods*, **51**, 5025-5047.
- [7] Noughabi, H. A. and Jarrahiferiz, J. (2019), On the estimation of extropy, *Journal of Nonparametric Statistics*, **31**(1), 88-99.
- [8] Dziubdziela, W. and Kopociński, B. (1976), Limiting properties of the k -th record values, *Applicationes Mathematicae*, **2**(15), 187-190.
- [9] Arnold, B. C., Balakrishnan, N. and Nagaraja, H. N. (1998), *Records*, Wiley Series in Probability and Statistics, John Wiley Sons, New York.
- [10] Ahsanullah, M. (1994), Record values, random record models and concomitants, *Journal of Statistical Research*, **28**, 89-109.
- [11] Vasicek, O. (1976), A test for normality based on sample entropy, *Journal of the Royal Statistical Society: Series B (Methodological)*, **38**, 54-59.

- [12] Xiong, P., Zhuang, W. and Qiu, G. (2020), Testing exponentiality based on the extropy of record values, *Journal of Applied Statistics*, **49**(4), 782-802.
- [13] Torabi, H., Montazeri, N. H. and Grané, A. (2018), A wide review on exponentiality tests and two competitive proposals with application on reliability, *Journal of Statistical Computation and Simulation*, **88**, 108-139.
- [14] Henze, N. and Meintanis, S. G. (2005), Recent and classical tests for exponentiality: A partial review with comparisons, *Metrika*, **61**, 29-45.
- [15] Zhang, J. and Boos, D. D. (1994), Adjusted power estimates in Monte-Carlo experiments, *Commun Statistics-Simulation and Computation* , **23**, 165-173.
- [16] Lawless, J. F. (2011), *Statistical Models and Methods for Lifetime Data*, vol. 362, Wiley, Hoboken.
- [17] Kandpal, G. and Gupta, N. (2024), A goodness-of-fit test for testing exponentiality based on normalized dynamic survival extropy, arXiv:2407.11634v1.
- [18] Wadsworth, HM. (1990), *Handbook of statistical methods for engineers and scientists*, McGraw-Hill, NewYork
- [19] Alizadeh Noughabi, H. (2015), Testing Exponentiality Based on the Likelihood Ratio and Power Comparison, DOI 10.1007/s40745-015-0041-0.
- [20] Thomas, P.Y. and Jose, J. (2020), A new bivariate distribution with Rayleigh and Lindley distributions as marginals, *Journal of Statistical Theory and Practice*, **14**(2), 1-33.
- [21] Jose, J., and Sathar, E.I.A. (2022), Symmetry being tested through simultaneous application of upper and lower k-records in extropy, *Journal of Statistical Computation and Simulation*, **92**(4), 830-846.
- [22] Chaudhary, S.K. and Gupta, N. (2023), Testing exponentiality using extropy of upper record values. arXiv:2301.02698.
- [23] Askari, M., Yousefzadeh, F., and Jomhoori, S. (2026), Generalized extropy of k-Records: properties and applications, *Journal of Mahani Mathematical Research*, **15**(2), 89-110. <https://doi.org/10.22103/jmmr.2025.24785.1761>.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Development of Repair Alert Model for Components Having Discrete Lifetimes in the Power Series Family

Atlekhani.M. ¹ Doostparast, M. ²

¹ Department of Statistics, Malayer University, Malayer, Iran

² Department of Statistics, Ferdowsi University of Mashhad, Iran

Abstract

This paper introduces a novel repair alert model designed for systems exhibiting discrete lifetime patterns, a common feature in maintenance data such as operational cycles, weekly failure reports, or counts of units produced before failure. We develop a robust framework under the assumption that system lifetimes follow a discrete distribution from the powerful and versatile power series family. The study tackles the critical challenge of parameter estimation for three prominent members of this family. Subsequently, we employ the Akaike Information Criterion (AIC) to identify the most suitable model and construct precise approximate confidence intervals for its parameters. The practicality and effectiveness of the proposed model are conclusively demonstrated through a comprehensive analysis of a real-world dataset, showcasing its significant potential for optimizing maintenance strategies.

Keywords: Maximum likelihood estimation, Power series family, Random sign censoring, Repair alert model.

¹Speaker. m.atlekhani@malayeru.ac.ir

²doustparast@um.ac.ir, doostparast@math.um.ac.ir

1 Introduction

Consider a system that begins operating at time zero. It may either fail at a random time X or be replaced by a new one at a random time Z , which may occur first. Typically, it is assumed that X and Z are independent. This policy is referred to as *random age replacement maintenance*; see, for example, [9, Chapter 2]. Thus, only the pair (Y, δ) is observed, where $Y = \min(X, Z)$ and $\delta = I(Z < X)$. Here $I(A)$ denotes the indicator function for the event A , such that $I(A) = 1$ when A occurs and $I(A) = 0$ otherwise. When Z and X are independent, the marginal distribution functions (DFs) of X and Z , denoted by $F_X(x)$ and $F_Z(z)$, respectively, can be identifiable based on the joint DF (Y, δ) ; For example, see [5, Theorem 1]. However, in practice, engineers often receive signals from an operating device that prompt them to replace it with a new device to avert the costs associated with random failure, suggesting that the independence assumption may not hold. In competing risks theory, it is established that the marginal DFs are unidentifiable based solely on observations (Y, δ) without considering the independence assumption between X and Z (see [10]). This unidentifiability issue becomes more complex when the failure processes are dependent. To address the identifiability challenge mentioned earlier, Cook [6] introduced the concept of *random sign censoring* (RSC), where X and Z are dependent, yet the marginal DF of X remains identifiable.

Definition 1.1. [7, 5] Let (X, Z) be a pair of lifetimes. Then $\{(Y = \min(X, Z), \delta = I(Z < X))\}$ is called a random signs censoring of X by Z if the event $\{Z < X\}$ is (stochastically) independent of X .

Definition 1.2. The subdistribution functions (sub-DFs) of X and Z are defined by $F_X^*(x) = P(X \leq x, X < Z)$ and $F_Z^*(z) = P(Z \leq z, Z < X)$, respectively. The conditional DFs of X and Z are defined by $\tilde{F}_X(x) = P(X \leq x | X < Z)$ and $\tilde{F}_Z(z) = P(Z \leq z | Z < X)$, respectively. Similarly, the sub-probability mass functions (sub-PMFs) and the conditional PMFs of X and Z are defined.

Lindqvist [7] proposed the Repair Alert Model (RAM) for analyzing the random sign censoring data (RSCD). The formal definition of the RAM is provided.

Definition 1.3. [7] The pair (X, Z) satisfies the requirements of the RAM if

- (i) Z is a random signs censoring of X ; that is, the event $\{Z < X\}$ is stochastically independent of X ;
- (ii) there exists an increasing function $G : [0, +\infty) \rightarrow [0, +\infty)$ with $G(0) = 0$, such that for all $x \in (0, +\infty)$,

$$P(Z \leq z | Z < X, X = x) = \frac{G(z)}{G(x)}, \quad 0 < z \leq x. \quad (1.1)$$

The function G is called the cumulative repair alert function (CRAF). Its derivative g (which we shall assume exists) is called the repair alert function (RAF). Statistical inferences based on RSCD have also been considered in the literature; for instance, see

[7, 1, 2]. They assumed that the lifetime X is continuous, but lifetimes can also be discrete. Examples of discrete lifetimes include the functioning of equipment in cycles, weekly field failure reports, and the number of pages printed by a device before failure. Recently, Atlekhani and Doostparast [3, 4] examined the RAM in situations where the discrete lifetime distribution X belongs to a broad class known as the telescopic family and power series family. They modified Definition 1.3 for discrete lifetimes. For the sake of brevity, the supports of X and Z are assumed to be $\mathbb{N} := \{0, 1, 2, \dots\}$, i. e. $S_x = S_z = \mathbb{N}$.

Definition 1.4. [3] The pair (X, Z) is called the discrete repair alert model (DRAM) if

- (i) the event $\{Z < X\}$ is stochastically independent of X ;
- (ii) there exists an increasing function $G : \mathbb{N} \rightarrow \mathbb{N}$ with $G(0) = 0$, such that for all $x \in \mathbb{N}$

$$P(Z \leq z | Z < X, X = x) = \frac{G(z)}{G(x)}, \quad z = 0, 1, \dots, x. \quad (1.2)$$

The next proposition provides the sub-DFs and conditional DFs under a DRAM. To do this, let

$$q = P(Z < X). \quad (1.3)$$

Proposition 1.5. [3] *Let the random vector (X, Z) follow the DRAM. The sub-DFs and the conditional DFs of X and Z are, respectively, given by*

$$\tilde{F}_X(x) = F_X(x), \quad F_X^*(x) = (1 - q)F_X(x), \quad (1.4)$$

$$\tilde{F}_Z(z) = F_X(z) + \sum_{k=z+1}^{\infty} \frac{G(z)}{G(k)} P(X = k), \quad F_Z^*(z) = q\tilde{F}_Z(z), \quad (1.5)$$

for all $x \in \mathbb{N}$ and $z \in \mathbb{N}$.

In this study, we assume that an individual has independently implemented a random age replacement policy N times. The available observations are denoted as $(Y_1, \delta_1); \dots; (Y_N, \delta_N)$, representing the Random Sign Censoring Data (RSCD). Our focus will be on the RAM when the lifetimes are represented as discrete random variables, with the distribution of device lifetimes $F_X(t)$ belonging to a class known as the *power series family*; see Section 3.1. This family encompasses a diverse range of discrete lifetime distributions, including logarithmic and negative binomial distributions, etc. Furthermore, we will explore the issue of parameter estimation and illustrate our findings through an analysis of a real dataset. Therefore, the remainder of this paper is organized as follows: In Section 3.1, the power series family and two important subclasses, i.e., logarithmic and negative binomial distributions, are studied in detail. In Section 3, the maximum likelihood estimates (MLEs) of the parameters and approximate confidence intervals are also derived. In Section 3.2, a real data set on ARC-1 VHF, reported by [8], is analyzed. Section 5 provides conclusions.

2 DRAM in the power series family

In this section, we assume that the DF of lifetime X under the DRAM belongs to the power series (PS) family. First, the formal definition of this family is given.

Definition 2.1. Suppose that the random variable X has the PMF as

$$f(x; \theta) = \frac{a(x)\theta^x}{c(\theta)}, \quad (2.1)$$

where $x \in \mathbb{N}$, $\theta > 0$, $a(x) > 0$, and $c(\theta) = \sum_{x=0}^{\infty} a(x)\theta^x$. Then X has the power series model and denoted by $X \sim PS(\theta)$.

Proposition 2.2. Under the DRAM, assume that $X \sim PS(\theta)$. Then, the sub-PMFs of X and Z are, respectively, given by

$$f_X^*(x; \theta) = (1 - q) \frac{a(x)\theta^x}{c(\theta)}, f_Z^*(z; \theta) = qg(z) \sum_{k=z}^{\infty} \frac{a(k)\theta^k}{c(\theta)G(k)}, \quad \forall x \in \mathbb{N}, z \in \mathbb{N}. \quad (2.2)$$

Remark 2.3. In Proposition 2.2, the conditional PMFs of X and Z are easily derived by dividing the corresponding sub-PMFs by the quantity q in Equation (1.3).

Now, we simplified the DRAM for three important members of the PS family.

2.1 Logarithmic distribution

In the PS distribution (2.1) if $\theta = p$, $c(\theta) = -\ln(1 - p)$, and $a(x) = x^{-1}$ that is $f_X(x; p) = \frac{-p^x}{x \ln(1 - p)}$ then, the logarithmic distribution is obtained.

Under the DRAM, assume that $G(t) = t^p$. Then $g(t) = G(t) - G(t - 1) = t^p - (t - 1)^p$. So, Equation (2.2) yields

$$f_X^*(x; p) = (1 - q) \frac{-p^x}{x \ln(1 - p)}, \quad \forall x \in \mathbb{N}, \quad (2.3)$$

Also,

$$f_Z^*(z; p) = q(z^p - (z - 1)^p) \sum_{k=z}^{\infty} \frac{-p^k}{k \ln(1 - p)k^p}, \quad \forall z \in \mathbb{N}. \quad (2.4)$$

2.2 Negative binomial distribution

In the PS distribution (2.1) if $\theta = p$, $c(\theta) = (1 - p)^{-r}$ and $a(x) = \Gamma(x + r)(x!\Gamma(r))^{-1}$ that is $f_X(x; r, p) = \Gamma(x + r)(x!\Gamma(r))^{-1}p^x(1 - p)^r$ then the negative binomial distribution is derived. In the DRAM, the sub-PMF of the lifetime X is

$$f_X^*(x; p) = (1 - q) \frac{\Gamma(x + r)}{x!\Gamma(r)} p^x (1 - p)^r, \quad \forall x \in \mathbb{N}. \quad (2.5)$$

If $G(t) = t^p$, then $g(t) = G(t) - G(t - 1) = t^p - (t - 1)^p$ and from Equation (2.2), we obtain

$$f_Z^*(z; p) = q(z^p - (z - 1)^p) \sum_{k=z}^{\infty} \frac{\Gamma(k + r)}{k!\Gamma(r)k^p} p^k (1 - p)^r, \quad \forall z \in \mathbb{N}. \quad (2.6)$$

2.3 Poisson distribution

In the PS distribution (2.1) if $\theta = \lambda$, $c(\theta) = \exp \lambda$ and $a(x) = x!^{-1}$ that is $f_X(x; \lambda) = \lambda^x e^{-\lambda} x!^{-1}$ then the Poisson distribution is derived. In the DRAM, the sub-PMF of the lifetime X is

$$f_X^*(x; \lambda) = (1 - q) \frac{\lambda^x e^{-\lambda}}{x!}, \quad \forall x \in \mathbb{N}. \quad (2.7)$$

If $G(t) = t^\lambda$, then $g(t) = G(t) - G(t - 1) = t^\lambda - (t - 1)^\lambda$ and from Equation (2.2), we obtain

$$f_Z^*(z; \lambda) = q(z^\lambda - (z - 1)^\lambda) \sum_{k=z}^{\infty} \frac{\lambda^k e^{-\lambda}}{k! k^\lambda}, \quad \forall z \in \mathbb{N}. \quad (2.8)$$

3 Parameter estimation

In this section, the estimates of parameters in a given DRAM are derived. Various approaches may be used to do this. In this paper, the ML method is considered. To end this, the likelihood function (LF) of the available RSCD $(\mathbf{x}, \mathbf{z})$ under the DRAM is

$$L(\theta, q; \mathbf{x}, \mathbf{z}) = \left(\prod_{i=1}^m f_X^*(x_i; \theta) \right) \left(\prod_{j=1}^n f_Z^*(z_j; \theta) \right). \quad (3.1)$$

The ML estimate of the parameter vector (θ, q) is then $L(\hat{\theta}, \hat{q}) = \arg \max_{\theta, q} L(\theta, q; \mathbf{x}, \mathbf{z})$. For the PS family (2.1), the LF of the RSCD $(\mathbf{x}, \mathbf{z})$ is obtained by substituting (2.2) into (3.1); that is,

$$L(\theta, q; \mathbf{x}, \mathbf{z}) = (1 - q)^m q^n \frac{\theta^{\sum_{i=1}^m x_i}}{(c(\theta))^m} \left(\prod_{i=1}^m a(x_i) \right) \left(\prod_{j=1}^n g(z_j) \sum_{k=z_j}^{\infty} \frac{a(k) \theta^k}{c(\theta) G(k)} \right). \quad (3.2)$$

The ML estimate for q is readily derived from Equation (3.2) as

$$\hat{q} = \frac{n}{m + n}. \quad (3.3)$$

So, the profile LF of the RSCD $(\mathbf{x}, \mathbf{z})$, denoted by L_P , is then obtained by substituting $\hat{q}$ in Equation (3.3) into Equation (3.2); that is,

$$L_P(\theta; \mathbf{x}, \mathbf{z}) = \left(\frac{m}{m + n} \right)^m \left(\frac{n}{m + n} \right)^n \frac{\theta^{\sum_{i=1}^m x_i}}{(c(\theta))^m} \left(\prod_{i=1}^m a(x_i) \right) \left(\prod_{j=1}^n g(z_j) \sum_{k=z_j}^{\infty} \frac{a(k) \theta^k}{c(\theta) G(k)} \right). \quad (3.4)$$

The profile ML estimate of θ is then derived by maximizing (3.4); that is,

$$\hat{\theta} = \arg \max_{\theta} \left[\frac{\theta^{\sum_{i=1}^m x_i}}{(c(\theta))^m} \prod_{i=1}^m a(x_i) \prod_{j=1}^n g(z_j) \sum_{k=z_j}^{\infty} \frac{a(k) \theta^k}{c(\theta) G(k)} \right]. \quad (3.5)$$

Notice that the ML estimate $\hat{q}$ in Equation (3.3) does not depend on the ML estimate $\hat{\theta}$ in Equation (3.5). The ML estimate of the parameters vector θ in the DRAM for special cases, mentioned in Section 3.1, are obtained in sequel.

3.1 The logarithmic distribution

For the logarithmic distribution, from Equations (2.3), (2.4) and (3.1), the logarithm of the likelihood function (LLF) is then

$$l(p, q; \mathbf{x}, \mathbf{z}) = n \ln(q) + m \ln(1q) + \left(\sum_{i=1}^m x_i \right) \ln p - (m+n) \ln(-\ln(1-p)) \\ - \sum_{i=1}^m \ln x_i + \sum_{j=1}^n \ln((z_j^p - (z_j - 1)^p) + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{p^k}{k^{p+1}} \right) \quad (3.6)$$

The logarithm of the profile LF, denoted by l_P , is obtained from Equations (3.3) and (3.6) as follows:

$$l_P(p, r; \mathbf{x}, \mathbf{z}) = n \ln \left(\frac{n}{m+n} \right) + m \ln \left(\frac{m}{m+n} \right) + \left(\sum_{i=1}^m x_i \right) \ln p - (m+n) \ln(-\ln(1-p)) \\ - \sum_{i=1}^m \ln x_i + \sum_{j=1}^n \ln((z_j^p - (z_j - 1)^p) + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{p^k}{k^{p+1}} \right) \quad (3.7)$$

The ML estimate of p is obtained by maximizing (3.6). For this purpose, one may use numerical methods; see Section 3.2.

3.2 The negative binomial distribution

For the negative binomial distribution, from Equations (2.5), (2.6) and (3.1), the LLF is then

$$l(p, r, q; \mathbf{x}, \mathbf{z}) = n \ln q + m \ln(1-q) + \left(\sum_{i=1}^m x_i \right) \ln p + r(m+n) \ln(1-p) \\ - (m+n) \ln \Gamma(r) + \sum_{i=1}^m (\ln \Gamma(x_i + r) - \ln x_i!) + \sum_{j=1}^n \ln (z_j^p - (z_j - 1)^p) \\ + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! k^p} p^k \right). \quad (3.8)$$

The logarithm of the profile LF, denoted by l_P , is obtained from Equations (3.3) and (3.8) as follows:

$$l_P(p, r; \mathbf{x}, \mathbf{z}) = n \ln \left(\frac{n}{m+n} \right) + m \ln \left(\frac{m}{m+n} \right) + \left(\sum_{i=1}^m x_i \right) \ln p \\ + r(m+n) \ln(1-p) - (m+n) \ln \Gamma(r) + \sum_{i=1}^m (\ln \Gamma(x_i + r) - \ln x_i!) \\ + \sum_{j=1}^n \ln (z_j^p - (z_j - 1)^p) + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! k^p} p^k \right). \quad (3.9)$$

The ML estimate of p is obtained by maximizing (3.9). To do this, one may use numerical methods; see Section 3.2.

3.3 The Poisson distribution

For the Poisson distribution, from Equations (2.7) - (3.1) the LLF is then

$$l(\lambda, q; \mathbf{x}, \mathbf{z}) = n \ln q + m \ln(1 - q) + (m + n)\lambda + \left(\sum_{i=1}^m x_i \right) \ln \lambda - \sum_{i=1}^m \ln x_i! \\ + \sum_{j=1}^n \ln (z_j^\lambda - (z_j - 1)^\lambda) + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\lambda^k}{k! k^\lambda} \right). \quad (3.10)$$

The ML estimate for q is given by Equation (3.3). So, the logarithm of the profile LF reads

$$l(\lambda, \hat{q}; \mathbf{x}, \mathbf{z}) = n \ln \left(\frac{n}{m + n} \right) + m \ln \left(\frac{m}{m + n} \right) + (m + n)\lambda + \left(\sum_{i=1}^m x_i \right) \ln \lambda - \sum_{i=1}^m \ln(x_i!) \\ + \sum_{j=1}^n \ln (z_j^\lambda - (z_j - 1)^\lambda) + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\lambda^k}{k! k^\lambda} \right). \quad (3.11)$$

The profile ML estimate of the parameter λ is derived numerically by maximizing Equation (3.11).

4 Case study

In this section a real data set is analysed to demonstrate applications of the findings of this paper.

4.1 The dataset

A real-world dataset pertaining to ARC-1 VHF communication transmitter-receivers used by a commercial airline. The dataset captures the number of operational hours until failure (see [8]). When failures occurred, the affected units were removed from the aircraft for maintenance. However, in some instances, the reported failures were unconfirmed, as the units demonstrated satisfactory functionality upon inspection at the maintenance center. Consequently, the failure times can be categorized into two distinct groups: **confirmed failures** ($\mathbf{x}$) and **unconfirmed failures** ($\mathbf{z}$). The dataset comprises $m = 218$ observations for confirmed failures (X) and $n = 107$ observations for unconfirmed failures (Z).

4.2 The repair alert model selection

For the analysis of this dataset the logarithmic, negative binomial, and Poisson distributions were considered for modeling the lifetime X . Numerical computations were performed using R (version 4.4.0) and Mathematica (version 9) software programs. The source code is publicly available at: <https://fumdribe.um.ac.ir/index.php/s/S8pBKfQFWYJKrmJ>. The empirical cumulative distribution function (ECDF) and the fitted parametric distribution functions as well the histogram of the data set and the corresponding fitted parametric PMFs are shown in Figures 2 and 1, respectively. These figures motivates that

Table 1: ML estimates of parameters based on ARC-1 VHF data set

Distribution	$G(t)$	$MLestimate$	$\widehat{LLF}$	AIC
logarithmic	t^p	$\hat{p} = 0.9996$	-2558.88	5119.76
Poisson	t^λ	$\hat{\lambda} = 3.535867$	-29783.56	59565.12
negative binomial	t^p	$r = 1, \hat{p} = 0.9963$	-2300.538	4603.076
		$r = 2, \hat{p} = 0.9923$	-2276.055	4554.11
		$r = 3, \hat{p} = 0.9884$	-2297.733	4597.466
		$r = 4, \hat{p} = 0.9845$	-2335.408	4672.816
		$r = 5, \hat{p} = 0.9806$	-2380.789	4763.578
		$r = 6, \hat{p} = 0.9767$	-2430.464	4862.928
		$r = 10, \hat{p} = 0.9616$	-2646.737	5295.474
		$r = 20, \hat{p} = 0.9257$	-3197.791	6397.582

the negative binomial distribution with parameters $r = 2$ and $\hat{p} = 0.9923$ is superior to the logarithmic and Poisson distributions for this data set; For more information about ECDF based on RSCD, see [7]. The Akaike Information Criterion (AIC) was employed to compare the fitted models. The AIC is defined as: $AIC = -2\widehat{LLF} + 2M$, where $\widehat{LLF}$ represents the maximized log-likelihood function (LLF) evaluated at the maximum likelihood (ML) estimates, and M denotes the number of estimated parameters. The results are summarized in Table 1. As illustrated in Table 1, the negative binomial distribution with parameters $r = 2$ and $\hat{p} = 0.9923$ exhibits a lower Akaike Information Criterion (AIC) value, indicating its superiority as the better-fitting model. An approximate 95 % confidence interval for p is $p \in (0.9917, 0.9929)$.

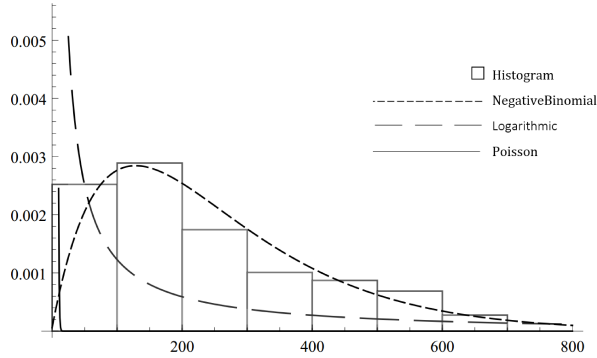


Figure 1

5 Conclusion

In this paper, we examined repair alert models for analyzing discrete lifetimes derived from maintenance programs. We focused on developing these models under the assumption that the lifetimes of the devices are discrete. We considered a broad class of discrete lifetime distributions known as the power series family. Furthermore, we addressed the challenges associated with parameter estimation for the logarithmic, poisson and negative

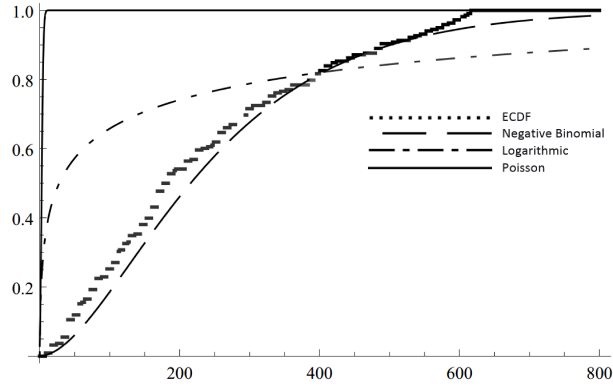


Figure 2

binomial distributions within this family. In analyzing a real dataset, RAM selection and its application in system availability evaluation were demonstrated. This paper may be extended in various directions. For example, RSCD may be analyzed by the Bayesian approach under symmetric and asymmetric loss functions such as square error and linear-exponential ones. Prediction of the component lifetime X based on the observed RSCD is also worth consideration.

References

- [1] Atlekhani, M., Doostparast, M. (2016), Inference under random maintenance policies with a polynomial form as the repair alert function. *Journal of Statistical Computation and Simulation*, **80**, 1415–1429.
- [2] Atlekhani, M., Doostparast, M. (2019), The repair alert models with fixed and random effects. *Communications in Statistics - Simulation and Computation*, **48**, 385–395.
- [3] Atlekhani, M., Doostparast, M. (2025), Repair alert model when the lifetimes are discretely distributed. *Communications in Statistics - Simulation and Computation*, **55**, 368–383.
- [4] Atlekhani, M., Doostparast, M. (2026), Repair alert model for one component systems with discrete lifetimes belonging to the power series family. *Statistics, Optimization and Information Computing*, to appear.
- [5] Cook, R. M. (1993), The total time on test statistics and age-dependent censoring. *Statistics and Probability Letters*, **18**, 307–312.
- [6] Cook, R. M. (1996), The design of reliability data bases, Part I and II: review of standard design concepts. *Reliability Engineering and System Safety*, **51**, 137–146, and 209–223.

- [7] Lindqvist, B. H., Støve, B., Langseth, H. (2006), Modelling of dependence between critical failure and preventive maintenance: The repair alert model. *Journal of Statistical Planning and Inference*, **136**, 1701–1717.
- [8] Mendenhall, W., Hader, R. (1958), Estimation of parameters of mixed exponentially distributed failure time distributions from censored life test data. *Biometrika*, **45**, 504–520.
- [9] Nakagawa, T. (2014), *Random maintenance policies*. Springer-Verlag, London.
- [10] Tsiatis, A. (1975), A nonidentifiability aspect of the problem of competing risks. *Proceedings of the National Academy of Sciences*, **72**, 20–22.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Stochastic Comparisons of Sequential $(n - r + 1)$ -out-of- n Systems Based on Relative Ageing

Esna-Ashari, M. ¹ Alimohammadi, M. ²

¹ Insurance Research Center, Tehran, Iran

² Department of Statistics, Faculty of Mathematical Sciences, Alzahra University, Tehran, Iran

Abstract

The hazard and reversed hazard rates are two important measures to study the lifetime random variables in reliability theory, survival analysis and stochastic modeling. In this article, we study the stochastic comparisons of sequential $(n - r + 1)$ -out-of- n systems based on these measures via relative ageing which is rather younger than classic ageing in one sample, as well as two samples cases.

Keywords: Hazard rate function, Reversed hazard rate function, Order statistics, Record values.

1 Introduction

A system with n independent components which functions if and only if at least $n - r + 1$ of the components are working is called a $(n - r + 1)$ -out-of- n system. Parallel and series systems are particular cases of such systems corresponding to $r = n$ and $r = 1$, respectively. These systems play an important role in reliability theory and life testing and in the real world. The lifetime of such a system is described by the r -th order

¹esnaashari@irc.ac.ir

²Speaker. m.alimohammadi@alzahra.ac.ir

statistic (OS) in a sample of size n when assuming that the remaining components are not affected by failures (cf. [7]). More generally, the failure of one component may influence the remaining components. Thus, a more flexible model for a $(n - r + 1)$ -out-of- n system should take the dependence structure into consideration. The failure of some component of the system can more or less strongly influence the life-length distributions of the remaining components. This can be thought of as damage caused by the i -th failure in the system. For example, the breakdown of an aircraft's engine will increase the load put on the remaining engines, such that their lifetime should tend to be shorter. For dealing with this type of problem, sequential order statistics (SOS) are proposed in [12, 3]. It is worth mentioning that one of the most important and applicable case is SOS under proportional hazard rate assumption which have a one to one correspondence to the generalized order statistics (GOS) model introduced by [12]. These models are closely connected to several other models of ordered random variables and, in particular they unify OS, progressively Type-II censored order statistics, record values, Pfeifer's record values, etc.

In the last three decades, a wide interest has been shown in investigating several stochastic orderings of OS and other ordered random variables. See e.g., recent articles [2, 8, 9] and references therein. In another direction, there is a few articles in the scene of relative aging such as [15] and [14] in which they argued just OS and record values. In this article, we provide some further results in this regard for flexible model of SOS.

The article is organized as follows. In Section 2, we recall some definitions and results which is used in the paper including the concept of SOS and aging notions. In Section 3, we study the relative ageing orderings among SOS in one sample, as well as two samples cases.

2 Preliminaries

We recall some definitions and well known notions and results which will be used in the sequel. The word increasing (decreasing) is used for non-decreasing (non-increasing) and all expectations are implicitly assumed to exist whenever they are written.

Let X and Y be two absolutely continuous nonnegative random variables. We denote, respectively, the cumulative distribution functions (cdf) by F and G , with survival functions (sf) $\bar{F} = 1 - F$ and $\bar{G} = 1 - G$, and probability density functions (pdf) f and g . We assume that $F^{-1}(0) = G^{-1}(0)$ (where F^{-1} , as customary, is the right-continuous inverse of F). Moreover, we denote the hazard rate (reversed hazard rate) functions of X and Y as $h_X = f/\bar{F}$ ($\kappa_X = f/F$) and $h_Y = g/\bar{G}$ ($\kappa_Y = g/G$), respectively. Also, let $M_X(t) = E[X - t|X > t]$ and $M_Y(t) = E[Y - t|Y > t]$ denote, respectively, the mean residual life functions of X and Y .

As a generalization of OS and record values, SOS and GOS were introduced by [12] as a unified approach to a variety of models of ordered random variables with different interpretations and statistical applications. SOS are defined by means of a triangular scheme of random variables where the r -th line contains $n - r + 1$ random variables with distribution function F_i , $i = 1, \dots, n$.

Definition 2.1. Let $F_1, \dots, F_n$ be continuous distribution functions with $F_1^{-1}(1) \leq \dots \leq$

$F_n^{-1}(1)$ and let $\{Y_{r,n}^{(j)}, 1 \leq j \leq n - r + 1\}$ be a sequence of independent and identically distributed random variables each distributed according to F_r , where $r = 1, \dots, n$. Let $X_{1,n}^{(j)} = Y_{1,n}^{(j)}, 1 \leq j \leq n$, and denote $X_{1,n}^* = \min_{j=1}^n X_{1,n}^{(j)}$. For $r = 2, \dots, n$, define $X_{r,n}^{(j)} = F_r^{-1}\{F_r(Y_{r,n}^{(j)})[1 - F_r(X_{r-1,n}^*)] + F_r(X_{r-1,n}^*)\}$ and denote $X_{r,n}^* = \min_{j=1}^{n-r+1} X_{r,n}^{(j)}$. Then, $X_{1,n}^*, \dots, X_{n,n}^*$ are called SOS based on $\{F_1, \dots, F_n\}$

From now on we consider a particular choice of the distribution functions $F_1, \dots, F_n$, namely

$$F_r(t) = 1 - (1 - F(t))^{\alpha_r}, \quad r = 1, \dots, n,$$

with some distribution function F and positive real numbers $\alpha_1, \dots, \alpha_n$. This is usually referred as the proportional hazard rate assumption. For example, if $\alpha_i = 1, \forall i$, or $\alpha_i = \frac{1}{n-i+1}, \forall i$, then the SOS would convert to the OS and record values, respectively (see Table 1 of [12]). In a sequential system, if the different lifetime distribution functions are denoted by $F_1, F_2, \dots$, then this leads to the interpretation that, after the i -th failure, the failure rate of remaining components will change from h_i to h_{i+1} , where $h_i = f_i/(1 - F_i)$. In the literature on SOS, it is usually assumed that there exists some common baseline hazard rate h with $h_i = \alpha_i \cdot h, i = 1, 2, \dots$, for positive constants $\alpha_1, \alpha_2, \dots$. Without changing the joint distribution in the above manner, the hazard rate remains $\alpha_1 \cdot h$ for any component working on its own. Thus, the introduction of unequal parameters α_i is to accommodate load sharing (cf. [3, 4]).

In this case, there exist several representations for the marginal pdf of SOS. Suppose that $n \in \mathbb{N}$, $m_i = (n - i + 1)\alpha_i - (n - i)\alpha_{i+1} - 1 \in \mathbb{R}, i = 1, \dots, n - 1$, $\tilde{m}_n = (m_1, \dots, m_{n-1})$, if $n \geq 2$ ($\tilde{m}_n \in \mathbb{R}$ is arbitrary, if $n = 1$) such that $\gamma_{(i,n,\tilde{m}_n)} = \alpha_n + n - i + \sum_{j=i}^{n-1} m_j > 0$ for all $i \in \{1, \dots, n - 1\}$, and $\gamma_{(n,n,\tilde{m}_n)} = \alpha_n$. In this case, we denote r -th SOS by $X_{(r,n,\tilde{m}_n)}$. Let $Y_{(r,n,\tilde{m}_n)}, r = 1, \dots, n$, be SOS based on G .

[19] obtained the expression

$$f_{X_{(r,n,\tilde{m}_n)}}(x) = c_{r-1}[\bar{F}(x)]^{\gamma_{(r,n,\tilde{m}_n)}-1} \xi_r(F(x))f(x), \quad x \in \mathbb{R}, \quad (2.1)$$

where $c_{r-1} = \prod_{i=1}^r \gamma_{(i,n,\tilde{m}_n)}, r = 1, \dots, n$, and ξ_r is a particular Meijer's G -function (cf. [5]).

According to Lemmas 2.1 of [1], we have the following recursive formula:

$$\xi_1(t) = 1, \quad \xi_r(t) = \int_0^t \xi_{r-1}(u)[1 - u]^{m_{r-1}} du, \quad 0 \leq t \leq 1, \quad r = 2, \dots, n. \quad (2.2)$$

Let us first recall some useful notions of stochastic orders (cf. [10]).

Definition 2.2. The random variable X is said to be smaller than Y in the

- likelihood ratio order (denoted by $X \leq_{lr} Y$) if $g(x)/f(x)$ is increasing in x ;
- hazard rate order (denoted by $X \leq_{hr} Y$) if $\bar{G}(x)/\bar{F}(x)$ is increasing in x or, equivalently, $h_Y(x) \leq h_X(x) \forall x$;
- reversed hazard rate order (denoted by $X \leq_{rh} Y$) if $G(x)/F(x)$ is increasing in x or, equivalently, $\kappa_X(x) \leq \kappa_Y(x) \forall x$;

- mean residual life order (denoted by $X \leq_{mrl} Y$) if $\bar{G}(t) \int_t^\infty \bar{F}(x) dx \leq \bar{F}(t) \int_t^\infty \bar{G}(x) dx$ $\forall t$ or, equivalently, $M_X(x) \leq M_Y(x) \forall x$;
- usual stochastic order (denoted by $X \leq_{st} Y$) if $\bar{F}(x) \leq \bar{G}(x) \forall x$.

We recall the following implications (cf. [10]):

$$\begin{array}{ccccc} X \leq_{lr} Y & \Rightarrow & X \leq_{hr} Y & \Rightarrow & X \leq_{mrl} Y \\ \downarrow & & \downarrow & & \downarrow \\ X \leq_{rh} Y & \Rightarrow & X \leq_{st} Y & \Rightarrow & E[X] \leq E[Y] \end{array}$$

The definitions of the ageing faster orders utilized in our paper are provided below (cf. [11, 17, 10]).

Definition 2.3. The random variable X is said to be ageing faster than Y in the

- hazard rate (denoted by $X \leq_c Y$) if $h_X(x)/h_Y(x)$ is increasing in x ;
- reversed hazard rate (denoted by $X \leq_b Y$) if $\kappa_Y(x)/\kappa_X(x)$ is increasing in x ;
- mean residual life (denoted by $X \leq_a Y$) if $M_X(x)/M_Y(x)$ is increasing in x .

Now, we recall the following definition about the very useful concept of total positivity (cf. [13]).

Definition 2.4. Let $\mathcal{X}$ and $\mathcal{Y}$ be subsets of the real line $\mathbb{R}$. A function $\lambda : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ is said to be totally positive of order 2 (TP_2) (reverse regular of order 2 (RR_2)) if

$$\lambda(x_1, y_1)\lambda(x_2, y_2) - \lambda(x_1, y_2)\lambda(x_2, y_1) \geq (\leq) 0,$$

for all $x_1 \leq x_2$ in $\mathcal{X}$ and all $y_1 \leq y_2$ in $\mathcal{Y}$.

Note that the TP_2 (RR_2) property is equivalent to $\lambda(x_2, y)/\lambda(x_1, y)$ is increasing (decreasing) in y when $x_1 \leq x_2$, whenever this ratio exists. Also note that the product of two TP_2 (RR_2) functions is TP_2 (RR_2). Moreover, if $\lambda(x, y)$ is TP_2 (RR_2) in (x, y) , then $\lambda_1(x)\lambda(x, y)\lambda_2(y)$ is TP_2 (RR_2) in (x, y) when λ_1 and λ_2 are two nonnegative functions (cf. [13]).

The following useful theorem was established by [13].

Theorem 2.5 (Basic composition theorem). *Let $\lambda_1 : \mathcal{X} \times \mathcal{Z} \rightarrow \mathbb{R}$, $\lambda_2 : \mathcal{Z} \times \mathcal{Y} \rightarrow \mathbb{R}$ and $\lambda : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ be Borel-measurable functions of two variables satisfying*

$$\lambda(x, y) = \int_{\mathcal{Z}} \lambda_1(x, z)\lambda_2(z, y)d\sigma(z),$$

where σ denotes a sigma-finite measure defined on $\mathcal{Z}$. If one of the functions λ_1 or λ_2 is RR_2 and the other TP_2 , then λ is RR_2 and, otherwise, if both of the them are RR_2 or TP_2 , then λ is TP_2 .

We also recall the following theorem from [13].

Theorem 2.6 (Variation diminishing property). *Let $\lambda : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ be a TP_2 Borel-measurable function, $d\sigma(y)$ be a sigma-finite measure defined on $\mathcal{Y}$, and $\lambda_1 : \mathcal{Y} \rightarrow \mathbb{R}$ be a bounded measurable function such that the integral*

$$\lambda_2(x) = \int_{\mathcal{Y}} \lambda(x, y) \lambda_1(y) d\sigma(y)$$

converges absolutely. If λ_1 changes sign at most $j \leq 1$ times on $\mathcal{Y}$, then λ_2 changes sign at most j times on $\mathcal{X}$. Moreover, if λ_2 changes sign j times, then it must have the same arrangement of signs as does λ_1 .

The lemma below, due to [16], is often used in establishing the monotonicity of a fraction in which the numerator and denominator are integrals or summations.

Lemma 2.7. *Assume that Θ is a subset of the real line $\mathbb{R}$, and let U be a nonnegative random variable having a cdf belonging to the family $\mathcal{P} = \{\Xi(\cdot|\theta), \theta \in \Theta\}$ which satisfies that, for $\theta_1, \theta_2 \in \Theta$,*

$$\Xi(\cdot|\theta_1) \leq_{st} (\geq_{st}) \Xi(\cdot|\theta_2), \text{ whenever } \theta_1 \leq \theta_2.$$

Let $\phi(u, \theta)$ be a real valued function defined on $\mathbb{R} \times \Theta$, which is measurable in u for each θ such that $E[\phi(U, \theta)]$ exists. Then, $E[\phi(U, \theta)]$ is

- (i) *increasing in θ , if $\phi(u, \theta)$ is increasing in θ and increasing (decreasing) in u ;*
- (ii) *decreasing in θ , if $\phi(u, \theta)$ is decreasing in θ and decreasing (increasing) in u .*

3 Main results

Now, we are ready to obtain our main results. At first, we argue the ageing faster orders in one-sample case.

Theorem 3.1. *Let $X_{(r,n,\tilde{m}_n)}$ be the SOS based on an absolutely continuous cdf F . If m_i is decreasing in i and $m_i \geq -1 \forall i$, then*

- (i) $X_{(r,n-1,\tilde{m}_{n-1})} \leq_c X_{(r,n,\tilde{m}_n)}$;
- (ii) $X_{(r+1,n+1,\tilde{m}_{n+1})} \leq_c X_{(r,n,\tilde{m}_n)}$;
- (iii) $X_{(r+1,n,\tilde{m}_n)} \leq_c X_{(r,n,\tilde{m}_n)}$.

Proof. (i) Using integration by parts and equations (2.1) and (2.2), we have

$$\bar{F}_{X_{(r,n,\tilde{m}_n)}}(x) = \sum_{j=1}^r \frac{c_{r-1}}{\prod_{i=j}^r \gamma_{(i,n,\tilde{m}_n)}} [\bar{F}(x)]^{\gamma_{(j,n,\tilde{m}_n)}} \xi_j(F(x))$$

and so

$$h_{X_{(r,n,\tilde{m}_n)}}(x) = h_X(x) \left[\sum_{j=1}^r \frac{1}{\prod_{i=j}^r \gamma_{(i,n,\tilde{m}_n)}} [\bar{F}(x)]^{\gamma_{(j,n,\tilde{m}_n)} - \gamma_{(r,n,\tilde{m}_n)}} \frac{\xi_j(F(x))}{\xi_r(F(x))} \right]^{-1}. \quad (3.1)$$

From (3.1), we have

$$\frac{h_{X(r,n-1,\tilde{m}_{n-1})}(x)}{h_{X(r,n,\tilde{m}_n)}(x)} = E[\phi_1(J_1, x)],$$

where

$$\begin{aligned}\phi_1(j, x) &= \frac{\prod_{i=j}^r \gamma(i, n-1, \tilde{m}_{n-1})}{\prod_{i=j}^r \gamma(i, n, \tilde{m}_n)} [\bar{F}(x)]^{(\gamma(j,n,\tilde{m}_n) - \gamma(r,n,\tilde{m}_n)) - (\gamma(j,n-1,\tilde{m}_{n-1}) - \gamma(r,n-1,\tilde{m}_{n-1}))} \\ &= \frac{\prod_{i=j}^r \gamma(i, n-1, \tilde{m}_{n-1})}{\prod_{i=j}^r \gamma(i, n, \tilde{m}_n)}\end{aligned}$$

and J_1 is a nonnegative random variable having a cdf belonging to the family $\mathcal{P}_1 = \{\mathcal{J}_1(\cdot|x), x \in \mathbb{R}_+\}$ with the probability mass function (pmf)

$$l_1(j|x) = c_1(x) I_{\{1 \leq j \leq r\}} \frac{1}{\prod_{i=j}^r \gamma(i, n-1, \tilde{m}_{n-1})} [\bar{F}(x)]^{\gamma(j,n-1,\tilde{m}_{n-1}) - \gamma(r,n-1,\tilde{m}_{n-1})} \frac{\xi_j(F(x))}{\xi_r(F(x))},$$

where

$$c_1(x) = \left[\sum_{j=1}^r \frac{1}{\prod_{i=j}^r \gamma(i, n-1, \tilde{m}_{n-1})} [\bar{F}(x)]^{\gamma(j,n-1,\tilde{m}_{n-1}) - \gamma(r,n-1,\tilde{m}_{n-1})} \frac{\xi_j(F(x))}{\xi_r(F(x))} \right]^{-1}$$

is the normalizing constant. Since $m_i \geq -1$, $\phi_1(j, x)$ is increasing in j . Also, it is constant with respect to x . We now show that $\mathcal{J}_1(\cdot|x_1) \leq_{lr} \mathcal{J}_1(\cdot|x_2)$ for $x_1 \leq x_2$. We have

$$\frac{l_1(j|x_2)}{l_1(j|x_1)} \propto \left[\frac{\bar{F}(x_2)}{\bar{F}(x_1)} \right]^{\sum_{i=j}^{r-1} (m_i+1)} \frac{\xi_j(F(x_2))}{\xi_j(F(x_1))}.$$

So, it is sufficient to show that $\xi_j(x)$ is TP_2 in (x, j) by induction. From (2.2), we have

$$\xi_j(x) = \int_{\mathbb{R}} \lambda_1(x, z) \lambda_2(j, z) dz,$$

where

$$\lambda_1(x, z) = I_{\{0 \leq z \leq x\}}, \quad \lambda_2(j, z) = \xi_{j-1}(z)(1-z)^{m_{j-1}}.$$

The function $\lambda_1(x, z)$ is TP_2 in (x, z) . Note that $\xi_j(z) = 1$ is TP_2 in (j, z) for $j = 1$. Since m_i is decreasing in i , $[1-z]^{m_{j-1}}$ is TP_2 in (j, z) . According to inductive assumption, $\lambda_2(j, z)$ is TP_2 in (j, z) and so, by Theorem 3.13, $\xi_j(x)$ is TP_2 in (x, j) . Therefore, $E[\phi_1(J_1, x)]$ is increasing in x according to part (i) of Lemma 3.11.

(ii) From (3.1), we have

$$\frac{h_{X(r+1,n+1,\tilde{m}_n)}(x)}{h_{X(r,n,\tilde{m}_n)}(x)} = E[\phi_2(J_2, x)],$$

where

$$\phi_2(j, x) = \frac{\prod_{i=j}^{r+1} \gamma_{(i, n+1, \tilde{m}_{n+1})}}{\prod_{i=j}^r \gamma_{(i, n, \tilde{m}_n)}} [\bar{F}(x)]^{-(m_r+1)} \frac{\xi_{r+1}(F(x))}{\xi_r(F(x))}$$

and J_2 is a nonnegative random variable having a cdf belonging to the family $\mathcal{P}_2 = \{\mathcal{J}_2(\cdot|x), x \in \mathbb{R}_+\}$ with the pmf

$$l_2(j|x) = c_2(x) \frac{I_{\{1 \leq j \leq r+1\}}}{\prod_{i=j}^{r+1} \gamma_{(i, n+1, \tilde{m}_{n+1})}} [\bar{F}(x)]^{\gamma_{(j, n+1, \tilde{m}_{n+1})} - \gamma_{(r+1, n+1, \tilde{m}_{n+1})}} \frac{\xi_j(F(x))}{\xi_{r+1}(F(x))},$$

where $c_2(x)$ is the normalizing constant. Since $m_i \geq -1$ and m_i is decreasing in i , $\phi_2(j, x)$ is increasing in x . Also, since $m_i \geq -1$, $\phi_2(j, x)$ is increasing in j . It is not difficult to see that $\mathcal{J}_2(\cdot|x_1) \leq_{lr} \mathcal{J}_2(\cdot|x_2)$ for $x_1 \leq x_2$. Therefore, $E[\phi_2(J_2, x)]$ is increasing in x according to part (i) of Lemma 3.11.

(iii) From part (ii) of theorem we have $X_{(r+1, n, \tilde{m}_n)} \leq_c X_{(r, n-1, \tilde{m}_{n-1})}$. So, the required result now follows from part (i).

The proof gets thus completed. $\square$

Theorem 3.2. Let $X_{(r, n, \tilde{m}_n)}$ be the SOS based on an absolutely continuous cdf F . If m_i is decreasing in i , $\gamma_{(i, n, \tilde{m}_n)} \geq 1$ and $m_i \geq 0 \forall i$, then

$$(i) \quad X_{(r, n, \tilde{m}_n)} \leq_b X_{(r, n-1, \tilde{m}_{n-1})};$$

$$(ii) \quad X_{(r, n, \tilde{m}_n)} \leq_b X_{(r+1, n+1, \tilde{m}_{n+1})};$$

$$(iii) \quad X_{(r, n, \tilde{m}_n)} \leq_b X_{(r+1, n, \tilde{m}_n)}.$$

Proof. From (2.1), we have

$$F_{X_{(r, n, \tilde{m}_n)}}(x) = c_{r-1} F(x) \int_0^1 (1 - tF(x))^{\gamma_{(r, n, \tilde{m}_n)} - 1} \xi_r(tF(x)) dt,$$

and so

$$\kappa_{X_{(r, n, \tilde{m}_n)}}(x) = \kappa_X(x) \left[\int_0^1 \left[\frac{(1 - tF(x))}{\bar{F}(x)} \right]^{\gamma_{(r, n, \tilde{m}_n)} - 1} \frac{\xi_r(tF(x))}{\xi_r(F(x))} dt \right]^{-1}. \quad (3.2)$$

Then, the rest of the proof follows along the lines of Theorem 3.1, and we omit it for the sake of brevity. $\square$

The following corollary is an immediate consequence of Theorems 3.1 and 3.2.

Corollary 3.3. Let $X_{(r, n, \tilde{m}_n)}$ and $X_{(r', n', \tilde{m}_{n'})}$ be the SOSs based on an absolutely continuous cdf F . If m_i is decreasing in i and $m_i \geq -1 \forall i$ ($\gamma_{(i, n, \tilde{m}_n)} \geq 1$ and $m_i \geq 0 \forall i$), then

$$X_{(r, n, \tilde{m}_n)} \geq_c (\leq_b) X_{(r', n', \tilde{m}_{n'})},$$

provided $r \leq r'$ and $n' - r' \leq n - r$.

Remark 3.4. [15] obtained Theorems 3.1 and 3.2 for special case of OS, i.e., for the case $\gamma_{(n,n,\tilde{m}_n)} = 1$ and $m_i = 0 \forall i$ (or, equivalently, $\alpha_i = 1 \forall i$).

We now shift our focus to the preservation of the ageing faster orderings among SOS from two distributions F and G .

Theorem 3.5. *Let $X_{(r,n,\tilde{m}_n)}$ and $Y_{(r,n,\tilde{m}_n)}$ be the SOSs based on absolutely continuous cdfs F and G , respectively, and $X \geq_{hr} Y$. If m_i is increasing in i and $m_i \leq -1 \forall i$, then, $X \leq_c Y$ implies*

$$X_{(r,n,\tilde{m}_n)} \leq_c Y_{(r,n,\tilde{m}_n)}.$$

Proof. From (3.1), we have

$$\frac{h_{X_{(r,n,\tilde{m}_n)}}(x)}{h_{Y_{(r,n,\tilde{m}_n)}}(x)} = \frac{h_X(x)}{h_Y(x)} E[\phi_3(J_3, x)],$$

where

$$\begin{aligned} \phi_3(j, x) &= \left[\frac{\bar{G}(x)}{\bar{F}(x)} \right]^{\gamma_{(j,n,\tilde{m}_n)} - \gamma_{(r,n,\tilde{m}_n)}} \frac{\xi_j(G(x))/\xi_r(G(x))}{\xi_j(F(x))/\xi_r(F(x))} \\ &= \left[\frac{\bar{G}(x)}{\bar{F}(x)} \right]^{\sum_{i=j}^{r-1} (m_i+1)} \frac{\xi_j(G(x))/\xi_j(F(x))}{\xi_r(G(x))/\xi_r(F(x))} \end{aligned}$$

and J_3 is a nonnegative random variable having a cdf belonging to the family $\mathcal{P}_3 = \{\mathcal{J}_3(\cdot|x), x \in \mathbb{R}_+\}$ with the pmf

$$l_3(j|x) = c_3(x) I_{\{1,\dots,r\}}(j) \frac{1}{\prod_{i=j}^r \gamma_{(i,n,\tilde{m}_n)}} [\bar{F}(x)]^{\gamma_{(j,n,\tilde{m}_n)} - \gamma_{(r,n-1,\tilde{m}_{n-1})}} \frac{\xi_j(F(x))}{\xi_r(F(x))},$$

where

$$c_3(x) = \left[\sum_{j=1}^r \frac{1}{\prod_{i=j}^r \gamma_{(i,n-1,\tilde{m}_{n-1})}} [\bar{F}(x)]^{\gamma_{(j,n-1,\tilde{m}_{n-1})} - \gamma_{(r,n-1,\tilde{m}_{n-1})}} \frac{\xi_j(F(x))}{\xi_r(F(x))} \right]^{-1}$$

is the normalizing constant.

Since $m_i \leq -1$ and $X \geq_{st} Y$, the term $[\bar{G}(x)/\bar{F}(x)]^{\sum_{i=j}^{r-1} (m_i+1)}$ is decreasing in j . Also, since $m_i \leq -1$ and $X \geq_{hr} Y$, it is increasing in x . As seen in the proof of Theorem 3.1, if m_i is increasing in i , then $\xi_j(x)$ is RR_2 in (x, j) and so $\xi_j(G(x))/\xi_j(F(x))$ is decreasing in j because $F(x) \leq G(x)$. We now show that the ratio

$$\eta_j(x) = (\xi_j(G(x))/\xi_j(F(x)))/(\xi_r(G(x))/\xi_r(F(x)))$$

is increasing in x by an inverse induction $j \in \{1, \dots, r\}$ starting from $j = r$. For $j = r$, the ratio is constant, i.e., $\eta_r(x) \equiv 1$, and, thus, the assertion is true. Consider now $j \in \{1, \dots, r\}$ and suppose that $\eta_j(x)$ is increasing in x . Then, from (2.2), we have

$$\frac{\xi_j(G(x))/\xi_j(F(x))}{\xi_r(G(x))/\xi_r(F(x))} = \frac{E[\psi_1(Z_2, x)]}{E[\psi_2(Z_1, x)]},$$

where

$$\begin{aligned}\psi_1(z, x) &= \frac{\xi_{j-1}(G(z))}{\xi_{j-1}(F(z))} \left[\frac{\bar{G}(z)}{\bar{F}(z)} \right]^{m_{j-1}} \frac{g(z)}{f(z)}, \\ \psi_2(z, x) &= \frac{\xi_{r-1}(G(z))}{\xi_{r-1}(F(z))} \left[\frac{\bar{G}(z)}{\bar{F}(z)} \right]^{m_{r-1}} \frac{g(z)}{f(z)},\end{aligned}$$

Z_1 and Z_2 are two nonnegative random variables with the pdfs

$$\begin{aligned}l_1(z|x, u) &= c_1(x, u) I_{\{0, \dots, x\}}(z) \xi_{r-1}(F(z)) [\bar{F}(z)]^{m_{r-1}} f(z), \\ l_2(z|x, u) &= c_2(x, u) I_{\{0, \dots, x\}}(z) \xi_{j-1}(F(z)) [\bar{F}(z)]^{m_{j-1}} f(z),\end{aligned}$$

in which $c_1(x, u)$ and $c_2(x, u)$ are the normalizing constants. It is shown below that the functions ψ_1 and ψ_2 satisfy the condition (iii) of Lemma 3.12. For each fixed x , we have

$$\frac{\psi_1(z, x)}{\psi_2(z, x)} = \frac{\xi_{j-1}(G(z))/\xi_{j-1}(F(z))}{\xi_{r-1}(G(z))/\xi_{r-1}(F(z))} \left[\frac{\bar{G}(z)}{\bar{F}(z)} \right]^{m_{j-1}-m_{r-1}}$$

is increasing in z according to the inductive assumption and $m_{j-1} \leq m_{r-1}$. As seen in the proof of Theorem 3.1, if m_i is increasing in i , then $\xi_j(x)$ is RR_2 in (x, j) . So, it is evident that $Z_1 \leq_{lr} Z_2$ since $m_{j-1} \leq m_{r-1}$. Thus, $\phi_3(j, x)$ is increasing in x . Finally, we show that $\mathcal{J}_3(\cdot|x_1) \geq_{lr} \mathcal{J}_3(\cdot|x_2)$ for $x_1 \leq x_2$. We have

$$\frac{l_3(j|x_2)}{l_3(j|x_1)} \propto \left[\frac{\bar{F}(x_2)}{\bar{F}(x_1)} \right]^{\sum_{i=j}^{r-1} (m_i+1)} \frac{\xi_j(F(x_2))}{\xi_j(F(x_1))}.$$

Since $m_i \leq -1$, $[\bar{F}(x_2)/\bar{F}(x_1)]^{\sum_{i=j}^{r-1} (m_i+1)}$ is decreasing in j . Therefore, $E[\phi_3(J_3, x)]$ is increasing in x according to part (i) of Lemma 3.11. $\square$

Theorem 3.6. *Let $X_{(r,n,\tilde{m}_n)}$ and $Y_{(r,n,\tilde{m}_n)}$ be the SOSs based on absolutely continuous cdfs F and G , respectively, and $X \leq_{hr} Y$. If m_i is decreasing in i , $\gamma_{(i,n,\tilde{m}_n)} \geq 1$ and $m_i \geq 0 \forall i$, then, $X \leq_b Y$ implies*

$$X_{(r,n,\tilde{m}_n)} \leq_b Y_{(r,n,\tilde{m}_n)}.$$

Proof. From (3.2), we have

$$\frac{\kappa_{Y_{(r,n,\tilde{m}_n)}}(x)}{\kappa_{X_{(r,n,\tilde{m}_n)}}(x)} = \frac{\kappa_Y(x)}{\kappa_X(x)} E[\phi_4(T, x)],$$

where

$$\phi_4(t, x) = \left[\frac{(1 - tF(x))/\bar{F}(x)}{(1 - tG(x))/\bar{G}(x)} \right]^{\gamma_{(r,n,\tilde{m}_n)}-1} \frac{\xi_r(tF(x))/\xi_r(F(x))}{\xi_r(tG(x))/\xi_r(G(x))}$$

and T is a nonnegative random variable having a cdf belonging to the family $\mathcal{P}_4 = \{\mathcal{T}(\cdot|x), x \in \mathbb{R}_+\}$ with the pdf

$$l_4(t|x) = c_4(x) I_{\{0 \leq t \leq 1\}} \left[\frac{(1 - tG(x))}{\bar{G}(x)} \right]^{\gamma_{(r,n,\tilde{m}_n)}-1} \frac{\xi_r(tG(x))}{\xi_r(G(x))},$$

where $c_4(x)$ is the normalizing constant. Since $X \leq_{st} Y$ and $\gamma_{(r,n,\tilde{m}_n)} \geq 1$, $\left[\frac{(1-tF(x))}{(1-tG(x))}\right]^{\gamma_{(r,n,\tilde{m}_n)}-1}$ is decreasing in t . We show that $\xi_r(tx_2)/\xi_r(tx_1)$ is decreasing in t for $x_1 \leq x_2$ by induction. This holds if, and only if, for every $c > 0$, the function $\xi_r(tx_2) - c\xi_r(tx_1)$ has at most one change of sign, and if so, it would be from $+$ to $-$. From (2.2), we have

$$\xi_r(tx_2) - c\xi_r(tx_1) = \int_{\mathbb{R}} \lambda(t, z) \lambda_1(z) dz,$$

where

$$\begin{aligned} \lambda(t, z) &= I_{\{0 \leq z \leq t\}}, \\ \lambda_1(z) &= \xi_r(zx_1) [1 - zx_1]^{m_r-1} x_1 \left(\frac{\xi_{r-1}(zx_2)}{\xi_{r-1}(zx_1)} \left[\frac{1 - zx_2}{1 - zx_1} \right]^{m_{r-1}} \frac{x_2}{x_1} - c \right). \end{aligned}$$

According to $m_i \geq 0$ and inductive assumption, the expression in parentheses of $\lambda_1(z)$ changes sign at most once, and if once from $+$ to $-$. In addition, $\lambda(t, z)$ is TP_2 in (t, z) . So, $\xi_r(tx_2) - c\xi_r(tx_1)$ changes sign at most once, and if once from $+$ to $-$ according to Theorem 3.12. We now show that the first term in $\phi_4(t, x)$ is increasing in x . We have

$$\begin{aligned} \frac{d}{dx} \ln \frac{(1-tF(x))/\bar{F}(x)}{(1-tG(x))/\bar{G}(x)} &= \frac{-tf(x)}{1-tF(x)} + h_X(x) + \frac{tg(x)}{1-tG(x)} - h_Y(x) \\ &= h_X(x) \left(1 - \frac{t\bar{F}(x)}{1-tF(x)}\right) - h_Y(x) \left(1 - \frac{t\bar{G}(x)}{1-tG(x)}\right) \\ &= \left(\frac{h_X(x)}{1-tF(x)} - \frac{h_Y(x)}{1-tG(x)}\right)(1-t) \geq 0, \end{aligned}$$

because of $X \leq_{hr} Y$ and $X \leq_{st} Y$. By a similar approach as in the proof of Theorem 3.5 and using Lemma 3.12, the second term in $\phi_4(t, x)$ is also increasing in x . One can see that

$$\frac{l_4(t|x_2)}{l_4(t|x_1)} \propto \left[\frac{(1-tG(x_2))}{(1-tG(x_1))} \right]^{\gamma_{(r,n,\tilde{m}_n)}-1} \frac{\xi_r(tG(x_2))}{\xi_r(tG(x_1))}$$

is decreasing in t and thus $\mathcal{T}(\cdot|x_1) \geq_{lr} \mathcal{T}(\cdot|x_2)$ for $x_1 \leq x_2$. Therefore, $E[\phi_4(T, x)]$ is increasing in x according to part (i) of Lemma 3.11. $\square$

Remark 3.7. (a) [14] obtained Theorem 3.5 for special case of record values, i.e., for the case $\gamma_{(n,n,\tilde{m}_n)} = 1$ and $m_i = -1 \forall i$ under the condition $X \geq_{st} Y$ instead of $X \geq_{hr} Y$. Note that if for all i , $m_i = -1$, then the terms including $\bar{G}(x)/\bar{F}(x)$ vanishes and therefore the condition $X \geq_{hr} Y$ can be weakened to $X \geq_{st} Y$;

(b) [15] obtained Theorem 3.6 for special case of OS.

Finally, it is worth mentioning that [15] showed that none of Theorems 3.1 and 3.2 is valid for the ordering of $\leq_a$ for OSs ($m_i = 0$). These cases are done similarly for different values of the parameter m_i 's, which is not listed here for brevity.

References

- [1] Alimohammadi, M. and Alamatsaz, M.H., 2011. Some new results on unimodality of generalized order statistics and their spacings. *Statistics and Probability Letters*, **81**(11), 1677-1682.
- [2] Alimohammadi, M., Esna-Ashari, M., Cramer, E. (2021). On dispersive and star orderings of random variables and order statistics. *Statistics and Probability Letters*, **170**, 109014.
- [3] Cramer, E. and Kamps, U., 1996. Sequential order statistics and k-out-of-n systems with sequentially adjusted failure rates. *Annals of the Institute of Statistical Mathematics*, **48**(3), 535-549.
- [4] Cramer, E. and Kamps, U., 2001. Sequential k -out-of- n systems. In: Balakrishnan, N. and Rao, C. R. (Eds.) *Handbook of Statistics: Advances in Reliability*, Vol. 20, pp. 301-372, Elsevier, Amsterdam.
- [5] Cramer, E. and Kamps, U., 2003. Marginal distributions of sequential and generalized order statistics. *Metrika*. **58**(3), 293-310.
- [6] Cramer, E., Kamps, U. and Rychlik, T., 2004. Unimodality of uniform generalized order statistics, with applications to mean bounds. *Annals of the Institute of Statistical Mathematics*. **56**(1), 183-192.
- [7] David, H.A. and Nagaraja, H.N., 2003. *Order Statistics*, Third edition. John Wiley & Sons, Hoboken, New Jersey.
- [8] Esna-Ashari, M., Alimohammadi, M., Cramer, E. (2022). Some new results on likelihood ratio ordering and aging properties of generalized order statistics. *Communications in Statistics-Theory and Methods*, **51**(14), 4667-4691.
- [9] Esna-Ashari, M., Balakrishnan, N. and Alimohammadi, M., 2023. HR and RHR orderings of generalized order statistics. *Metrika*, **86**(1), 131-148.
- [10] Finkelstein, M., 2006. On relative ordering of mean residual lifetime functions. *Statistics and Probability Letters*, **76**(9), 939-944.
- [11] Kalashnikov, V.V., Rachev, S.T., 1986. Characterization of queuing models and their stability. In: Prochorov, Y.K. (Ed.), *Probability Theory and Mathematical Statistics*, 37-53.
- [12] Kamps, U. 1995a. *A concept of generalized order statistics*. Teubner, Stuttgart.
- [13] Karlin, S., 1968. *Total Positivity*. Stanford University Press, Stanford, California.
- [14] Kayid, M., 2025. Stochastic comparisons of record values based on their relative aging. *Ricerche di Matematica*, **74**(4), 2261-2282.

- [15] Misra, N. and Francis, J., 2015. Relative ageing of $(n-k+1)$ -out-of- n systems. *Statistics and Probability Letters*, **106**, 272-280.
- [16] Misra, N. and van der Meulen, E.C., 2003. On stochastic properties of m -spacings. *Journal of Statistical Planning and Inference*, **115**(2), 683-697.
- [17] Rezaei, M., Gholizadeh, B. and Izadkhah, S., 2015. On relative reversed hazard rate order. *Communications in Statistics-Theory and Methods*, **44**(2), 300-308.
- [18] Shaked, M. and Shanthikumar, J.G., 2007. *Stochastic Orders*. Springer, New York.
- [19] Tavangar, M. and Asadi, M., 2012. Some unified characterization results on the generalized Pareto distributions based on generalized order statistics. *Metrika*, **75**, 997-1007.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Stochastic Ordering and Reliability Properties of a Discrete Lévy-type Distribution

Farbod, D.¹

Department of Mathematics, Faculty of Engineering Science, Quchan University of Technology, Quchan, Iran

Abstract

In this paper, we study stochastic ordering and reliability properties of a discrete distribution generated by the Lévy density (discrete Lévy-type distribution). We show that the family is ordered with respect to its parameter in the likelihood ratio order, which implies both hazard rate ordering and the usual stochastic order. These ordering relations provide a natural interpretation of the parameter in terms of lifetime behavior: larger parameter values correspond to stochastically larger lifetimes and lower hazard rates. We also investigate the asymptotic behavior of the main reliability characteristics. In particular, the hazard rate decreases asymptotically to zero, while the mean residual life increases linearly for large ages. These results provide a unified understanding of how the parameter of the discrete Lévy-type distribution governs both stochastic ordering and long-term reliability behavior, offering insight into the dependence between lifetime characteristics, hazard rate decay, and mean residual life growth.

Keywords: discrete Lévy-type distribution, stochastic ordering, hazard rate, mean residual life.

¹Speaker. d.farbod@qiet.ac.ir

1 Introduction

Discrete analogues of continuous heavy-tailed distributions have attracted considerable attention in recent years due to their usefulness in modeling count data exhibiting large variability and heavy tails (see, for example, [1, 6, 7]). Among these models, the *one-parameter discrete distribution generated by the Lévy density (discrete Lévy-type distribution)* provides a flexible framework for describing phenomena where extreme values occur with non-negligible probability (see, for example, [3, 4, 5]).

An important aspect in the study of parametric distribution families is the comparison of models through stochastic orderings. Such orderings provide a rigorous way to compare random variables in terms of their magnitude, variability, or failure behaviour, and they play a central role in reliability theory, survival analysis, and risk analysis.

In this paper we investigate stochastic ordering properties of the discrete Lévy-type distribution with respect to its parameter. In particular, we establish a likelihood ratio ordering and derive the corresponding hazard rate and stochastic orderings. We further examine reliability characteristics of the model by studying the asymptotic behaviour of the hazard rate and the mean residual life function.

1.1 Discrete Lévy-type distribution

Motivated by the continuous Lévy distribution, consider the Lévy probability density function with zero location parameter, given by (see, for example, [9])

$$f_x(\gamma) = \sqrt{\frac{\gamma}{2\pi}} x^{-\frac{3}{2}} \exp\left(-\frac{\gamma}{2x}\right), \quad \gamma > 0, x > 0. \quad (1.1)$$

Extracting the structural kernel from (1.1), define

$$s_x(\gamma) = x^{-\frac{3}{2}} \exp\left(-\frac{\gamma}{2x}\right). \quad (1.2)$$

The probability mass function of the discrete Lévy-type distribution is given by [3]

$$p_x(\gamma) = \frac{1}{c(\gamma)} s_x(\gamma), \quad x = 1, 2, \dots, \quad (1.3)$$

where the normalization constant is

$$c(\gamma) = \sum_{k=1}^{\infty} k^{-\frac{3}{2}} \exp\left(-\frac{\gamma}{2k}\right). \quad (1.4)$$

The discrete Lévy-type distribution exhibits heavy-tailed behavior, characterized by the moment condition $E[X^r] < \infty$ if and only if $r < \frac{1}{2}$, (see [5]).

The distribution (1.3) can be viewed as a discrete analogue of the Lévy kernel of the $\frac{1}{2}$ -stable law. Motivated by the continuous Lévy density, the discrete Lévy-type model is defined on the support $\{1, 2, \dots\}$ and retains the same tail exponent while reflecting discretization effects [5].

2 Main Results

2.1 Stochastic ordering

The stochastic ordering of X_γ with respect to γ helps clarify how the tail behaviour and reliability characteristics change as the parameter varies. We briefly recall several stochastic orders commonly used in reliability theory (see, for example, [2, 10]).

Definition 2.1. [2] A random variable X_1 is said to be smaller than X_2 in the usual stochastic order, written $X_1 \leq_{st} X_2$, if

$$P(X_1 > t) \leq P(X_2 > t), \quad t \in \mathbb{R}.$$

Definition 2.2. [10] A random variable X_1 is said to be smaller than X_2 in the likelihood ratio order, written $X_1 \leq_{lr} X_2$, if

$$\frac{P(X_2 = x)}{P(X_1 = x)}$$

is increasing in x over the union of their supports.

Definition 2.3. [10] A random variable X_1 is said to be smaller than X_2 in the hazard rate order, written $X_1 \leq_{hr} X_2$, if

$$\frac{P(X_1 \geq x)}{P(X_2 \geq x)}$$

is decreasing in x .

We now state the main ordering result for the model.

Theorem 2.4. *Let X_{γ_1} and X_{γ_2} be discrete Lévy-type random variables. If $\gamma_1 < \gamma_2$, then*

$$X_{\gamma_1} \geq_{lr} X_{\gamma_2}.$$

Consequently,

$$X_{\gamma_1} \geq_{hr} X_{\gamma_2} \quad \text{and} \quad X_{\gamma_1} \geq_{st} X_{\gamma_2}.$$

Proof. Consider

$$R(x) = \frac{p_x(\gamma_2)}{p_x(\gamma_1)}, \quad x = 1, 2, \dots$$

Using (1.3),

$$R(x) = \frac{c(\gamma_1)}{c(\gamma_2)} \exp\left(-\frac{\gamma_2 - \gamma_1}{2x}\right).$$

Hence

$$\frac{R(x+1)}{R(x)} = \exp\left(\frac{\gamma_2 - \gamma_1}{2x(x+1)}\right) > 1,$$

so $R(x)$ is increasing and therefore

$$X_{\gamma_1} \geq_{lr} X_{\gamma_2}.$$

The remaining implications follow from standard ordering results [10]. □

2.2 Reliability interpretation

In reliability analysis the discrete hazard rate, corresponding to the parameter γ , is defined as

$$h_x(\gamma) = \frac{p_x(\gamma)}{P(X_\gamma \geq x)}, \quad x = 1, 2, \dots, \quad (2.1)$$

From the likelihood ratio ordering, if $\gamma_1 < \gamma_2$ then $X_{\gamma_1} \geq_{st} X_{\gamma_2}$, implying that smaller values of γ correspond to stochastically larger lifetimes.

Proposition 2.5. *For a fixed $\gamma > 0$, the hazard rate $h_x(\gamma)$ of the discrete Lévy-type distribution (1.3) asymptotically behaves as*

$$h_x(\gamma) \sim \frac{1}{2x}, \quad \text{as } x \rightarrow \infty. \quad (2.2)$$

where $\sim$ denotes asymptotic equivalence as $x \rightarrow \infty$.

Proof. Since $\exp(-\frac{\gamma}{2x}) \rightarrow 1$ as $x \rightarrow \infty$, the probability mass function satisfies [5]

$$p_x(\gamma) \sim \frac{1}{c(\gamma)} x^{-\frac{3}{2}}. \quad (2.3)$$

Consequently, the survival function can be approximated by

$$P(X_\gamma \geq x) \sim \frac{1}{c(\gamma)} \sum_{j=x}^{\infty} j^{-\frac{3}{2}} \sim \frac{2}{c(\gamma)} x^{-\frac{1}{2}}. \quad (2.4)$$

Substituting (2.3) and (2.4) into (2.1) yields (2.2). $\square$

Remark 2.6. The asymptotic relation $h_x(\gamma) \sim (2x)^{-1}$ indicates that the hazard rate decreases at large values of x , which is consistent with heavy-tailed lifetime behavior.

2.3 Mean residual life

The mean residual life function is a fundamental concept in reliability theory [2]. For $x \in \mathbb{N}$, it is defined by

$$m_x(\gamma) = E[X_\gamma - x \mid X_\gamma \geq x].$$

We now establish the asymptotic behavior of the mean residual life function.

Proposition 2.7. *For the discrete Lévy-type model (1.3), if (2.4) holds, then as $x \rightarrow \infty$*

$$m_x(\gamma) \sim 2x. \quad (2.5)$$

Proof. For integer-valued lifetimes,

$$m_x(\gamma) = \frac{\sum_{j=x+1}^{\infty} P(X_\gamma \geq j)}{P(X_\gamma \geq x)}.$$

Assume that $P(X_\gamma \geq j) \sim Aj^{-\frac{1}{2}}$, where $A = \frac{2}{c(\gamma)}$. By the integral test, we have

$$\sum_{j=x+1}^{\infty} P(X_\gamma \geq j) \sim A \int_x^{\infty} t^{-\frac{1}{2}} dt = 2A\sqrt{x}.$$

Since $P(X_\gamma \geq x) \sim Ax^{-\frac{1}{2}}$, it follows that

$$m_x(\gamma) \sim \frac{2A\sqrt{x}}{Ax^{-\frac{1}{2}}} = 2x.$$

□

Remark 2.8. The asymptotic relation (2.5) shows that the expected remaining lifetime grows linearly with age for the distribution defined in (1.3).

3 Conclusion

In this study, we established the fundamental reliability framework and stochastic ordering properties for the discrete Lévy-type distribution. Our analysis demonstrated that the parameter γ serves as a crucial determinant of the model's tail behavior and lifetime characteristics. By proving that the distribution belongs to the likelihood ratio ordering family with respect to γ , we provided a mathematical basis for its stochastic monotonicity; specifically, an increase in γ consistently leads to stochastically larger lifetimes and a reduction in the hazard rate.

Furthermore, our asymptotic investigation revealed that the model exhibits a vanishing hazard rate, decaying at the rate of $(2x)^{-1}$. This property, coupled with the linear growth of the mean residual life ($m_x(\gamma) \sim 2x$), characterizes the distribution as a suitable candidate for modeling systems with long-term survival or heavy-tailed count data where the risk of failure decreases over time. These findings not only offer a unified understanding of the discrete Lévy-type distribution but also enhance its applicability in reliability engineering and survival analysis, where capturing such limiting behaviors is essential.

References

- [1] Astola, J. and Danielian, E. (2007), *Frequency Distributions in Biomolecular Systems and Growing Networks*, Tampere International Center for Signal Processing (TICSP), Series no. 31, Tampere, Finland.
- [2] Barlow, R. E. and Proschan, F. (1975), *Statistical Theory of Reliability and Life Testing*, Holt, Rinehart and Winston.
- [3] Farbod, D. (2008), The asymptotic properties of some discrete distributions generated by Levy's law, *Far East Journal of Theoretical Statistics*, **26**(1), 121-128.

- [4] Farbod, D. (2014), Some statistical inferences for two frequency distributions arising in bioinformatics, *Applied Mathematics E-Notes*, **14**, 151-160.
- [5] Farbod, D. and Basirat, M. (2024), On the discrete distribution generated by Levy probability, *Journal of Mathematical Sciences*, **280**(2), 198-211.
- [6] Kuznetsov, V. A. (2017), *Mathematical modeling of avidity distribution and estimating general binding properties of transcription factors from genome-wide binding profiles*, In: Tatarinova, T. V. and Nikolsky, Y. (eds), *Biological Networks and Pathway Analysis*, pp. 193-276, Springer, New York.
- [7] Kuznetsov, V. A., Grageda, A. and Farbod, D. (2022), Generalized hypergeometric distributions generated by birth-death process in bioinformatics, *Markov Processes and Related Fields*, **28**(2), 303-327.
- [8] Marshall, A. W., Olkin, I. and Arnold, B. C. (2010), *Stochastic Ordering. In: Inequalities: Theory of Majorization and Its Applications*. Springer Series in Statistics. Springer, New York.
- [9] Nolan, J. P. (2020), *Univariate Stable Distributions: Models for Heavy-Tailed Data*, Springer Series in Operations Research and Financial Engineering, Springer.
- [10] Shaked, M. and Shanthikumar, J. G. (2007), *Stochastic Orders*, Springer.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Agent-Informed Discretization: A Knowledge-Informed Reinforcement Learning Framework for Smart Steel Production Maintenance

Gholami, Z. ¹ Haerian, N. ² Safari, A. ³ Haghighi, F. ⁴

Department of Science, Faculty of Mathematics, Statistics and Computer Science, University
of Tehran, Tehran, Iran

Abstract

In a scrap-based steel production line, traditional maintenance policies fail due to complex interactions between hammer wear, buffer limits, demand, and RUL. Existing DRL approaches are reactive and rely on pre-defined degradation models. This paper proposes Agent-Informed Discretization: a DDQN agent is first trained in a continuous space with 21 actions (0-20 hammer replacements). Knowledge extraction via sensitivity and action-frequency analysis identifies strategic break-points and prunes actions to 5 effective ones (0,3,5,15,20), forming an interpretable 96-state discrete space. Zero-shot evaluation shows our model achieves higher mean reward (−1549.2 vs. −1534.6 and −1275.21) and lower standard deviation (22.96 vs. 16.11 and 12.67). Unlike black-box DRL, the resulting policy is transparent, expressed as simple if-then rules from learned behavior. This work demonstrates that combining expressive DRL training with principled knowledge extraction yields high-performance, low-cost, deployable maintenance strategies for real industrial systems.

¹Speaker. gholami.z78@ut.ac.ir

²Speaker. nhaerian@ut.ac.ir (^{1,2} These authors contributed equally to this work.)

³a.Safari@ut.ac.ir

⁴fhaghighi@ut.ac.ir

Keywords: Reinforcement learning, Maintenance optimization, Scrap-based steel production line, Remaining useful life, Multi-component system.

1 Introduction

The steel industry is one of the most energy-intensive and carbon-emitting sectors globally. Scrap-based steel production, which recycles ferrous scrap through shredding and melting, offers a more sustainable alternative to primary production. However, the efficiency of such a production line heavily depends on the reliable operation of its core components—particularly the shredder, whose hammers undergo progressive wear, and the downstream workstation, which faces variable demand and stochastic failures. Between these two stations lies a buffer of limited capacity, creating complex material flow interdependencies. A failure or unnecessary stoppage at any point can propagate through the system, leading to production loss, high corrective costs, and unmet demand [1].

Maintenance policy optimization in such multi-component systems is therefore critical but challenging. Traditional approaches include time-based maintenance (TBM), which follows fixed schedules, and corrective maintenance (CM), which acts only after failure. Both ignore real-time system condition and dynamic interactions among components. Condition-based maintenance (CBM) improves upon these by using sensor or operational data to trigger maintenance only when needed. However, implementing CBM in complex industrial lines remains difficult because direct wear data (e.g., hammer thickness sensors) is often expensive or unavailable, and degradation processes are stochastic with non-trivial component interactions [2].

Recent advances in deep reinforcement learning (DRL) have shown promise in overcoming some of these limitations. DRL agents learn optimal maintenance policies through direct interaction with a simulated environment, without requiring an explicit analytical model of system dynamics [3]. Nevertheless, current DRL-based CBM solutions suffer from two major gaps. First, most models rely on either pre-defined degradation models (e.g., gamma processes with known parameters) or direct state measurements that are rarely available in real industrial settings. Second—and more critically—existing DRL agents are fundamentally *reactive*: they base decisions solely on the current observed state, missing the opportunity to anticipate future events [1]. For example, if the second workstation is about to fail (low remaining useful life, RUL), a reactive agent cannot proactively schedule preventive maintenance on the first workstation to coincide with the upcoming unavoidable stoppage. This lack of forward-looking coordination leads to multiple separate stoppages, higher downtime, and increased costs.

To address these gaps, this paper builds upon the foundational model of Ferreira Neto, Cavalcante, and Do [1], who first formulated a scrap-based steel production line as a Markov Decision Process and applied a Double Deep Q-Network (DDQN) in a discrete state space. We extend their work by introducing a novel framework called **Agent-Informed Discretization**. The core idea is to first train a high-capacity DRL agent in a *continuous state space* with a *rich action space* (replacing 0 to 20 hammers) without any engineering simplifications. This agent learns an optimal, though black-box,

policy. Then, instead of deploying this computationally expensive continuous agent directly, we *extract knowledge* from it: we perform sensitivity analysis on the learned value function to identify strategic breakpoints (phase transitions) where the optimal decision changes sharply, and we analyze action frequency to identify which actions are actually used by the optimal policy. These extracted insights are then used to construct a *pruned action space* (only 5 effective actions: replace 0, 3, 5, 15, or 20 hammers) and an *informed discrete state space* (96 interpretable states). The resulting policy is transparent, low-complexity, and suitable for real-time industrial controllers, while retaining the performance of the original continuous DRL model.

The main contributions of this work are:

- A novel methodology that bridges the gap between expressive continuous DRL and interpretable, deployable maintenance policies through knowledge extraction, building directly on the foundational model of [1].
- Integration of prognostic information (RUL of the downstream workstation) as a forward-looking signal, enabling opportunistic maintenance coordination.
- Demonstration that action-frequency analysis can naturally prune a large action space (21 actions) to a small set of strategically meaningful actions.
- Construction of an informed discrete state space (96 states) that preserves critical system dynamics while being fully interpretable.
- Quantitative evaluation in a zero-shot setting, showing that the final model achieves higher mean reward and significantly lower variance compared to baseline models.

The remainder of this paper is organized as follows. Section 2 describes the scrap-based steel production line and its mathematical modeling as a Markov Decision Process. Section 3 presents the proposed Agent-Informed Discretization framework in detail. Section 4 reports experimental results and comparisons. Section 5 concludes the paper with a discussion of limitations and future work.

2 Methodology

The proposed Agent-Informed Discretization framework consists of three main phases: (i) development of a continuous simulation environment and training of a high-capacity DDQN agent, (ii) knowledge extraction from the trained agent via sensitivity analysis and action-frequency analysis, and (iii) construction of an informed discrete state-action space that is interpretable and suitable for real-time industrial deployment.

2.1 Advanced Simulation Environment

Unlike previous studies that used a discrete state space with only three coarse variables (production rate, buffer level, demand), we develop a continuous-state simulation environment that captures the full dynamics of the scrap-based steel production line. The

environment models the shredder (first workstation), a buffer, and the downstream workstation with stochastic failures and variable demand.

The system state at time step t is represented by a continuous vector of five key variables:

$$\mathbf{s}_t = [P_t, b_t, d_t, RUL_t, Def_t] \quad (2.1)$$

where:

- P_t : instantaneous production rate of the shredder (tons/hour), affected by hammer wear and input scrap quality.
- b_t : actual buffer inventory level (tons), $0 \leq B_t \leq K$ with $K = 1000$ tons.
- d_t : demand from the second workstation (tons/hour), sampled from a normal distribution $\mathcal{N}(15, 1)$.
- RUL_t : remaining useful life of the second workstation (hours), which decreases deterministically over time and resets after a failure or periodic inspection.
- Def_t : number of defective hammers (0 to 20), where a hammer is considered defective when its degradation level exceeds 0.2.

The action space is significantly expanded compared to classical binary decisions (continue / stop for PM). The agent can choose to replace any integer number of hammers from 0 to 20. Thus, the action set contains $|\mathcal{A}| = 21$ discrete actions:

$$a_t \in \{0, 1, 2, \dots, 20\} \quad (2.2)$$

where a_t indicates how many hammers are replaced during a maintenance intervention. If $a_t = 0$, the shredder continues production without stoppage; if $a_t > 0$, a maintenance stop is triggered, all hammers with degradation above 0.2 are replaced (at least the number specified), and the associated costs and downtime are incurred.

Hammer wear is modeled using independent gamma processes. Each hammer i has a degradation level $d_i(t)$ evolving as:

$$d_i(t + \Delta t) = d_i(t) + G_a(\alpha, \Delta t, \beta) \quad (2.3)$$

where α and β are shape and scale parameters that depend on the production speed (fast mode T_1 or slow mode T_2). The total production rate is:

$$P_t = \rho_t \cdot \sum_{i=1}^n p_i(d_i(t)) \quad (2.4)$$

with $\rho_t \sim \mathcal{N}(1, 0.1)$ representing scrap heterogeneity, and $p_i(d_i) \in [0, 1]$ the individual hammer efficiency.

The reward function is designed to minimize long-term costs and is normalized to the range $[-10, 0]$:

$$r_t = -c_u U_t - c_s B_t - \mathbb{I}_{a_t > 0} (c_i + c_r(a_t)) - \mathbb{I}_{CM}(c_f) \quad (2.5)$$

where U_t is unmet demand, $c_u = 100$ (cost per unit unmet), $c_s = 0.01$ (holding cost), $c_i = 50$ (fixed inspection cost), $c_r = 230$ (cost per hammer replacement), $c_f = 3000$ (catastrophic failure cost).

For clarity, Table 1 summarizes the parameters and their values used in the simulation.

Table 1: Parameters and values used in the simulation

Par	Description	Value
n	Number of hammers	20 components
h_1	Buffer level for decreasing production pace to T_2	16 ton
h_2	Buffer level for increasing production pace to T_1	4 ton
α_{T_1}	Shape parameter for T_1	1
β_{T_1}	Scale parameter for T_1	0.1
α_{T_2}	Shape parameter for T_2	0.5
β_{T_2}	Scale parameter for T_2	0.12
c_1	Production constant for T_1	1
c_2	Production constant for T_2	0.8
μ_{ρ_t}	Mean parameter for iron percentage	1 ton
σ_{ρ_t}	Standard deviation for iron percentage	0.1 ton
λ_{F_s}	Distribution parameter for CM duration in the second workstation	3 hours
μ_{d_t}	Mean parameter for demand	15
σ_{d_t}	Standard deviation for demand	1
λ_{F_f}	Distribution parameter for failure events at the second workstation	336 hours
λ_{F_a}	Distribution parameter for T_a	10 hours
c_r	Cost per hammer replacement	230 un./component
c_i	Cost per inspection	50 un.
c_s	Inventory carrying cost	0.01 un./ton.hours
c_u	Cost per unit of unmet demand	100 un./ton
c_f	Cost per failure event of the shredder	3000 un.
T_P	Time for inspection (fixed)	4 hours
T_u	Time per hammer replacement	1 hour
T_a	Additional time for CM	-
T_s	Second workstation time interval for PM	24 hours
T_d	PM duration for the second workstation	2 hours
K	Buffer capacity	1000 ton

2.2 Phase 1: Training a Continuous DDQN Agent

We train a Double Deep Q-Network (DDQN) agent directly on the continuous state space (5-dimensional input) with the full 21 action outputs. The architecture is a fully connected neural network with two hidden layers, each containing 64 neurons with ReLU activation. The output layer uses linear activation to produce Q-values for each of the 21 actions.

Key hyperparameters are:

- Replay buffer size: 10^5 transitions.
- Batch size: 128.
- Discount factor $\gamma = 0.91$.
- Learning rate: 2.24 (Adam optimizer).
- Target network update interval: every 100 steps.
- Exploration: ϵ -greedy with ϵ decaying from 1.0 to 0.5 at a decay rate of 0.3.

The agent is trained for 25 million time steps (approximately 25000 episodes with 1000 steps). The training curve shows convergence after about 300 episodes. However, this continuous DDQN model, while high-performing, is computationally expensive and operates as a black box.

2.3 Phase 2: Knowledge Extraction from the Trained Agent

Instead of deploying the continuous agent directly, we extract strategic knowledge from it. This extraction proceeds in two parallel analyses.

2.3.1 2.1 Sensitivity Analysis and Breakpoint Identification

After training, we fix the agent’s policy and compute the sensitivity of the Q-value function with respect to each continuous state variable. For a given state $\mathbf{s}$, the sensitivity is approximated by the mean absolute gradient over a local neighbourhood:

$$\text{Sensitivity}(s_j) = \mathbb{E}_{\mathbf{s} \sim \mathcal{D}} \left[\left| \frac{\partial Q(\mathbf{s}, a^*)}{\partial s_j} \right| \right] \quad (2.6)$$

where $a^* = \arg \max_a Q(\mathbf{s}, a)$.

We identify *breakpoints* — values of a state variable where the derivative changes abruptly, indicating a phase transition in the optimal policy. For example, the analysis reveals a sharp change in the optimal action when the number of defective hammers exceeds 5, another at 12, and again at 15. Similarly, the RUL variable shows critical thresholds at 400 minutes (6.5 hours) and 1200 minutes (20 hours) remaining.

These breakpoints are not chosen arbitrarily; they emerge directly from the agent’s learned behaviour and reflect the underlying physical and economic trade-offs of the system.

2.3.2 2.2 Action Pruning via Frequency Analysis

We run the trained continuous agent for 500 evaluation episodes (without exploration) and record the frequency of each of the 21 actions. The resulting histogram shows that the optimal policy uses only a small subset of actions. This indicates that the agent naturally discovers that only five distinct levels of maintenance intensity are strategically meaningful. Consequently, we prune the action space from 21 to 5 effective actions:

$$\mathcal{A}_{\text{pruned}} = \{0, 3, 5, 15, 20\} \quad (2.7)$$

2.4 Phase 3: Construction of the Informed Discrete State–Action Space

Using the breakpoints identified from sensitivity analysis, we discretize each continuous state variable into meaningful intervals. For example:

- Number of defective hammers (*Def*): intervals $[0, 5], [6, 12], [13, 15], [16, 20] \rightarrow 4$ levels.

- RUL: intervals $[0, 400]$, $[401, 1200]$, $[1201, 3000] \rightarrow 3$ levels.
- Production rate (P): intervals derived from the two operating modes (low: $[0, 5]$, high: $[6, 20]$) $\rightarrow 2$ levels.
- Buffer level (B): based on the critical thresholds $[0, 600]$, $[601, 1000] \rightarrow 2$ levels (low/medium/high buffer level can be used; we use 2 for simplicity after analysis).
- Demand (D): $[0, 5]$, $d_t > 5 \rightarrow 2$ levels.

The total number of discrete states is the product: $4 \times 3 \times 2 \times 2 \times 2 = 96$. However, not all combinations are reachable under the system dynamics. After removing physically impossible combinations, the theoretical informed state space contains 96 states, of which only approximately 57 are actually visited during simulation. The agent is trained on the full 96-state space, but the unreachable states do not affect learning.

Each state is now a tuple of discrete indices, and the action space is the pruned set of 5 actions. This discrete MDP is extremely compact and can be represented as a lookup table or a simple rule-based system.

2.5 Evaluation Protocol

To fairly compare the proposed informed model against baselines—specifically, (i) a continuous model with a 3-dimensional state space and binary actions (0: continue, 1: stop for PM), and (ii) a continuous model with a 4-dimensional state space (including RUL) and the same binary actions—we use a zero-shot evaluation setup. All policies are evaluated on the same simulation environment (identical stochastic seeds, 500 independent episodes). No additional learning occurs during evaluation; we use a fully greedy policy ($\epsilon = 0$).

Performance metrics are:

- Mean cumulative reward per episode (normalized to the range $[-10, 0]$ as during training).
- Standard deviation of cumulative reward (measuring operational stability).

The results, presented in the next section, demonstrate that the informed model achieves higher mean reward and significantly lower variance than both the 3-dimensional state space model and the 4-dimensional state space model, while being fully interpretable and computationally lightweight.

3 Results

We evaluate the proposed Agent-Informed Discretization framework against two baseline models: (i) the first baseline model is a continuous 3-dimensional state space with binary actions (continue/stop), and (ii) the second baseline model a 4-dimensional state space, which operates on the full continuous state space but with the same binary action space.

All models are compared in a zero-shot setting on the same stochastic simulation environment over 500 independent episodes, using a fully greedy policy ($\epsilon = 0$) to isolate policy quality from exploration noise.

Table 2 summarizes the quantitative results. Our final informed discrete model (96 states, 5 pruned actions) achieves the highest mean reward (-1275.21), outperforming the first baseline model (-1549.2) and the second one (-1534.6). More importantly, the standard deviation of the reward drops dramatically from 22.96 (first baseline model) and 16.11 (second baseline model) to only 12.67 in our model. This reduction in variance indicates that the informed policy is not only more cost-effective on average but also significantly more stable and predictable—a critical property for real-world industrial deployment where consistency matters as much as average performance. The following subsections analyze the learning behavior, policy structure, and sensitivity analyses that lead to these improvements.

Table 2: Comparison of model performance

Model	Mean Reward	Std Dev
3-dimensional state space model	-1549.2	22.96
4-dimensional state space model	-1534.6	16.11
5-dimensional state space model (smart)	-1275.21	12.67

3.1 Training Convergence of the Continuous DDQN Agent

Before extracting knowledge, we first trained a high-capacity DDQN agent directly on the continuous state space with 21 possible actions (replacing 0 to 20 hammers). Figure 1 shows the training progression over 25,000 episodes (each episode simulates up to 1000 time steps). The blue line represents the total undiscounted reward per episode, while the red line is a moving average to reveal the underlying trend.

- **Exploration and Rapid Improvement Phase (Episodes 0–3,000):** At the onset of training, the agent starts with a Total Reward of approximately $-12,500$. As the exploration rate decays, the moving Average (red line) exhibits a sharp upward trajectory, reaching a near-zero plateau by episode 3,000. This indicates that the agent quickly identifies a policy that avoids catastrophic failures and unnecessary maintenance costs.
- **Convergence and Policy Stability (Episodes 3,000–21,500):** For the remainder of the training, the moving average remains remarkably stable, hovering close to the zero-baseline (estimated between $-1,000$ and $-3,000$). This extended plateau over 18,000 episodes demonstrates that the continuous DDQN architecture has successfully converged to a robust and near-optimal maintenance policy.
- **Environmental Stochasticity and Extreme Outliers:** Despite the stable moving average, the **Episode Reward (light blue)** shows intermittent sharp dips throughout the training process. The most significant outlier occurs around episode

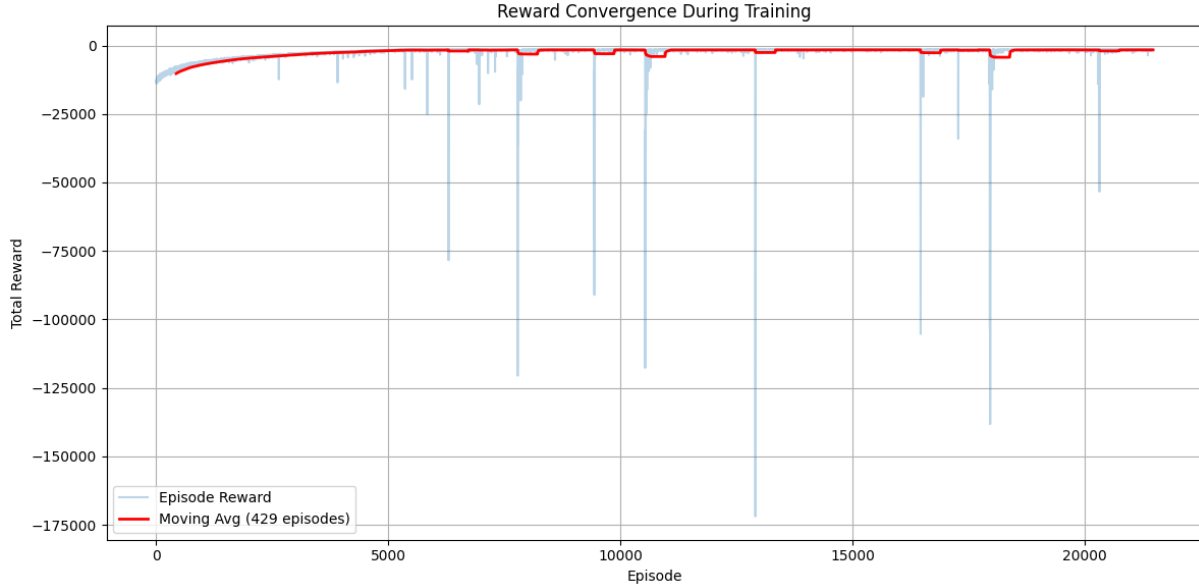


Figure 1: Training convergence of the continuous DDQN agent. The blue line shows episode reward; the red line shows moving average (429-episode window).

13,000, where the reward drops to approximately $-172,000$. These fluctuations are attributed to the inherent stochastic nature of the environment—such as sudden equipment failures or extreme variations in scrap quality—which can occasionally lead to high-cost episodes even under an optimized policy.

The training convergence confirms that the continuous DDQN architecture is capable of learning an effective maintenance policy in this complex environment. However, the computational cost (over 25,000 episodes, each requiring up to 1000 steps) and the black-box nature of the neural network motivate the need for knowledge extraction and simplification, which we present next.

3.2 Learned Policy of the Continuous DDQN Agent

After training convergence, we extract the deterministic greedy policy from the continuous DDQN agent by fixing $\epsilon = 0$ and recording the action with the highest Q-value for each representative state. Figure 2 visualizes this policy as a matrix (heatmap) where rows correspond to the number of defective hammers (grouped into intervals 0, 5, 10, 15, 20) and columns represent different operating conditions (primarily varying the remaining useful life, RUL, of the second workstation and the production rate). The entries indicate the selected action: 1 (replace 0 hammers, i.e., continue), 2 (replace 3 hammers), 3 (replace 5 hammers), 4 (replace 15 hammers), or 5 (replace 20 hammers, full overhaul).

The following sections detail the strategic behaviors developed by the agent across different operational regimes, illustrating how the policy adapts to production pressure and system health.

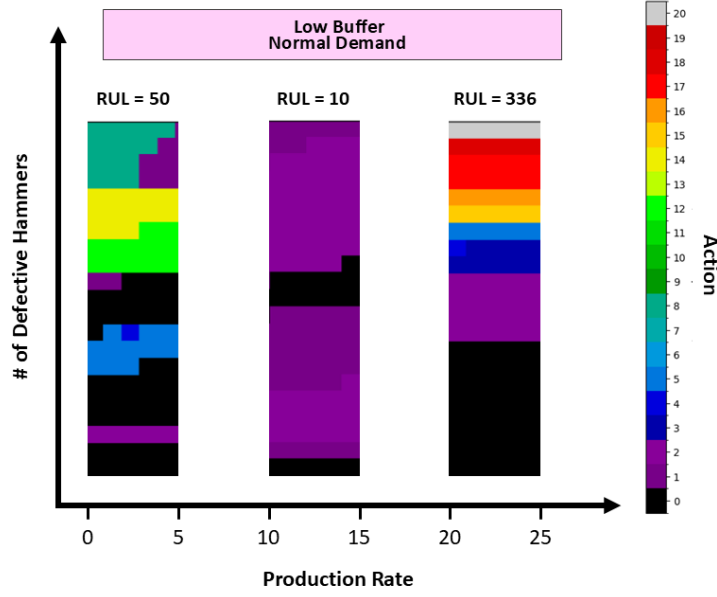


Figure 2: Learned greedy policy of the continuous DDQN agent after convergence. Each cell shows the optimal action (1–5) for a given number of defective hammers (row) and operating condition (column). The agent adapts maintenance intensity based on both wear level and prognostic information (RUL).

3.2.1 Low Production Regime (Production Rate: 0–5), Flexible Preventive Maintenance

In this regime, typically observed when $RUL = 50$, the system is in a "transition" state (neither brand new nor near failure).

- **Analysis:** At low production rates, the agent exhibits high sensitivity to the number of defective hammers. The high variance in action selection (represented by diverse color mapping in the policy plot) indicates the frequent use of *Intermediate Actions*.
- **Conclusion:** The agent learns to exploit low-pressure periods to perform minor replacements (e.g., 3 or 5 hammers). This improves production quality without incurring the heavy costs of a full system shutdown.

3.2.2 Medium Production Regime (Production Rate: 10–15), Risk Management and Survival Strategy

In this scenario, where $RUL = 10$ (indicating a critical health state), the agent's behavior shifts fundamentally.

- **Analysis:** Despite the moderate production demand, the severe degradation of the machine ($RUL = 10$) leads the agent to favor minimal intervention, represented by the dominance of dark-shaded regions (lowest replacement intensity).

- **Conclusion:** The agent correctly identifies that investing in hammer replacements is economically inefficient when the entire unit is on the verge of a major breakdown. The priority shifts from "tool quality" to *Sunk Cost Minimization*, as the workstation is expected to stop for a total overhaul shortly regardless of minor repairs.

3.2.3 High Production Regime (Production Rate: 20–25), Maximum Productivity and Overhaul Strategy

At a high $RUL = 336$ (representing a healthy downstream state), the policy demonstrates the highest level of decisiveness.

- **Analysis:** The policy mapping shows highly organized, horizontal stratification. As defective hammers increase, the agent aggressively transitions toward heavy interventions (e.g., replacing 15 to 20 hammers).
- **Conclusion:** When the machine is healthy and production demand is high, the agent adopts a *Maximum Reliability* strategy. It prevents the defect count from exceeding critical thresholds by performing major maintenance, thereby guaranteeing continuous high-rate production. Due to low buffer levels at high RUL, the agent becomes hyper-sensitive to tool quality to prevent any unscheduled stoppages.

3.3 Sensitivity Analysis and Policy Breakpoints

To understand the decision logic of the converged agent, we analyze the relationship between the *Max Q-value* and key state variables. These curves reveal strategic "breakpoints"—thresholds where the agent’s valuation of the state changes significantly, signaling a shift in the optimal policy.

3.3.1 Broken Hammers: A Multi-Stage Degradation Analysis

The analysis of the defective hammers variable (Figure 3) reveals four distinct levels of valuation that dictate the agent’s maintenance intensity:

- **Initial Wear [0, 5]:** The Q-value shows a minor initial dip but remains relatively high (above -9.0). The agent views the system as healthy, tolerating minor wear without urgent intervention.
- **Developing Damage [6, 10]:** The valuation stabilizes and begins a slight recovery. In this interval, the agent is likely monitoring the wear while prioritizing production, as the "point of no return" has not yet been reached.
- **Critical Transition [11, 15]:** This is the most volatile region. A sharp, precipitous drop occurs at approximately 13 hammers, where the Q-value plunges from -8.5 to -13.5 . This identifies 13 as the definitive "action threshold" where the risk of failure becomes economically unacceptable.

- **Saturation/Failure [16, 20]:** Beyond 15 hammers, the curve flatlines at its minimum valuation. The agent considers the workstation effectively failed or in a state of maximum inefficiency, necessitating a full overhaul.

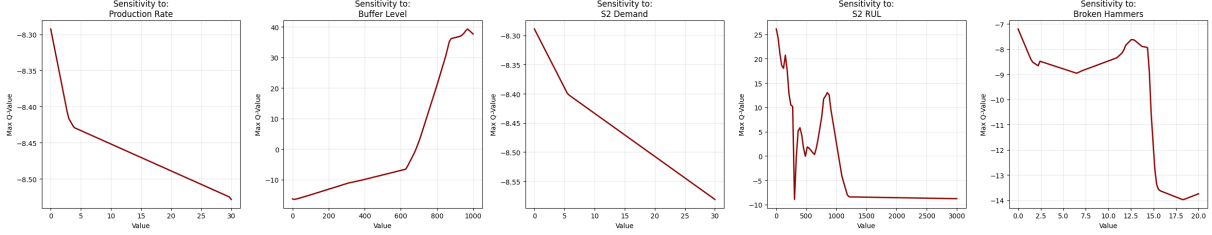


Figure 3: Sensitivity analysis of the Max Q-Value across key state variables: Production Rate, Buffer Level, Demand, S2 RUL, and Broken Hammers.

3.3.2 Buffer Level and Production Dynamics

The buffer level plot shows a massive positive correlation with the agent’s valuation, indicating its role as the system’s ”economic heartbeat”:

- The Q-value remains negative until the buffer reaches approximately 650 units.
- A steep surge in valuation occurs between 650 and 900 units, where the Q-value jumps from -5 to $+35$. This identifies 650 units as a strategic ”safety threshold” that the agent strives to maintain to ensure downstream continuity.

3.3.3 Remaining Useful Life (RUL)

For the RUL of the second workstation, the agent’s valuation is sensitive primarily to low-range health:

- Significant volatility and peaks are observed when the RUL is below 1,200 units, reflecting the agent’s sophisticated coordination of opportunistic maintenance to align with downstream health.
- Beyond 1,200 units, the Q-value flatlines, indicating that once the downstream workstation is sufficiently healthy, its exact RUL value no longer influences immediate maintenance decisions.

3.3.4 Production Rate and Demand

Both variables show a downward trend in valuation as values increase, though for different reasons. For the production rate, a notable inflection point occurs at 4 units. Below 4, the valuation drops sharply, suggesting the agent treats 4 units as the minimum viable rate required to avoid heavy system penalties.

3.3.5 Informed Discretization for Explainability

These data-driven breakpoints allow us to discretize the continuous state space into meaningful, interpretable intervals. Specifically:

- **Broken Hammers** is partitioned into four levels: *Healthy* [0–5], *Degraded* [6–10], *Critical* [11–15], and *Failed* [16–20].
- **RUL** is partitioned into three intervals based on the 1,200-unit stability threshold: *Urgent* [0–400], *Opportunistic* [401–1,200], and *Stable* [$>1,200$].

This ensures that the simplified model extracted from the neural network preserves the sophisticated economic logic learned during training.

3.4 Action Frequency Analysis and Pruning

After the continuous DDQN agent has converged, we evaluate its greedy policy ($\epsilon = 0$) over 500 independent episodes and record the frequency of each action (number of hammers replaced). Figure 4 presents the resulting histogram. The x-axis shows the action (0 to 20 hammers replaced), and the y-axis shows the total count of times that action was selected.

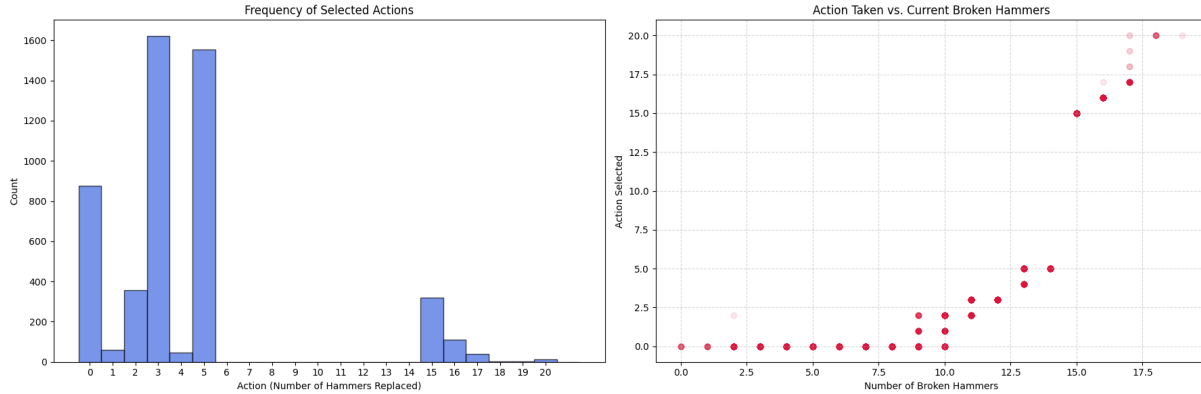


Figure 4: Frequency of each action selected by the converged continuous DDQN agent during 500 evaluation episodes. Only a small subset of actions (0, 3, 5, 15, 20) are used significantly; all intermediate actions have negligible counts.

3.4.1 Observed Action Distribution

The histogram reveals a highly strategic distribution. Contrary to simple policies, action $a = 3$ is the most frequent, selected over 1,600 times, followed closely by $a = 5$. This indicates that the agent’s primary strategy is *Proactive Maintenance*—performing frequent light repairs to maintain high system health. Action $a = 0$ (continue production) is the third most frequent, utilized when the defect count is below the critical threshold.

Heavy maintenance actions, specifically $a = 15$, appear when the system state degrades significantly. As shown in the scatter plot (Figure 4, right), there is a clear phase transition: the agent strictly avoids maintenance ($a = 0$) until approximately 9 broken

hammers, transitions to light repairs between 10–14, and triggers heavy interventions once defects exceed 15.

3.4.2 Negligible Actions and Pruning

Remarkably, the vast majority of the 21 possible actions have near-zero frequencies. The agent naturally ignores "middle-ground" actions (like 7 or 12 hammers) that do not align with the system’s economic breakpoints. This validates our action space pruning strategy. By restricting the final model to the five identified effective actions $\{0, 3, 5, 15, 20\}$ (or the specific clusters observed), we eliminate redundant decision noise while preserving the optimal strategic logic learned by the continuous DDQN.

3.4.3 Implication: Action Space Pruning

This observation directly motivates the action pruning step in our framework. The full action space of 21 possibilities contains many redundant or suboptimal actions that the optimal policy never uses. By pruning the action space to only the five effective actions $\{0, 3, 5, 15, 20\}$, we reduce the complexity of the decision problem without sacrificing performance. The pruned action space is used to construct the final informed discrete model, which retains the agent’s strategic behaviour while being compact and interpretable.

3.5 Training Convergence of the Informed Discrete Model

After constructing the informed discrete state space (96 states) and the pruned action space ($\{0, 3, 5, 15, 20\}$), we train a DDQN agent on this reduced MDP. Figure 5 shows the training progression over episodes. The blue line represents the raw undiscounted reward per episode, and the red line is a moving average that reveals the underlying trend.

Figure 5: Training convergence of the informed discrete Q-learning agent. The agent learns rapidly due to the compact state–action space and stabilizes with low variance.

Compared to the continuous DDQN training (Figure 1), the informed discrete model converges much faster—within approximately 2,000 episodes instead of 12,000. This speedup is a direct result of the knowledge extraction step: the state space has been reduced from a continuous 6-dimensional space to only 96 discrete states, and the action space from 21 to 5 strategic actions. The agent no longer needs to explore irrelevant actions or wasteful state regions.

The moving average shows a smooth, monotonic improvement without the extreme volatility observed during continuous training. After convergence, the episode reward stabilises around a value consistent with the zero-shot evaluation reported in Table 2 (approximately -1521 mean reward). The remaining small fluctuations reflect the inherent stochasticity of the environment (random scrap quality, demand, and failure timing), not poor policy quality.

This fast and stable convergence confirms that the informed discrete model retains the strategic optimality of the original continuous agent while being computationally lightweight and suitable for real-time industrial deployment.

3.6 Structural Comparison and Benchmarking

To evaluate the superiority of the proposed model, a zero-shot assessment was conducted. This benchmark compares the evolution of the models—from the initial thesis baseline to the continuous architecture and finally to the current smart discrete model—under a standardized reference environment. Table 3 summarizes the key structural differences between these approaches.

Table 3: Structural Comparison of RL Frameworks for Maintenance Optimization

	>p2.5cm >X >X >X			
Feature	Model-1	Model-2	Smart-Model (Proposed)	
State Space	Coarse	Discrete	Continuous	Informed Discrete
Action Space	2 Binary Actions (Repair / No Repair)		2 Binary Actions (Repair / No Repair)	
	Repair	5 Strategic Actions (Pruned Actions)		
Decision Method	DDQN	DDQN (Large Space)	Smart DDQN	
Computational Complexity	High	Very High	Balanced (Optimal)	

3.6.1 Robustness Evaluation

The zero-shot evaluation highlights how the models handle unseen environmental stochasticity:

- **Model-1:** While computationally efficient, its coarse state space leads to significant performance degradation in boundary conditions, as seen in the high variance of its rewards.
- **Model-2:** Offers high precision but suffers from extreme computational overhead and occasional instability due to the vast action space, evidenced by the sharp reward dips during training.
- **Proposed smart model:** By utilizing **Informed Discretization** based on strategic breakpoints (e.g., 13 broken hammers), the model achieves superior stability. It filters out irrelevant actions, resulting in a robust policy that consistently outperforms its predecessors in both mean reward and variance reduction.

3.7 Quantitative Comparison of Models

Figure 6 presents the zero-shot evaluation results for the three models. The y-axis shows the cumulative normalized reward per episode (higher is better), and the x-axis lists the three compared models: model-1 (3-dimensional state space model), model-2 (4-dimensional state space model), and smart mode; (our proposed 5-dimensional informed model with 5 strategic actions).

Figure 6: Zero-shot comparison of cumulative normalized reward. Model-2 achieves the highest mean reward and the smallest standard deviation, indicating both superior cost-effectiveness and operational stability.

3.7.1 Mean Reward

The informed model (smart-model) achieves the highest mean reward of -1275.21 , significantly outperforming model-2 (-1534.6) and model-1 (-1549.2). This improvement corresponds to a reduction in long-term operating costs of approximately **18%** relative to the baseline Model-1. This substantial gain demonstrates that the 5-state informed discretization captures critical system dynamics that coarser models fail to represent.

3.7.2 Standard Deviation and Stability

More importantly, the smart-model exhibits the smallest standard deviation, as evidenced by the significantly tighter box and whiskers in Figure ???. While model-1 shows high variance and extreme outliers (reaching as low as -1700), the Smart-Model’s rewards are tightly clustered between -1250 and -1310 . This reduction in variance is a critical advantage for real-world deployment: a policy that is both high-performing on average and highly predictable is essential for reliable production planning and workforce scheduling.

3.7.3 Interpretability vs. Performance Trade-off

Notably, the smart-Model achieves superior performance despite using a simplified representation compared to continuous architectures. This confirms that the knowledge extraction process does not simply compress the policy—it refines it. By focusing on the 5 strategic state-action breakpoints, the model eliminates the “decision noise” present in model-1 and model-2. The result is a policy that is not only fully interpretable as a set of logical rules but also objectively superior in both cost-effectiveness and operational robustness.

References

- [1] Ferreira Neto, W. A., Cavalcante, C. A. V., & Do, P. (2024). Deep reinforcement learning for maintenance optimization of a scrap-based steel production line. *Reliability Engineering & System Safety*, 249, 110199.
- [2] Andersen, J. F., & Nielsen, B. F. (2025). A comparative study of time-based maintenance and condition-based maintenance for multi-component systems. *Reliability Engineering & System Safety*, 256, 110759.
- [3] Sutton, R. S., & Barto, A. G. (2018). *Reinforcement Learning: An Introduction* (2nd ed.). MIT Press.
- [4] Hao, S., Zheng, J., Yang, J., Sun, H., Zhang, Q., Zhang, L., Jiang, N., & Li, Y. (2023). Deep reinforcement learning for joint optimization of condition-based maintenance and spare ordering. *Information Sciences*, 634, 85–100.
- [5] Skordilis, E., & Moghaddass, R. (2020). A deep reinforcement learning approach for real-time sensor-driven decision making and predictive analytics. *Computers & Industrial Engineering*, 147, 106600.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Multi-State Stress-Strength Reliability Models

Hajialiakbar, S. ¹ Tavangar, M. ²

Department of Statistics, Faculty of Mathematics and Statistics, University of Isfahan,
Isfahan 81746-73441, Iran

Abstract

A system comprising n independent components (or subsystems) is considered. The stress-strength Reliability (SSR) of the system is analytically evaluated, and the states of the system are classified considering the number of components whose strengths above (below) the multiple stresses available in an environment. That is, consideration of m stresses create $m + 1$ states for the system and its components. In this study, multi-state systems, in particular, multi-state k -out-of- $n:G$ and multi-state consecutive k -out-of- $n:G$, are considered in a stress-strength setup. Results are illustrated for various stress-strength distributions. Maximum likelihood estimators of state probabilities are also presented. It is observed that ignoring various states of components leads to significant bias in reliability estimation. Finally, key conclusions and generalizations are discussed.

Keywords: Exponential distribution, Maximum likelihood estimate, Multi-state k -out-of- $n:G$ system, Order statistics, Pareto distribution, Stress-strength reliability.

¹Speaker. s.hajialiakbar@sci.ui.ac.ir

²m.tavangar@sci.ui.ac.ir

1 Introduction

One of the most widely used models in reliability analysis is the stress-strength model, which was first introduced by Birnbaum [1]. This model is one of the basic approaches in reliability analysis of engineering components and systems, in which the system is functioning if and only if its strength Y exceeds the applied stress X . In the classical stress-strength model, reliability is defined as $R = P(X < Y)$, where X and Y are the stress and strength variables of the component, respectively. Given the importance of coherent systems in various industries and their role in improving the efficiency and safety of engineering systems, it is necessary to develop stress-strength models within the framework of multi-state systems. Therefore, the considering of stress-strength model in a setup of multi-state system framework enables us to make more accurate and sensitive inferences about the system.

In most previous research works, only two possible states for the system and its components have been considered: either working (if $Y_i > Z$) or failed (if $Y_i \leq Z$). While many engineering systems in practice have multiple operating states in which the system and its components can operate at multiple levels of efficiency. Estimating reliability of such systems may be incorrect if various states are ignored. Lisnianski et al. [8] studied a power supply system consisting of generating units, in which the system and its components may be modeled as multi-state. Eryilmaz [4] studied the SSR under multi-state system model by using the size of $X - Y$ to define the state of the multi-state system, and evaluated system state probabilities under various assumptions on the system. Qin et al. [10] used the ratio X/Y of the strength and stress values to define the state of multi-state system and presented the estimation of the multi-state SSR function in detail when the stress and strength variables followed independent exponential distributions. Eryilmaz and Tuncel [6] addressed multi-state system with multi-state components in particular.

More recently, Deveci and Iscioglu [3] investigated the mean remaining strength of multi-state components and multi-state systems under a stress-strength framework. Furthermore, Bai et al. [1] extended the multi-state stress-strength model to systems with multiple component types and dependent stress-strength variables through copula-based modeling.

In the present study, we develop a multi-state stress-strength reliability model for coherent systems under two common stresses. The proposed framework allows both the components and the system to operate in three distinct performance states, providing a more informative reliability assessment than the classical binary stress-strength model. Suppose that a system consists of n independent components subjected to two common stresses Z_1 (temperature) and Z_2 (friction).

The rest of this paper is organized as follows. In Section 2, the system used in the paper is described. The state distributions associated with the system are presented in section 3. Illustrations and Numerical Results are shown in section 4. Eventually, Discussion and Generalizations are stated in section 4.

2 Model Description

Consider a system consisting of n independent components whose strengths $Y_1, Y_2, \dots, Y_n$ which represent the strength of the components are independent identically distributed (i.i.d) with continuous cumulative distribution function (cdf) $F(x) = P\{Y_i \leq x\}, i = 1, 2, \dots, n$. Suppose these components are subjected to two types of stresses Z_1 and Z_2 over time. The states 0, 1 and 2 represent completely failure, partially working, and completely working states, respectively and a component is assigned to be “completely failure” when the strength of the component is below both of the stresses, “partially working” if the strength of the component is above only one stress, and “completely working” when the strength of the component is above both of the stresses. Let X_i denote the state of the i th component. Define

$$X_i = \begin{cases} 0 & \text{if } Y_i < Z_{1:2} \\ 1 & \text{if } Z_{1:2} < Y_i < Z_{2:2}, \\ 2 & \text{if } Y_i > Z_{2:2} \end{cases} \quad i = 1, 2, \dots, n \quad (2.1)$$

where $Z_{1:2} = \min(Z_1, Z_2)$ and $Z_{2:2} = \max(Z_1, Z_2)$.

The random variables $X_1, \dots, X_n$ representing the states of the components are not independent due to common random stresses.

The specification and determination of system states generally depends on engineering and system considerations. Various structures for system states may appear based on the conditions in the operating environment. We consider a model where the system is in state j ($j = 0, 1, 2$) or above if and only if at least k components are in state j or above. This system is known as a simple multi-state k -out-of- n : G system and in this case the system and its components are both multi-state. We will use ϕ to denote the structure function of this system. Another multi-state system called simple multi-state consecutive k -out-of- n : G system is in state j ($j = 0, 1, 2$) or above if and only if at least k consecutive components are in state j or above. We use ϕ_c to denote the structure function of this system. For more details about the multi-state k -out-of- n : G systems, the interested reader may refer to Kuo and Zuo [7].

3 State Distributions

Considering the distributions of $Z_{1:2}$ and $Z_{2:2}$. we initially assume stresses Z_1 and Z_2 to be independent random variables with continuous cdfs G_1 and G_2 respectively. Then we have

$$P\{Z_{1:2} \leq z\} = 1 - \bar{G}_1(z)\bar{G}_2(z),$$

where $\bar{G}(z) = 1 - G(z)$ Similarly,

$$P\{Z_{2:2} \leq z\} = G_1(z)G_2(z),$$

and

$$P\{Z_{1:2} \leq z_1, Z_{2:2} \leq z_2\} = G_1(z_2)G_2(z_2) - (G_1(z_2) - G_1(z_1))(G_2(z_2) - G_2(z_1)),$$

For $z_1 < z_2$.

Now, Let $p_{i,j}$ denote the probability that component i is in state j and $P_{i,j}$ denote the probability that component i is in state j or above, $1 \leq i \leq n, 0 \leq j \leq 2$. Then:

$$\begin{aligned} p_{i,0} &= P\{X_i = 0\} = \int F(z)(g_2(z)\bar{G}_1(z) + g_1(z)\bar{G}_2(z)) \, dz, \\ p_{i,1} &= P\{X_i = 1\} = \iint_{z_1 < z_2} (F(z_2) - F(z_1))(g_1(z_2)g_2(z_1) + g_1(z_1)g_2(z_2)) \, dz_1 \, dz_2, \\ p_{i,2} &= P\{X_i = 2\} = \int \bar{F}(z)(g_1(z)G_2(z) + g_2(z)G_1(z)) \, dz. \end{aligned}$$

As a special case if Z_1 and Z_2 are i.i.d with common cdf G and pdf g , then we have

$$\begin{aligned} p_{i,0} &= 2 \int F(z)\bar{G}(z)g(z) \, dz, \\ p_{i,1} &= 2 \iint_{z_1 < z_2} (F(z_2) - F(z_1))g(z_1)g(z_2) \, dz_1 \, dz_2, \\ p_{i,2} &= 2 \int \bar{F}(z)G(z)g(z) \, dz, \quad i = 1, 2, \dots, n. \end{aligned}$$

Clearly

$$P_{i,0} = 1, \quad P_{i,1} = p_{i,1} + p_{i,2}, \quad P_{i,2} = p_{i,2}.$$

In the following, we present the state distributions for multi-state k -out-of- n : G and multi-state consecutive k -out-of- n : G systems, where the components states are denoted by $X_1, X_2, \dots, X_n$.

Theorem 3.1. Let $R_j^{k,n}$ denote the probability that the multi-state k -out-of- n : G system is in state j or above, i.e., $R_j^{k,n} = P\{\phi(X) \geq j\}$. Then:

- (a) $R_0^{k,n} = 1$;
- (b) $R_1^{k,n} = \sum_{m=k}^n \binom{n}{m} E((F(Z_{1:2}))^{n-m} (\bar{F}(Z_{1:2}))^m)$;
- (c) $R_2^{k,n} = \sum_{m=k}^n \binom{n}{m} E((F(Z_{2:2}))^{n-m} (\bar{F}(Z_{2:2}))^m)$.

The probability that the system is in state j can be derived using

$$r_j^{k,n} = P\{\phi(X) = j\} = R_j^{k,n} - R_{j+1}^{k,n}.$$

Proof. The proof can be found in [2]. □

Theorem 3.2. Let $R_{j,c}^{k,n}$ denote the probability that the multi-state consecutive k -out-of- n : G system is in state j or above. Then:

- (a) $R_{0,c}^{k,n} = P\{\phi_c(X) \geq 0\} = 1$;
- (b) $R_{1,c}^{k,n} = 1 - \sum_{l=0}^n \sum_{j=0}^a \sum_{i=0}^{n-l} (-1)^j (-1)^i \binom{n-l+1}{j} \binom{n-kj}{n-l} \binom{n-l}{i} E(\bar{F}(Z_{1:2}))^{l+i}$;

$$(c) R_{2:c}^{k,n} = 1 - \sum_{l=0}^n \sum_{j=0}^a \sum_{i=0}^{n-l} (-1)^j (-1)^i \binom{n-l+1}{j} \binom{n-kj}{n-l} \binom{n-l}{i} E(\bar{F}(Z_{2:2}))^{l+i},$$

where $a = \min\left(\left\lfloor \frac{1}{k} \right\rfloor, n-l+1\right)$ and $[x]$ denotes the integer part of x .

Proof. The proof can be found in [2]. □

Now consider the case where the stresses are dependent, i.e., the random variables Z_1 and Z_2 have the joint cdf $G(z_1, z_2)$ and pdf $g(z_1, z_2)$. In this case, the distributions associated with $Z_{1:2}$ and $Z_{2:2}$ are obtained as follows:

$$P\{Z_{1:2} \leq z\} = 1 - \bar{G}(z, z) = G_1(z)G_2(z) - G(z, z), \quad (3.1)$$

$$P\{Z_{2:2} \leq z\} = G(z, z), \quad (3.2)$$

$$P\{Z_{1:2} \leq z_1, Z_{2:2} \leq z_2\} = G(z_1, z_2) + G(z_2, z_1) - G(z_1, z_1), \quad (3.3)$$

where $\bar{G}(z_1, z_2) = P\{Z_1 > z_1, Z_2 > z_2\}$ represents the bivariate survival function.

4 Illustrations and Numerical Results

In this section, the state distributions of the multi-state systems introduced above are computed and estimated under both independent and dependent stresses using parametric statistical models.

4.1 Independent Stresses

Let Y_i s and Z_1, Z_2 are all independent exponential variates with cdfs

$$F(x) = 1 - e^{-\lambda x}, \quad G_i(x) = 1 - e^{-\theta_i x}; \quad i = 1, 2, \quad x \geq 0.$$

Then, using the formulas used in the previous section,

$$\begin{aligned} p_{i,0} &= \frac{\lambda}{\lambda + \theta_1 + \theta_2} \\ p_{i,1} &= \frac{\theta_1 + \theta_2}{\lambda + \theta_1 + \theta_2} - \frac{\theta_1 \theta_2}{(\lambda + \theta_1 + \theta_2)(\lambda + \theta_2)} - \frac{\theta_1 \theta_2}{(\lambda + \theta_1 + \theta_2)(\lambda + \theta_1)} \\ p_{i,2} &= \frac{\theta_1}{\lambda + \theta_1} + \frac{\theta_2}{\lambda + \theta_2} - \frac{\theta_1 + \theta_2}{\lambda + \theta_1 + \theta_2}. \end{aligned}$$

The system state probabilities for a multi-state k -out-of- n : G system can also be calculated as

$$\begin{aligned} R_1^{k,n} &= \sum_{m=k}^n \sum_{i=0}^{n-m} (-1)^i \binom{n}{m} \binom{n-m}{i} \frac{\theta_1 + \theta_2}{(m+i)\lambda + \theta_1 + \theta_2}, \\ R_2^{k,n} &= \sum_{m=k}^n \sum_{i=0}^{n-m} (-1)^i \binom{n}{m} \binom{n-m}{i} \left[\frac{\theta_1}{(m+i)\lambda + \theta_1} + \frac{\theta_2}{(m+i)\lambda + \theta_2} - \frac{\theta_1 + \theta_2}{(m+i)\lambda + \theta_1 + \theta_2} \right]. \end{aligned}$$

Tables 1 and 2 contain some numerical results for 2-out-of-3 and 3-out-of-5 systems for selected values of the parameters λ, θ_1 and θ_2 .

Table 1: State distribution of multi-state 2-out-of-3: G system

λ	θ_1	θ_2	$R_1^{2,3}$	$R_2^{2,3}$	$r_0^{2,3}$	$r_1^{2,3}$	$r_2^{2,3}$
0.1	0.2	0.3	0.893	0.607	0.107	0.286	0.607
0.1	0.2	0.5	0.933	0.660	0.067	0.273	0.660
0.5	0.2	0.3	0.500	0.124	0.500	0.376	0.124
0.5	0.8	0.3	0.724	0.271	0.276	0.453	0.271

Table 2: State distribution of multi-state 3-out-of-5: G system

λ	θ_1	θ_2	$R_1^{3,5}$	$R_2^{3,5}$	$r_0^{3,5}$	$r_1^{3,5}$	$r_2^{3,5}$
0.1	0.2	0.3	0.917	0.619	0.083	0.298	0.619
0.1	0.2	0.5	0.955	0.676	0.045	0.279	0.676
0.5	0.2	0.3	0.500	0.110	0.500	0.390	0.110
0.5	0.8	0.3	0.742	0.259	0.258	0.483	0.259

In the following, we compute system state probabilities for the multi-state consecutive k -out-of- n : G system. The expected values $E(\bar{F}(Z_{1:2}))^m$ and $E(\bar{F}(Z_{2:2}))^m$ are enough for computing the system state probabilities.

$$\begin{aligned}
E(\bar{F}(Z_{1:2}))^m &= P\{Y_1 > Z_{1:2}, \dots, Y_m > Z_{1:2}\} \\
&= \int (\bar{F}(z))^m dP\{Z_{1:2} \leq z\} \\
&= \int_0^\infty (e^{-\lambda z})^m (\theta_2 e^{-\theta_2 z} e^{-\theta_1 z} + \theta_1 e^{-\theta_1 z} e^{-\theta_2 z}) dz \\
&= \frac{\theta_1 + \theta_2}{\theta_1 + \theta_2 + m\lambda},
\end{aligned} \tag{4.1}$$

and

$$\begin{aligned}
E(\bar{F}(Z_{2:2}))^m &= P\{Y_1 > Z_{2:2}, \dots, Y_m > Z_{2:2}\} \\
&= \int (\bar{F}(z))^m dP\{Z_{2:2} \leq z\} \\
&= \int_0^\infty (e^{-\lambda z})^m (\theta_1 e^{-\theta_1 z} (1 - e^{-\theta_2 z}) + (1 - e^{-\theta_1 z}) \theta_2 e^{-\theta_2 z}) dz \\
&= \frac{\theta_1 \theta_2}{\theta_1 + \theta_2 + m\lambda} \left[\frac{1}{\theta_1 + m\lambda} + \frac{1}{\theta_2 + m\lambda} \right],
\end{aligned} \tag{4.2}$$

Tables 3 and 4 contain some numerical results for 2-out-of-3 and 3-out-of-5 consecutive systems for the same values of the parameters used in Tables 1 and 2. When the parameters λ, θ_1 and θ_2 are known then the system state probabilities are simply computed as above, otherwise we have to estimate them. Undoubtedly, the most flexible and widely used method for estimating reliability is the maximum likelihood method. Because the maximum likelihood estimator (mle) of a function of parameters is the same as the mle function of the parameters. Let $(Y_1, \dots, Y_m), (Z_1^{(1)}, \dots, Z_{m_1}^{(1)})$ and $(Z_1^{(2)}, \dots, Z_{m_2}^{(2)})$ be independent random samples from exponential distributions with means $\frac{1}{\lambda}, \frac{1}{\theta_1}$, and $\frac{1}{\theta_2}$,

Table 3: State distribution of multi-state consecutive 2-out-of-3: G system

λ	θ_1	θ_2	$R_{1,c}^{2,3}$	$R_{2,c}^{2,3}$	$r_{0,c}^{2,3}$	$r_{1,c}^{2,3}$	$r_{2,c}^{2,3}$
0.1	0.2	0.3	0.804	0.496	0.196	0.308	0.496
0.1	0.2	0.5	0.856	0.548	0.144	0.308	0.548
0.5	0.2	0.3	0.417	0.094	0.583	0.323	0.094
0.5	0.8	0.3	0.625	0.211	0.375	0.414	0.211

Table 4: State distribution of multi-state consecutive 3-out-of-5: G system

λ	θ_1	θ_2	$R_{1,c}^{3,5}$	$R_{2,c}^{3,5}$	$r_{0,c}^{3,5}$	$r_{1,c}^{3,5}$	$r_{2,c}^{3,5}$
0.1	0.2	0.3	0.764	0.412	0.236	0.352	0.412
0.1	0.2	0.5	0.827	0.470	0.173	0.357	0.470
0.5	0.2	0.3	0.350	0.060	0.650	0.290	0.060
0.5	0.8	0.3	0.560	0.152	0.440	0.408	0.152

respectively. Then the mle of $(\lambda, \theta_1, \theta_2)$ is given by

$$\tilde{\lambda} = 1/\bar{Y}_m, \quad \tilde{\theta}_1 = \frac{1}{\bar{Z}_{m_1}^{(1)}} \quad \text{and} \quad \tilde{\theta}_2 = \frac{1}{\bar{Z}_{m_2}^{(2)}},$$

where $\bar{Y}_m$, $\bar{Z}_{m_1}^{(1)}$, and $\bar{Z}_{m_2}^{(2)}$ represent the corresponding sample means.

4.2 Dependent Stresses

Let (Z_1, Z_2) has a bivariate Pareto distribution with survival function and pdf given, respectively, by

$$\begin{aligned} \bar{G}(z_1, z_2) &= (z_1 + z_2 - 1)^{-\theta}, \\ g(z_1, z_2) &= \theta(\theta + 1)(z_1 + z_2 - 1)^{-(\theta+2)}, \quad \theta > 0, \quad z_1 > 1, \quad z_2 > 1. \end{aligned}$$

Then, using (2.2) and (2.3), we have:

$$\begin{aligned} P\{Z_{1:2} \leq z\} &= 1 - (2z - 1)^{-\theta}, \\ P\{Z_{2:2} \leq z\} &= 1 - 2z^{-\theta} + (2z - 1)^{-\theta}. \end{aligned}$$

Let the strengths of the components has a Pareto distribution with survival function

$$\bar{F}(x) = x^{-\lambda}, \quad \lambda > 0, \quad x > 1.$$

Then using parts (b) and (c) of Theorem 3.1 we find the probabilities of the k -out-of- n : G system being in states “1” and “2”, respectively, as

$$R_1^{k,n} = 2\theta \sum_{m=k}^n \binom{n}{m} \int_1^\infty z^{-\lambda m} (1 - z^{-\lambda})^{n-m} (2z - 1)^{-(\theta+1)} dz$$

and

$$R_2^{k,n} = 2\theta \sum_{m=k}^n \binom{n}{m} \int_1^\infty z^{-\lambda m} (1 - z^{-\lambda})^{n-m} (z^{-(\theta+1)} - (2z - 1)^{-(\theta+1)}) dz.$$

These probabilities can be easily calculated if the parameters θ and λ are known as in Table 5. When the parameter values are unknown, then we can use their mles to estimate state probabilities. For this, let $(Y_1, \dots, Y_m)$ and $(Z_1^{(1)}, Z_1^{(2)}), \dots, (Z_{m_1}^{(1)}, Z_{m_1}^{(2)})$ be independent random samples from the distributions of Y and (Z_1, Z_2) . Then the mle of (λ, θ) is given by

$$\tilde{\lambda} = \frac{m}{\sum_{i=1}^m \ln Y_i}, \quad \tilde{\theta} = (C^{-1} - \frac{1}{2}) + (C^{-2} + \frac{1}{4})^{\frac{1}{2}},$$

Where $C = m_1^{-1} \sum_{i=1}^{m_1} \ln(Z_i^{(1)} + Z_i^{(2)} - 1)$; see, e.g., Mardia [9].

Table 5: State distribution of multi-state k -out-of- n : G system for dependent stresses

n	k	θ	λ	$R_1^{k,n}$	$R_2^{k,n}$	$r_0^{k,n}$	$r_1^{k,n}$	$r_2^{k,n}$
10	7	2	2	0.543	0.185	0.457	0.358	0.185
10	7	3	2	0.677	0.290	0.323	0.387	0.290
15	7	3	2	0.864	0.541	0.136	0.323	0.541
15	10	3	2	0.699	0.300	0.301	0.399	0.300
15	10	2	3	0.438	0.106	0.562	0.332	0.106

5 Discussion and Generalizations

In this article, we considered multi-state k -out-of- n : G systems and multi-state consecutive k -out-of- n : G systems and studied stress-strength reliability. Under the proposed model, the states of the system are assigned considering the number of components whose strengths above (below) the multiple stresses available in an environment. That is, consideration of m stresses create $m+1$ states for the system and its components. In this study, we considered the case of two stresses $m = 2$. Extending the proposed framework to the general case of multiple stresses $m > 2$ is a promising direction for future research. In this case, the state of the i th component can be defined in terms of order statistics as $X_i = r$ iff $Z_{r:m} < Y_i < Z_{r+1:m}$, $i = 1, \dots, n$ and $r = 0, 1, \dots, m$, where $Z_{0:m} \equiv 0$, $Z_{m+1:m} \equiv \infty$.

We considered the simple multi-state k -out-of- n : G systems and multi-state consecutive k -out-of- n : G systems whose components' states are $X_1, X_2, \dots, X_n$. These multi-state system were chosen for two main reasons. First, these structures are widely used in engineering applications and second they generalize multi-state series and parallel systems. Future research may focus on extending the proposed framework to multiple-stress $m > 2$ settings as well as more general classes of multi-state coherent systems.

References

- [1] Bai, X., Wang, Y., Zhang, C., Cai, J., Zhou, H. (2025), Stress-strength reliability inference of multi-state system with multi-type components based on copula theory, *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **239**(6), 1556–1567.

- [2] Birnbaum, Z. (1956), On a use of the Mann-Whitney statistics, *in: Proceedings of the Third Berkeley Symposium in Mathematics, Statistics and Probability*, **1**, 13-17.
- [3] Deveci, F., Iscioglu, F. (2024), The mean remaining strength evaluation of a multi-state system under a stress-strength model, *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **238**(6), 1118–1135.
- [4] Eryilmaz, S., İşçioğlu, F. (2011), Reliability evaluation for a multi-state system under stress-strength setup, *Communications in Statistics - Theory and Methods*, **40** (3), 547-558.
- [5] Eryilmaz, S. (2011), A new perspective to stress-strength models, *Annals of the Institute of Statistical Mathematics*, **63** (1), 101-115.
- [6] Eryilmaz, S., Tuncel, A. (2016), *Generalizing the survival signature to unrepairable homogeneous multi-state systems*, *Naval Research Logistics*, **63** (8), 593-599.
- [7] Kuo, W., Zuo, M. J. (2003), *Optimal Reliability Modeling, Principles and Applications*, New York: John Wiley & Sons.
- [8] Lisnianski, A., Frenkel, L., Ding, Y. (2010), *Multi-state system reliability analysis and optimization for engineers and industrial managers*, Springer Science & Business Media.
- [9] Mardia, K. V. (1962), Multivariate Pareto distributions, *Annals of Mathematical Statistics*, **33**, 1008–1015.
- [10] Qin, H., Jana, N., Kumar, S., Chatterjee, K. (2017), Stress-strength models with more than two states under exponential distribution, *Communications in Statistics-Theory and Methods*, **46** (1), 120-132.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



A Statistical Model for Real Lifespan Under Environmental Conditions: A Gamma-Dirichlet Approach

Homei, H.¹ Hedayati, A.² Kamali Alamdari, M.³

^{1,3} Department of Mathematics, Statistics and Computer Science, Faculty of Mathematical Sciences, University of Tabriz, Tabriz, Iran

² Department of Statistics, Mathematics and Computer Science, Faculty of Mathematical Sciences, Allameh Tabataba'i University, Tehran, Iran

Abstract

This paper introduces a novel statistical framework for modeling the actual lifespan of components under real-world environmental conditions. Unlike traditional laboratory-based assessments, our model incorporates environmental variability by defining the actual lifespan as the product of a Gamma-distributed laboratory lifetime and a Dirichlet-distributed environmental impact vector. We derive the joint distribution of this product and explore its fundamental properties, including moments and covariance structure. The model is further extended to scenarios involving multiple independent positions, where the aggregate lifespan is expressed as a sum of such products. For cases where exact distributions are intractable, we propose an approximation method based on moment matching. Numerical simulations and a real-data application using field records from an Iranian tractor manufacturer illustrate the model's accuracy and practical relevance. The proposed framework offers a flexible tool for reliability analysis in fields such as civil engineering and manufacturing.

¹homei@tabrizu.ac.ir

²Poster. Alireza.heda0914@gmail.com

³Maryam.Kamali@gmail.com

Keywords: Environmental lifetime, Gamma distribution, Dirichlet distribution, Product of random variables, Approximation

1 Introduction

The assessment of product lifespan under standard laboratory conditions often fails to capture the complexities of real-world environments. While traditional reliability analyses provide a useful baseline, factors such as temperature fluctuations, mechanical stress, and material degradation significantly influence actual longevity. This disparity between idealized testing and practical performance has motivated the development of more sophisticated statistical models that account for environmental variability [16].

In this paper, we propose a framework where the "real lifetime" of a component is conceptualized as the product of two distinct elements: a laboratory-measured baseline lifetime and an environmental coefficient vector that modulates this baseline according to external conditions. This approach extends the existing literature on generalized distributions [3, 14] and addresses questions regarding the approximation of multivariate random variables raised in recent studies [5].

A key contribution of our work is the development of an approximation technique capable of handling multivariate cases, in contrast to previous univariate methods [13]. Section 5 includes numerical simulations and a real-data analysis based on field observations from the Tabriz Tractor Manufacturing Company. The paper is structured as follows: Section 2 introduces fundamental definitions and assumptions. Section 3 presents the main distributional properties of the proposed model. Section 4 extends the framework to multiple independent positions. Section 5 provides numerical comparisons, evaluates the proposed approximation method, and presents a real-data application.

2 Preliminaries and Definitions

Traditional approaches typically treat lifetime as a univariate random variable measured under controlled conditions. However, real-world applications require a more nuanced perspective. Let Y denote the *laboratory lifetime*, a non-negative random variable representing the lifespan under ideal conditions. To incorporate environmental effects, we introduce the *environmental coefficient vector* $\mathbf{X} = (X_1, X_2, \dots, X_n)$, where each component $X_i \geq 0$ represents the relative environmental impact on different aspects of the component, satisfying $\sum_{i=1}^n X_i = 1$.

Definition 2.1 (Real Lifetime). The *real lifetime* $\mathbf{Z}$ is defined as the component-wise product of the environmental coefficient vector and the laboratory lifetime:

$$\mathbf{Z} = \mathbf{X}Y = (X_1Y, X_2Y, \dots, X_nY). \quad (2.1)$$

This formulation scales the ideal laboratory lifetime by environmental factors, producing a vector whose components represent lifespans under different environmental conditions.

For applications involving multiple stages or locations, we consider an aggregate measure.

Definition 2.2 (Aggregate Real Lifetime). Let $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_r$ be independent environmental coefficient vectors and $Y_1, Y_2, \dots, Y_r$ be independent laboratory lifetimes, with independence between the sequences. The *aggregate real lifetime* across r independent positions is:

$$\mathbf{T} = \sum_{i=1}^r \mathbf{X}_i Y_i. \quad (2.2)$$

Throughout this paper, we assume parametric forms that are both flexible and mathematically tractable. Specifically, we assume that the environmental coefficient vector $\mathbf{X}$ follows a Dirichlet distribution with parameters $(\alpha_1, \alpha_2, \dots, \alpha_k)$, denoted as $\mathbf{X} \sim \text{Dir}(\alpha_1, \dots, \alpha_k)$, and that the laboratory lifetime Y follows a Gamma distribution with shape parameter α and scale parameter β , denoted as $Y \sim \text{Gamma}(\alpha, \beta)$. The independence of $\mathbf{X}$ and Y is a fundamental assumption throughout our analysis.

3 Distribution of the Real Lifetime

The joint distribution of the real lifetime vector $\mathbf{Z} = \mathbf{X}Y$ can be derived through a standard transformation technique. The following theorem presents the closed-form density.

Theorem 3.1. *Let $\mathbf{X} \sim \text{Dir}(\alpha_1, \alpha_2, \dots, \alpha_r)$ and $Y \sim \text{Gamma}(\alpha, \beta)$ be independent. Then the joint probability density function of $\mathbf{Z} = (Z_1, Z_2, \dots, Z_r)$ is given by:*

$$f(z_1, \dots, z_r) = \frac{\beta^{-\alpha} \Gamma(\sum_{i=1}^r \alpha_i)}{\Gamma(\alpha) \prod_{i=1}^r \Gamma(\alpha_i)} \left(\sum_{i=1}^r z_i \right)^{\alpha - \sum_{i=1}^r \alpha_i} \exp \left\{ -\frac{\sum_{i=1}^r z_i}{\beta} \right\} \prod_{i=1}^r z_i^{\alpha_i - 1}, \quad (3.1)$$

where $z_i > 0$ for $i = 1, 2, \dots, r$.

Proof. Consider the transformation $y = \sum_{i=1}^r z_i$ and $x_i = z_i / \sum_{i=1}^r z_i$ for $i = 1, \dots, r-1$. The Jacobian of this transformation is $|J| = (\sum_{i=1}^r z_i)^{-(r-1)}$. The joint density of $\mathbf{Z}$ is obtained by substituting the Gamma density of Y and the Dirichlet density of $\mathbf{X}$ into the standard transformation formula:

$$f_{\mathbf{Z}}(z_1, \dots, z_r) = f_Y \left(\sum_{i=1}^r z_i \right) \cdot f_{\mathbf{X}} \left(\frac{z_1}{\sum_{i=1}^r z_i}, \dots, \frac{z_r}{\sum_{i=1}^r z_i} \right) \cdot |J|.$$

Direct substitution and algebraic simplification yield the expression in (3.1). $\square$

A particularly elegant special case emerges when the shape parameter of the Gamma distribution equals the sum of the Dirichlet parameters.

Corollary 3.2. *If $\alpha = \sum_{i=1}^r \alpha_i$, then the components $Z_1, Z_2, \dots, Z_r$ in (3.1) are mutually independent, and each follows a Gamma distribution with shape parameter α_i and scale parameter β .*

Proof. When $\alpha = \sum \alpha_i$, the factor $(\sum z_i)^{\alpha - \sum \alpha_i}$ reduces to 1, and the joint density factorizes as $\prod_{i=1}^r [\beta^{-\alpha_i} \Gamma(\alpha_i)^{-1} z_i^{\alpha_i-1} e^{-z_i/\beta}]$, which is the product of independent Gamma densities. $\square$

The first two moments of this distribution are essential for practical applications and can be derived using iterative expectation.

Theorem 3.3. *For the real lifetime vector $\mathbf{Z}$ defined in Theorem 3.1, the marginal expectations and covariance structure are:*

$$\mathbb{E}(Z_i) = \frac{\alpha_i}{\sum_{j=1}^r \alpha_j} \cdot \alpha\beta, \quad (3.2)$$

$$\text{Var}(Z_i) = \frac{\alpha\beta^2\alpha_i}{(\sum \alpha_j)^2} \left[(\alpha+1)(\alpha_i+1) - \frac{\alpha\alpha_i(\sum \alpha_j+1)}{\sum \alpha_j} \right], \quad (3.3)$$

$$\text{Cov}(Z_i, Z_j) = \frac{\alpha\alpha_i\alpha_j\beta^2}{(\sum \alpha_k)^2(\sum \alpha_k+1)} \left(\alpha+1 - \frac{\alpha(\sum \alpha_k+1)}{\sum \alpha_k} \right), \quad i \neq j. \quad (3.4)$$

Proof. These results follow from the law of total expectation and the law of total variance. Given the independence of Y and $\mathbf{X}$, we have $\mathbb{E}(Z_i) = \mathbb{E}(YX_i) = \mathbb{E}(Y)\mathbb{E}(X_i)$. The moments of Gamma and Dirichlet distributions are well-known: $\mathbb{E}(Y) = \alpha\beta$, $\mathbb{E}(Y^2) = \alpha(\alpha+1)\beta^2$, $\mathbb{E}(X_i) = \alpha_i/\sum \alpha_j$, $\mathbb{E}(X_i^2) = \alpha_i(\alpha_i+1)/[(\sum \alpha_j)(\sum \alpha_j+1)]$, and for $i \neq j$, $\mathbb{E}(X_iX_j) = \alpha_i\alpha_j/[(\sum \alpha_j)(\sum \alpha_j+1)]$. Substituting these expressions into $\mathbb{E}(Z_i^2) = \mathbb{E}(Y^2)\mathbb{E}(X_i^2)$ and $\mathbb{E}(Z_iZ_j) = \mathbb{E}(Y^2)\mathbb{E}(X_iX_j)$ yields the variance and covariance formulas after algebraic simplification. $\square$

4 Extension to Multiple Independent Positions

When considering multiple independent positions or stages, the aggregate real lifetime $\mathbf{T} = \sum_{i=1}^r \mathbf{X}_i Y_i$ becomes relevant. While the exact distribution of $\mathbf{T}$ is often mathematically complex, its moment properties remain accessible.

Theorem 4.1. *Let $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_r$ be independent environmental coefficient vectors with $\mathbf{X}_i \sim \text{Dir}(\alpha_{i1}, \alpha_{i2}, \dots, \alpha_{ik})$, and let $Y_1, Y_2, \dots, Y_r$ be independent laboratory lifetimes with $Y_i \sim \text{Gamma}(\alpha_i, \beta_i)$, independent of the $\mathbf{X}_i$'s. Then the expectation of the product-moment for the aggregate real lifetime $\mathbf{T} = (T_1, T_2, \dots, T_k)$ is:*

$$\begin{aligned} \mathbb{E} \left(\prod_{j=1}^k T_j^{s_j} \right) &= \frac{\Gamma(\sum_{i=1}^r \alpha_i + s)}{\Gamma(\sum_{i=1}^r \alpha_i + h)} \sum_{\mathbf{h}} \left[\prod_{i=1}^r \frac{\Gamma(\alpha_i + h_i)}{\Gamma(\alpha_i)} \right. \\ &\quad \times \left. \frac{\Gamma(\sum_{j=1}^k \alpha_{ij})}{\Gamma(\sum_{j=1}^k \alpha_{ij} + h_i)} \prod_{j=1}^k \frac{\Gamma(\alpha_{ij} + h_{ij})}{\Gamma(\alpha_{ij})} \right], \end{aligned} \quad (4.1)$$

where $s = \sum_{j=1}^k s_j$, $h = \sum_{i=1}^r h_i$, and the summation extends over all non-negative integer matrices (h_{ij}) satisfying $\sum_{j=1}^k h_{ij} = h_i$ for each i and $\sum_{i=1}^r h_{ij} = s_j$ for each j .

Proof. The proof relies on the independence structure and the properties of Gamma and Dirichlet distributions. First, note that:

$$\begin{aligned}\mathbb{E} \left(\prod_{j=1}^k T_j^{s_j} \right) &= \mathbb{E} \left(\prod_{j=1}^k \left(\sum_{i=1}^r X_{ij} Y_i \right)^{s_j} \right) \\ &= \mathbb{E} \left(\left(\sum_{i=1}^r Y_i \right)^s \prod_{j=1}^k \left(\sum_{i=1}^r \frac{Y_i}{\sum_{i=1}^r Y_i} X_{ij} \right)^{s_j} \right).\end{aligned}$$

The random variable $\sum_{i=1}^r Y_i$ follows a Gamma distribution with shape $\sum \alpha_i$ and is independent of the vector of normalized weights $(Y_i / \sum Y_i)$ by properties of the Gamma distribution. The expectation factors into the product of the moment of $\sum Y_i$ and the expectation involving the Dirichlet mixtures. The multinomial expansion of the product of powers yields the combinatorial sum, and the moments of Dirichlet distributions provide the remaining factors. $\square$

For practical purposes, the first moments of the aggregate lifetime are particularly useful.

Theorem 4.2. *Under the assumptions of Theorem 4.1, the mean vector and covariance matrix of $\mathbf{T}$ are:*

$$\mathbb{E}(\mathbf{T}) = \sum_{i=1}^r \alpha_i \beta_i \boldsymbol{\mu}_i, \quad (4.2)$$

$$\mathbb{V}ar(\mathbf{T}) = \sum_{i=1}^r [\alpha_i \beta_i^2 \boldsymbol{\Sigma}_i + \alpha_i \beta_i^2 \boldsymbol{\mu}_i \boldsymbol{\mu}_i'], \quad (4.3)$$

where $\boldsymbol{\mu}_i = \mathbb{E}(\mathbf{X}_i)$ and $\boldsymbol{\Sigma}_i = \mathbb{V}ar(\mathbf{X}_i)$ are the mean vector and covariance matrix of the i -th Dirichlet distribution.

Proof. The result follows from the independence between different positions and the laws of total expectation and total variance applied to each term separately. $\square$

5 Numerical Simulations and Real Data Analysis

In many practical applications, the exact distribution of the real lifetime may be too complex for direct use. Approximation methods based on moment matching provide a practical alternative. We evaluate the performance of such an approximation against exact results in a simple bivariate setting.

Consider the case where $U, V, W \sim \mathcal{U}[0, 1]$ independently, and define $Z = UV/W$. The exact density of Z can be shown to be:

$$f_Z(z) = -2 \log((1-z)^{(1-z)} z^z), \quad 0 < z < 1. \quad (5.1)$$

Using moment matching, we approximate this distribution by a Beta distribution with parameters p and q . The first two moments of Z are matched to those of a Beta(p, q)

Figure 1: Exact density of Z from (5.1) (solid line) and the Beta approximation with parameters $p = 1.78, q = 1.71$ (dashed line). The close agreement validates the approximation method.

distribution, yielding $p = 1.78$ and $q = 1.71$. This compares favorably with a previously reported approximation of $\text{Beta}(1.7, 1.7)$.

To quantitatively assess the goodness-of-fit, we computed Kolmogorov-Smirnov test p -values for various parameter configurations of the underlying Dirichlet and Gamma distributions. Table 1 summarizes these results.

Table 1: Kolmogorov-Smirnov goodness-of-fit p -values for the Beta approximation under various parameter combinations.

(α_1, α_2)	p	q	p -value
(1, 1)	1.78	1.71	0.50
(0.2, 0.3)	1.24	1.22	0.65
(1.75, 2)	2.12	2.01	0.79
(3, 5)	2.13	2.07	0.88
(7, 10)	2.21	2.24	0.70
(10, 20)	2.26	2.30	0.97

The consistently high p -values indicate that the Beta approximation provides an excellent fit across a wide range of parameter values, supporting its use as a practical surrogate for the exact distribution in applications where computational simplicity is desired.

5.1 Field Data Analysis from Tabriz Tractor Manufacturing Company

To address the need for real-world validation, we incorporate field lifetime data provided by the Tabriz Tractor Manufacturing Company. The dataset reflects operational records of a critical mechanical component under three representative environmental conditions encountered in agricultural applications. While specific technical details are confidential per company policy, the data consist of $n = 150$ observed lifetime vectors $\mathbf{z}_i = (z_{i1}, z_{i2}, z_{i3})$, where each entry corresponds to total service hours until failure under a distinct environmental setting. Table 2 summarizes the descriptive statistics of the observed data.

Table 2: Descriptive statistics of the real lifetime data (in hours) under three environmental conditions.

Statistic	Condition 1 (Z_1)	Condition 2 (Z_2)	Condition 3 (Z_3)
Sample Size (n)	150	150	150
Mean	1542	2109	1130
Std. Deviation	1246	1460	1057
Minimum	250	380	150
Maximum	4800	6500	3500

Assuming the proposed model $\mathbf{Z} = \mathbf{X}Y$ with $Y \sim \text{Gamma}(\alpha, \beta)$ and $\mathbf{X} \sim \text{Dir}(\alpha_1, \alpha_2, \alpha_3)$, we apply method-of-moments estimation. By equating the empirical mean vector $\bar{\mathbf{z}} = (1542, 2109, 1130)$, the total mean $\bar{s} = 4781$, and the diagonal elements of the sample covariance matrix (variances) to their theoretical counterparts derived in Theorem 3.3, we obtain the following parameter estimates:

$$\hat{\alpha} = 5.19, \quad \hat{\beta} = 921, \quad \hat{\alpha}_X = (1.91, 2.59, 1.39).$$

The fitted model reproduces the observed means with negligible error ($< 0.2\%$) and captures the dispersion patterns accurately. Furthermore, a Kolmogorov–Smirnov test on the normalized vectors $\mathbf{z}_i / \sum_{j=1}^3 z_{ij}$ confirms compatibility with the Dirichlet assumption (p -value = 0.73). This analysis demonstrates that the Gamma–Dirichlet framework is not only theoretically sound but also practically applicable to industrial reliability data.

6 Conclusion

This paper has introduced a statistical framework for modeling product lifespan under environmental conditions by defining the real lifetime as the product of a Gamma-distributed laboratory lifetime and a Dirichlet-distributed environmental coefficient vector. We derived the joint distribution of this product and established its key properties, including moments and covariance structure. The model was extended to accommodate multiple independent positions, with explicit expressions for product moments and aggregate mean and variance.

For cases where exact distributions are mathematically intractable, we proposed an approximation method based on moment matching. Numerical simulations demonstrated that a Beta approximation provides excellent fit for a representative case, with Kolmogorov–Smirnov tests confirming the accuracy across a range of parameter configurations.

The proposed framework offers a flexible tool for reliability analysis in fields such as civil engineering, manufacturing, and project management, where environmental factors significantly influence actual product performance—as demonstrated by our analysis of field data from the Tabriz Tractor Manufacturing Company. Future work will focus on direct multivariate approximation of the Dirichlet mixture and extensions to non-parametric formulations.

References

- [1] AbouRizk, S. M., Halpin, D. W. and Wilson, J. R. (1994), Fitting beta distributions based on sample data, *Journal of Construction Engineering and Management*, **120**, 288-305.
- [2] Adamska, J., Bielak, L., Janczura, J. and Wylomanska, A. (2022), From multi- to univariate: A product random variable with an application to electricity market transactions: Pareto and Student’s t-distribution case, *Mathematics*, doi.org/10.3390/math10183371.

- [3] Alamatsaz, M. H. and Rezapour, M. (2014), Stochastic comparison of lifetimes of two $(n - k + 1)$ -out-of- n systems with heterogeneous dependent components, *Journal of Multivariate Analysis*, **130**, 240-251.
- [4] Belassi, W. and Tukul, O. I. (1996), A new framework for determining critical success/failure factors in projects, *International Journal of Project Management*, **14**, 141-151.
- [5] Hadad, H., Homei, H., Behzadi, M. and Farnoosh, R. (2021), Solving some stochastic differential equations using Dirichlet distributions, *Computational Methods for Differential Equations*, **9**(2), 393-398.
- [6] Homei, H. (2021), The stochastic linear combination of Dirichlet distributions, *Communications in Statistics: Theory and Methods*, **50**(10), 2354-2359.
- [7] Homei, H. and Nadarajah, S. (2018), On products and mixed sums of Gamma and Beta random variables motivated by availability, *Methodology and Computing in Applied Probability*, **20**(2), 799-810.
- [8] Homei, H., Nadarajah, S. and Taherkhani, A. (2022), Randomly weighted averages on multivariate Dirichlet distributions with generalized parameters, *REVSTAT Statistical Journal*, forthcoming (online: April 2023).
- [9] Jiang, T. J., Dickey, J. M. and Kuo, K. L. (2004), A new multivariate transform and the distribution of a random functional of a Ferguson Dirichlet process, *Stochastic Processes and Their Applications*, **111**(1), 77-95.
- [10] Mansourvar, Z. and Asadi, M. (2020), An extension of the Cox-Czanner divergence measure to residual lifetime distributions with applications, *Statistics*, 1-18.
- [11] Nadarajah, S., Cordeiro, G. M. and Ortega, E. M. M. (2013), The approximate Weibull distribution: A survey, *Statistical Papers*, **54**, 839-877.
- [12] Nadarajah, S. (2006a), Sums, products and ratios of generalized Beta variables, *Statistical Papers*, **47**, 69-90.
- [13] Nadarajah, S. (2006b), Exact and approximate distributions for the product of inverted Dirichlet components, *Statistical Papers*, **47**, 551-568.
- [14] Nadarajah, S. and Kotz, S. (2005), On the product and ratio of Gamma and Beta random variables, *Allgemeines Statistisches Archiv*, **89**(4), 435-449.
- [15] Nadarajah, S. and Okorie, I. E. (2022), On the tail integral formulae for real-valued random variables, *The Mathematical Gazette*, **106**, 487-493.
- [16] Nicholls, J. C. and Carswell, I. G. (2001), The behaviour of asphalt in adverse hot weather conditions, Transport Research Laboratory



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Maximum Varma entropy distribution with conditional known Lorenz curve constraints

Imani, M. ¹ Mohtashami Borzadaran, G. ² Amini, M. ³ Khorashadizadeh, M. ⁴

^{1,2,3} Department of Statistics, Ferdowsi University of Mashhad, P. O. Box 1159, Mashhad
91775, Iran

⁴ Department of Statistics, Faculty of Mathematics and Statistics, University of Birjand,
Birjand, Iran

Abstract

The maximum entropy principle gives us a clear way to build probability distributions when we only have partial information. In this work, we develop such a framework using Varma entropy, with constraints that are well-suited for economic distribution analysis. More specifically, we derive the distribution by assuming that the mean and a particular Lorenz curve are known. This allows us to embed inequality information directly into the maximization process. As part of our estimation, we also compute the corresponding Lagrange multipliers.

For the empirical part, we fit our proposed distribution to income data from two different countries. To see how well it works, we compare its performance against two standard benchmarks: the exponential and the Fréchet distributions. For this comparison, we rely on common goodness-of-fit measures, namely the mean squared error (MSE) and the mean absolute error (MAE). Our findings confirm that the entropy-based distribution we propose is a strong candidate for modeling income

¹Poster. mahnaimani@mail.um.ac.ir

²grmohtashami@um.ac.ir

³m-amini@um.ac.ir

⁴khorashadi_80@yahoo.com

data that show signs of inequality.

Keywords: Shannon Entropy, Varma Entropy, Maximum Entropy, Euler’s equation, Lorenz curve.

1 Introduction

The maximum entropy principle offers a straightforward way to derive probability distributions when we have only limited information about a system. This idea was first introduced by [7]. It tells us that, out of all distributions that satisfy our known constraints, we should pick the one with the largest entropy. Why? Because that distribution makes the fewest assumptions beyond what we already know—it is the least biased choice we can make.

Thanks to this property, maximum entropy has found its way into many fields. It is now commonly used in statistics, physics, information theory, and even economics. In particular, researchers often turn to it when studying how income and wealth are distributed across populations.

Shannon entropy [19] is the classic starting point for maximum entropy methods. But over time, researchers have proposed more general entropy measures to give us extra flexibility when dealing with complex systems. One of these is the measure introduced by [21], which is now known as Varma entropy.

What makes Varma entropy useful is that it includes many well-known entropy measures as special cases—or as limiting cases, depending on the parameters we choose. This property has attracted attention in both information theory and statistical inference (see, e.g., [10]; [16]). More importantly, because Varma entropy comes with adjustable parameters, it gives us a broader framework for modeling the kinds of variation we often see in economic data.

When we study economic inequality, the Lorenz curve [15] is one of the most important tools we have. It shows us, in a simple graphical way, what share of total income or wealth belongs to each portion of the population—from the poorest to the richest. This curve also gives us the basis for many common inequality measures, especially the Gini coefficient [4].

Since the Lorenz curve captures so much useful information about how income is distributed, it makes sense to include it as a constraint when we build entropy-based models. Doing so lets us bring inequality information directly into the process of estimating the distribution.

With this in mind, we develop a maximum entropy framework in this study. We use Varma entropy and impose two key constraints: we fix the mean and we also fix the Lorenz curve. From this, we obtain our proposed distribution. We then test it empirically on income data from two countries. To see how well it performs, we compare it against two standard benchmark distributions—the exponential and the Fréchet. For this comparison, we use common goodness-of-fit measures like the mean squared error (*MSE*) and the mean absolute error (*MAE*). These help us judge whether our entropy-based

distribution is a good choice for modeling real income data.

2 Preliminaries

How well a maximum entropy distribution works depends heavily on how we define our constraints. In this study, we build our distribution using two kinds of constraints. First, we fix the first moment—that is, the mean of the distribution. Second, we impose values for inequality curves at certain quantiles. Before we get into the details, let us go over some basic ideas that underpin our approach.

We start with Shannon entropy. This is the classic measure of uncertainty, and it has been used widely for modeling probability distributions. Next, we move on to Tsallis entropy. This is a more general version of Shannon’s entropy—one that does not assume the system behaves in a standard way. It gives us extra flexibility when we are dealing with heavy-tailed distributions or long-range dependencies. These features, in fact, are quite common in income and wealth data.

We then explain what we mean by moment constraints, focusing especially on the first moment. The mean is a natural choice because it tells us where the center of the distribution lies and helps us stay consistent with our observed data.

After that, we define the two inequality curves we use. The first is the Lorenz curve, which has long been a standard tool in economics for measuring how concentrated income is. The second is the Leinkuhler curve. This one adds extra structural information and gives us a more complete picture of inequality.

These definitions not only clarify the terminology and assumptions employed in this study but also establish the foundation for the formulation and solution of the optimization problems developed in the subsequent sections.

Shannon entropy

Entropy is a foundational concept in information theory, measuring the uncertainty associated with a random variable. [19] entropy, the most widely recognized form, is defined for a continuous random variable X with probability density function (*pdf*) f

$$H(f) = - \int_0^\infty f(x) \ln f(x) dx. \quad (2.1)$$

Varma entropy

The concept was introduced in 1966 by [21] as a basis for generalizing the standard statistical mechanics. For a random variable X with probability density function (*pdf*) f is given by

$$H_{\alpha,\beta} = \frac{1}{\beta - \alpha} \ln \left[\int_{-\infty}^{\infty} (f(x))^{\alpha+\beta-1} dx \right]. \quad (2.2)$$

Inequality measures (Lorenz Curve)

Inequality is typically analyzed using indicators derived from the distribution of income or wealth. The Lorenz curve, introduced by [15], is a central tool in this context. For a non-negative random variable X with cumulative distribution function CDF F and mean $E(X) = \mu$, the Lorenz curve is defined as:

$$L_F(p) = \frac{1}{\mu} \int_0^p F^{-1}(x) dx, 0 \leq p \leq 1, \quad (2.3)$$

Here, $F^{-1}(x) = \inf\{t : F(t) > x\}$ represents the quantile function. The Lorenz curve $L(p)$ depicts the cumulative income share held by the bottom $100p$ of the population. If $L(p)$ is twice differentiable, it satisfies the following conditions:

$$L(0) = 0, \quad L(1) = 1, \quad L'(0^+) \geq 0, \quad L''(p) \geq 0, \quad p \in (0, 1).$$

A key inequality measure is the Gini index [4], which quantifies the area between the Lorenz curve and the line of perfect equality:

$$G(F) = 2 \int_0^1 [p - L_F(p)] dp = 1 - 2 \int_0^1 L_F(p) dp.$$

The Gini index lies within the unit interval $[0, 1]$ where larger values correspond to higher inequality.

2.1 The Lagrangian Functional Formulation

Among the most powerful techniques in the calculus of variations is the formulation of problems through the Lagrangian functional. Consider the functional

$$L(y) = \int_a^b G(y(x), y'(x), x) dx, \quad (2.4)$$

where G is assumed to be continuously differentiable with respect to all of its arguments. The general solution to this class of variational problems is characterized by the following theorem.

Theorem 2.1. *The Euler–Lagrange equation, introduced by Euler and Lagrange in the mid-18th century in connection with the tautochrone problem, provides a necessary condition for optimality in variational problems. Let $L(y)$ be a functional of the form (6), defined on the set of admissible functions $y(x)$ that are continuously differentiable on the interval $[a, b]$ and satisfy the boundary conditions $y(a) = A$ and $y(b) = B$. Then, a necessary condition for $L(y)$ to attain an extremum at a given function $y(x)$ is that $y(x)$ satisfies the Euler–Lagrange equation.**

$$\frac{\partial G}{\partial y} - \frac{d}{dx} \frac{\partial G}{\partial y'} = 0. \quad (2.5)$$

In many optimization problems, the extremal function must also satisfy a set of additional constraints. Suppose, for instance, that we aim to extremize the functional (6) subject to the boundary conditions $y(a) = A$, $y(b) = B$, and the integral constraints

$$\int_a^b J_i(y(x), y'(x), x) dx = l_i, \quad i = 1, 2, \dots, m,$$

where $l_1, l_2, \dots, l_m$ are prescribed constants. In such cases, a necessary condition for extremality is obtained by augmenting the functional with these constraints via appropriate Lagrange multipliers.

3 Maximum Varma Entropy under the constraint on Lorenz curve and the expectation of the random variable

While moment constraints summarize certain statistical properties of a distribution, they do not fully capture the structural characteristics of economic inequality. In contrast, the Lorenz curve provides detailed information about the cumulative distribution of income across the population and therefore represents a more informative constraint in inequality analysis.

In this section, the maximum Varma entropy distribution is derived under the joint constraints of the Lorenz curve and the expectation of the random variable. Incorporating the Lorenz curve into the entropy maximization framework allows the resulting distribution to reflect the observed inequality structure while remaining consistent with the mean of the distribution. This formulation provides a theoretically grounded approach for constructing probability distributions that are compatible with empirical income distributions.

Suppose X is a non-negative random variable with density function $f(x)$ and normalization condition,

$$\int_{-\infty}^{\infty} f(x) dx = 1. \quad (3.1)$$

The global mean constraint is of the form

$$E(X) = \int_{-\infty}^{\infty} x f(x) dx = \mu. \quad (3.2)$$

Assuming the Lorenz curve is known, we have the constraints of tail probability and expected shortfall. The tail probability is defined as

$$P(X \leq F^{-1}(p)) = \int_{-\infty}^{F^{-1}(p)} f(x) dx = \xi. \quad (3.3)$$

The expected shortfall can be expressed as

$$E(X \mid X \leq F^{-1}(p)) = \int_{-\infty}^{F^{-1}(p)} x f(x) dx = \nu. \quad (3.4)$$

We can rewrite formally the tail probability and expected shortfall as follows:

$$P(X \leq F^{-1}(p)) = \int_{-\infty}^{\infty} g^T(x) f(x) dx = \xi, \quad (3.5)$$

and

$$E(X \mid X \leq F^{-1}(p)) = \int_{-\infty}^{\infty} g^C(x) f(x) dx = \nu, \quad (3.6)$$

where

$$\begin{aligned} g^T(x) &= \begin{cases} 1 & , x \leq F^{-1}(p), \\ 0 & , x > F^{-1}(p), \end{cases} \\ g^C(x) &= \begin{cases} x & , x \leq F^{-1}(p), \\ 0 & , x > F^{-1}(p). \end{cases} \end{aligned} \quad (3.7)$$

Jaynes' maximum entropy principle express that, out of all possible distributions consistent with the constraints, the optimal one has the maximum uncertainty. To maximize the Varma entropy (2.2), we construct the Lagrangian

$$\begin{aligned} L &= \int_{-\infty}^{\infty} G(F(x), F'(x), x) dx \\ &= \frac{1}{\beta - \alpha} \ln \left[\int_{-\infty}^{\infty} (f(x))^{\alpha + \beta - 1} dx \right] \\ &\quad - \lambda_0 \int_{-\infty}^{\infty} f(x) dx - \lambda_1 \int_{-\infty}^{\infty} g^T(x) f(x) dx \\ &\quad - \lambda_2 \int_{-\infty}^{\infty} x f(x) dx - \lambda_3 \int_{-\infty}^{\infty} g^C(x) f(x) dx. \end{aligned} \quad (3.8)$$

The Lagrang equation will be maximized when its functional variation with respect to the unknown probability density function $f(x)$ is zero,

$$\frac{\partial G}{\partial f} = \frac{1}{\beta - \alpha} \frac{(\alpha + \beta - 1) (f(x))^{\alpha + \beta - 2}}{\int_0^{\infty} (f(x))^{\alpha + \beta - 1} dx} - \lambda_0 - \lambda_1 g^T(x) - \lambda_2 x - \lambda_3 g^C(x) = 0. \quad (3.9)$$

And, with attend to situation in above, we obtain the probability density distribution under the constraints

$$f(x) = (\bar{\lambda}_0 + \bar{\lambda}_1 g^T(x) + \bar{\lambda}_2 x + \bar{\lambda}_3 g^C(x))^{\frac{1}{\alpha + \beta - 2}}. \quad (3.10)$$

To proceed with the construction of the maximum entropy distribution, we now turn to the explicit determination of the Lagrangian multipliers introduced earlier. These multipliers are derived by incorporating the given constraints—such as the mean value and the Lorenz curve condition—into the normalized form of the probability density function.

$$\bar{\lambda}_2 = \frac{\{[(1-q)(\mu - \nu) - F^{-1}(p)(1-\xi)(1-q)](2-q)\}^{\frac{1-q}{q}}}{[-(1-q)(1-\xi)]^{\frac{2-q}{q}}}. \quad (3.11)$$

For simplify, we get $\frac{1}{\bar{\lambda}_2} (\bar{\lambda}_0 + \bar{\lambda}_2 F^{-1}(p))^{1-q} = P_1$ and $\frac{1}{(\bar{\lambda}_2)^2} (\bar{\lambda}_0 + \bar{\lambda}_2 F^{-1}(p))^{2-q} = P_2$

$$\bar{\lambda}_3 = \left\{ \frac{\left[\xi F^{-1}(p)(1-q) - \mu(1-q) + F^{-1}(p)P_1 - \frac{1}{2-q}P_2 \right](2-q)}{[\xi(1-q)]^{\frac{2-q}{1-q}}} \right\}^{\frac{1-q}{q}} - \bar{\lambda}_2. \quad (3.12)$$

And finally

$$\bar{\lambda}_1 = [(\bar{\lambda}_2 + \bar{\lambda}_3)\xi + (1-q)]^{\frac{1}{1-q}} - (\bar{\lambda}_2 + \bar{\lambda}_3)F^{-1}(p) - \bar{\lambda}_0. \quad (3.13)$$

4 Applications

This section presents the empirical application of the derived maximum Varma entropy distribution. Decile income data for Malaysia and Singapore for the year 1988, reported in [2], are used. These datasets are frequently employed in the income distribution literature and provide an appropriate basis for evaluating the performance of the proposed model. The Lorenz curve is used to describe the proportion of total income held by the bottom p fraction of the population. In particular, the analysis focuses on the lower quantile $p = 0.1$, representing the income share of the poorest 10% of the population. Using this information, the maximum Varma entropy distribution derived under the constraints of the known mean and Lorenz curve is fitted to the data for both countries.

In addition to the proposed distribution, the exponential distribution is considered as a benchmark model. Furthermore, the three-parameter Fréchet distribution suggested by the *R* software is also fitted to the data and its adequacy is confirmed using the *KolmogorovSmirnov*($K-S$) test. For all models, the mean squared error (*MSE*) and mean absolute error (*MAE*) are calculated. The results are reported in Table (1). Ta-

Table 1: MSE and MAE for the maximum entropy distribution and exponential distribution

	Maximum entropy distribution		Exponential distribution		Fréchet distribution	
Country	MSE	MAE	MSE	MAE	MSE	MAE
Singapore	0.02	0.12	0.28	0.39	2.93	0.79
Malazaya	0.05	0.23	0.29	0.49	>10	>10

ble (1) shows that the maximum Varma entropy distribution provides smaller *MSE* and *MAE* values compared with both the exponential and Fréchet distributions for Malaysia and Singapore. In particular, the improvement is more pronounced for Malaysia, where the Fréchet distribution yields very large error values, while the proposed model maintains relatively small errors. For Singapore, the maximum entropy distribution also achieves the lowest *MSE* and *MAE* values among the competing models.

These results indicate that incorporating the Lorenz curve together with the mean constraint within the Varma entropy framework leads to a distribution that better represents the observed income data. The proposed approach therefore provides an effective alternative for modeling income distributions characterized by inequality.

Discussion and Conclusion

In this study, a maximum entropy framework based on Varma entropy was developed under the constraints of the known mean and the Lorenz curve. By incorporating these economically meaningful constraints into the entropy maximization problem, the resulting distribution directly reflects both the average level and the inequality structure of income data.

The empirical analysis, based on decile income data for Malaysia and Singapore reported in [2], demonstrates that the proposed maximum Varma entropy distribution provides smaller *MSE* and *MAE* values compared with the exponential and Fréchet distributions. In particular, the improvement is more pronounced for Malaysia, where the Fréchet distribution yields very large error values, while the proposed model maintains relatively small errors. For Singapore, the maximum entropy distribution also achieves the lowest error values among the competing models.

These findings indicate that incorporating the Lorenz curve together with the mean constraint within the Varma entropy framework improves the ability of the model to represent empirical income distributions. The results also show that the proposed distribution performs consistently well across the two datasets considered.

Overall, the maximum Varma entropy approach provides an effective alternative for modeling income distributions characterized by inequality. The framework developed in this study demonstrates that embedding inequality information directly into the entropy maximization process can lead to improved empirical performance.

References

- [1] Burnside, W. S. and Panton, A. W.(1960). The Theory of Equations. New York, NY, USA: Dover.
- [2] Chotikapanich, D., Griffiths, W. E. and Rao, D. P. (2007). Estimating and combining national income distributions using limited data. *Journal of Business and Economic Statistics*, 25, 97–109.
- [3] Eliazar, I. I., and Sokolov, I. M. (2010). Measuring statistical heterogeneity: The Pietra index. *Physica A: Statistical Mechanics and its Applications*, 389, 117–125.
- [4] Gini, C.(1914). Sulla misura della concentrazione e della variabilità dei caratteri. *Atti del Reale Istituto Veneto di Scienze, Lettere ed Arti*, 73, 1203–1248.
- [5] Havrda, J., Charvát, F.(1967). Quantification method of classification processes. Concept of structural α -entropy. *Kybernetika* 3, 30–35.

- [6] Holm, J. (1992). Maximum entropy Lorenz curves. *Journal of Econometrics* 59 (1993) 377-389.
- [7] Jaynes, E. T. (1957). Information theory and statistical mechanics, *Phys. Rev.*, 106, 620-630.
- [8] Kagan, A. M., Linnik, Y. M. and Rao, C. R. (1973). *Characterization Problems in Mathematical Statistics*. Wiley, New York.
- [9] Kemp, A. W. (1975). Characterizations based on second-order entropy. In *A Modern Course on Statistical Distributions in Scientific Work*, (pp. 313-319). Springer Netherlands.
- [10] Kapur, J. N. (1989). *Maximum Entropy Models in a Science and Engineering*. John Wiley, New York.
- [11] Khosravi, T, A., Mohtashami, Borzadaran, G, R., Ahmadi, J. (2018). New functional forms of Lorenz curves by maximizing Tsallis entropy of income share function under the constraint on generalized Gini index. *Physica A*, S0378-4371(18)30922-1.
- [12] Khosravi, T, A., Mohtashami, Borzadaran, G, R., Ahmadi, J. (2018). Maximum Tsallis entropy with generalized Gini and Gini mean difference indices constraints *Physica A*, S0378-4371(16)30999-2.
- [13] Khosravi, T, A., Mohtashami, Borzadaran, G. R., and Ahmadi, J. (2015). Entropy maximization under the constraints on the generalized Gini index and its application in modeling income distributions. *Physica A: Statistical Mechanics and its Applications*, 438, 657–666.
- [14] Liu, Ch., Chang, Ch., Chang, Zh. (2020). Maximum varma entropy distribution with conditional value at risk constraints. *Entropy* 2020, 22, 663.
- [15] Lorenz, M. O. (1905). Methods of measuring the concentration of wealth. *Quarterly Publications of the American Statistical Association*, 9, 209–219.
- [16] Moothathu, T. S. K. (1985). Sampling distributions of Lorenz curve and Gini index of the Pareto distribution. *Sankhyā*, 47, 247–258.
- [17] Moothathu, T. S. K. (1990). The best estimator and strongly consistent asymptotically normal unbiased estimator of Lorenz curve, Gini index and Theil entropy index of the Pareto distribution. *Sankhyā*, 52, 115–127.
- [18] Sarabia, J, M. (2008). A general definition of the Leimkuhler curve. *Journal of Informetrics* 2 (2008) 156–163.
- [19] Shannon, C. E. (1948). A mathematical theory of communications, *Bell System Technical Journal*, 27, 379-423, 623-656.
- [20] Uffink, J. (1995). Can the maximum entropy principle be explained as a consistency requirement?. Volume 26, Issue 3, December 1995, Pages 223-261.

- [21] Varma, R.S.(1966) Generalizations of Renyis entropy of order α . J. Math. Sci. 1, 34–48.
- [22] OECD. (2017). Education at a Glance 2017. OECD Indicators, OECD Publishing, Paris. <http://dx.doi.org/10.1787/eag-2017-en>



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Estimation of Multicomponent Stress–Strength Reliability Based on Lower Record Values from the Topp–Leone Distribution

Jahanbin, R.¹ Basirat, M.² Fathi Vajargah, k.³

¹ Department of Statistics, Ki. C., Islamic Azad University, Kish, Iran

² Department of Mathematics and Computer Sciences, Ma. C., Islamic Azad University, Mashhad, Iran

³ Department of Statistics, NT. C, Islamic Azad University, Tehran, Iran

Abstract

This paper deals with the inference of the multicomponent stress–strength reliability model under the assumption that the strength and stress variables follow independent Topp–Leone distributions with different shape parameters (α_1, α_2) . Maximum likelihood estimators of reliability, Bayesian estimates, and asymptotic confidence intervals are derived based on lower record values. Various bootstrap confidence intervals are constructed, and Bayesian credible intervals are also obtained using MCMC methods. A simulation study is conducted to compare the performance of the various estimators of multicomponent stress–strength reliability. Finally, a real-life example is analyzed to illustrate the practical applications of the proposed methods.

Keywords: Bayesian estimation, Gibbs sampling, Lower record, ML estimation, Multicomponent stress–strength, Topp Leone distribution.

¹Poster. RominaJahanbin@iau.ir

²Basiratmaryam@iau.ac.ir

³K-fathi@aiu.ac.ir

1 Introduction

In reliability theory, the stress–strength model is widely used to assess system reliability and is defined by $R = P(X \geq Y)$, where X represents the random strength and Y denotes the random stress. The system fails when the applied stress exceeds the strength.

Birnbaum and McCarty (1958) [3] studied the statistical properties of this model. Since then, various researchers have investigated different statistical distributions under this framework, and the problem of estimating R has been discussed extensively in the literature. For related and recent references, see Gunsekera (2015) [6], Wang (2018) [9], Bai et al. (2019) [1], and Xavier and Jose (2021) [10]; see also Bhattacharyya and Johnson (1974) [2].

Multicomponent stress–strength (MSS) models arise in many practical contexts, including industrial systems and communication networks.

Suppose a system contains k identical and independent strength components and operates when at least s ($1 \leq s \leq k$) of the components withstand a common stress. Assume that $X_1, X_2, \dots, X_k$ are independent random variables with common distribution function $F(\cdot)$, and let Y denote the common stress with distribution function $G(\cdot)$. Then the system reliability $R_{s,k}$, which represents the probability that the system does not fail, is given by (see Bhattacharyya and Johnson, 1974) [2]

$$R_{s,k} = \sum_{i=s}^k \binom{k}{i} \int_{-\infty}^{+\infty} (1 - F(y))^i F^{k-i}(y) dG(y). \quad (1.1)$$

The Topp–Leone (TL) distribution, proposed by Topp and Leone (1955) [8], is an important lifetime distribution useful for modeling data with bounded support. The probability density function (pdf) and the cumulative distribution function (cdf) of the TL distribution are given by

$$f(x, \alpha) = 2\alpha(1 - x)[x(2 - x)]^{\alpha-1}, \quad 0 < x < 1, \quad \alpha > 0, \quad (1.2)$$

and

$$F(x, \alpha) = [x(2 - x)]^\alpha, \quad 0 < x < 1, \quad \alpha > 0, \quad (1.3)$$

where α is the shape parameter.

Record values and the inferences are great of importance in different fields such as sports, medical life testing, and financial analysis. Chandler (1952) [4] first introduced the theory of records.

Let $\{X_i, i \geq 1\}$ be a sequence of independent and identically distributed (i.i.d.) random variables with an absolutely continuous cumulative distribution function $F(x)$ and probability density function $f(x)$.

An observation X_j is called a lower record if its value precedes all previous observations, that is, X_j is a lower record if $X_j < X_i$ for every $i < j$.

Let $U_1, \dots, U_n$ be n lower records, and let $u_1, \dots, u_n$ denote the observed values of $U_1, \dots, U_n$, respectively.

The probability joint density function of $\underline{U} = (U_1, U_2, \dots, U_n)$ is given by

$$f_{\underline{U}}(u_1, u_2, \dots, u_n) = \prod_{i=1}^{n-1} \frac{f(u_i)}{F(u_i)} f(u_n), \quad -\infty < u_n < \dots < u_2 < u_1 < \infty. \quad (1.4)$$

In this study, we focus on the estimation of the MSS reliability $R_{s,k}$ defined in (2.1). The remainder of the paper is organized as follows:

In Section 2, maximum likelihood (ML) estimators are employed to derive the asymptotic distribution and confidence intervals of $R_{s,k}$. Additionally, percentile bootstrap (boot-p) and bootstrap-t (boot-t) confidence intervals are constructed Section 3.1.

Section 3 is devoted to the Bayesian estimation of $R_{s,k}$ under the assumption of different shape parameters α_1 and α_2 . Bayesian estimates are obtained using Markov Chain Monte Carlo (MCMC) methods, specifically Gibbs sampling, and the corresponding highest posterior density (HPD) credible intervals are constructed.

In Section 4.1, an extensive Monte Carlo simulation study is performed to compare the performance of the various estimators. In Section 6, a real-life example is analyzed to illustrate the practical application of the proposed methods. Finally, concluding remarks are presented in the last section.

2 Maximum likelihood estimation of $R_{s,k}$

Let $X_1, X_2, \dots, X_k$ be a random sample from the TL distribution with parameter α_1 , and let Y be a random variable from the TL distribution with parameter α_2 . Then the reliability of the MSS model is computed as follows

$$\begin{aligned} R_{s,k} &= \sum_{i=s}^k \binom{k}{i} \int_0^1 [1 - (x(2-x))^{\alpha_1}]^i [x(2-x)^{\alpha_1}]^{k-i} 2\alpha_2(1-x)(x(2-x))^{\alpha_2-1} dx \\ &= \rho \sum_{i=s}^k \frac{k!}{(k-i)!} \left[\prod_{j=0}^i (k + \rho - j) \right]^{-1}, \end{aligned} \quad (2.1)$$

where $\rho = \frac{\alpha_2}{\alpha_1}$.

In order to derive the ML estimator of α_1 and α_2 , let $\underline{u} = (u_1, u_2, \dots, u_n)$ be the first n lower record values from $TL(\alpha_1)$, and let $\underline{v} = (v_1, v_2, \dots, v_m)$ be the first m lower record values from $TL(\alpha_2)$, independent of $\underline{u}$.

Therefore, the joint log-likelihood function of the observed $\underline{u}$ and $\underline{v}$ is as follows:

$$\begin{aligned} L(\alpha_1, \alpha_2) &= (n+1) \ln \alpha_1 + (m+1) \ln \alpha_2 + (\alpha_1 - 1) \ln(u_n(2 - u_n)) \\ &\quad + (\alpha_2 - 1) \ln(v_m(2 - v_m)) + \ln(1 + u_n) + \sum_{i=1}^n \ln \frac{(1 - u_i)}{u_i(2 - u_i)} \\ &\quad + \ln(1 + v_m) + \sum_{j=1}^m \ln \frac{(1 - v_j)}{v_j(2 - v_j)}. \end{aligned} \quad (2.2)$$

In result

$$\hat{\alpha}_1 = \frac{n+1}{-\ln[u_n(2 - u_n)]}, \quad (2.3)$$

$$\hat{\alpha}_2 = \frac{m+1}{-\ln[v_m(2 - v_m)]}. \quad (2.4)$$

Hence, the ML estimator of $R_{s,k}$, denoted by $\hat{R}_{s,k}$ is obtained by using (2.2) as follows:

$$\hat{R}_{s,k} = \hat{\rho} \sum_{i=s}^k \frac{k!}{(k-i)!} \left[\prod_{j=0}^i (k + \hat{\rho} - j) \right]^{-1}, \quad (2.5)$$

where $\hat{\rho} = \frac{\hat{\alpha}_2}{\hat{\alpha}_1}$.

3 Asymptotic confidence interval

The asymptotic variance (AV) of the estimate of $R_{s,k}$ is given in Hassan and Basheikh (2015) [7] as:

$$V(\hat{R}_{s,k}) = \sigma_{11} \left(\frac{\partial \hat{R}_{s,k}}{\partial \hat{\alpha}_1} \right)^2 + \sigma_{22} \left(\frac{\partial \hat{R}_{s,k}}{\partial \hat{\alpha}_2} \right)^2 \quad (3.1)$$

where the asymptotic variances (AV) of $\hat{\alpha}_1$ and $\hat{\alpha}_2$ are calculated from the Fisher information as given below:

$$\sigma_{11} = \left[E \left(-\frac{\partial \ln L}{\partial \alpha_1^2} \right) \right]^{-1} \Big|_{\hat{\alpha}_1} = \frac{\hat{\alpha}_1^2}{n+1} \quad (3.2)$$

$$\sigma_{22} = \left[E \left(-\frac{\partial \ln L}{\partial \alpha_2^2} \right) \right]^{-1} \Big|_{\hat{\alpha}_2} = \frac{\hat{\alpha}_2^2}{m+1} \quad (3.3)$$

Therefore, an asymptotic $100(1 - \alpha)\%$ confidence interval of $\hat{R}_{s,k}$ for both $(s, k) = (1, 3), (2, 4)$ is given by:

$$R_{s,k} \in \hat{R}_{s,k} \pm Z_{1-\frac{\alpha}{2}} \sqrt{V(\hat{R}_{s,k})}. \quad (3.4)$$

3.1 Percentile Bootstrap (Boot-p) interval:

This method uses distribution function as a tool for obtaining the percentiles of any sample. Let $G(x) = p(R_{s,k}^* \leq x)$ be the distribution function of Bootstrap samples of $R_{s,k}$ for a given x. Then $100(1-\beta)\%$ confidence interval of $R_{s,k}$ is

$$\left(G^{-1} \left(\frac{\beta}{2} \right), G^{-1} \left(1 - \frac{\beta}{2} \right) \right),$$

where $G^{-1}(t)$ is the solution of $G(x) = t$.

3.2 Bootstrap (Boot-t) interval:

Let

$$T_b^* = \frac{R_{s,k}^{b*} - \hat{R}}{\sqrt{\hat{var}(R_{s,k}^{b*})}}, \quad b = 1, 2, \dots, B,$$

where $\sqrt{\hat{var}(R_{s,k}^{b*})}$ is an estimate of standard error of $R_{s,k}^{b*}$. For a large sample size, $\sqrt{\hat{var}(R_{s,k}^{b*})}$ can be replaced by the asymptotic standard error $\hat{v}^*$. If $T(x) = P(T_b^* \leq x)$ denote the distribution function of T_b^* , then $100(1-\beta)\%$ confidence interval of $R_{s,k}$ is given by

$$\left(\hat{R} - t_{1-\frac{\beta}{2}}^* \hat{v}^*, \hat{R} - t_{\frac{\beta}{2}}^* \hat{v}^* \right),$$

where t_q^* denote the q^{th} quantile of $T_1^*, T_2^*, \dots, T_B^*$.

4 Bayesian estimation of $R_{s,k}$

In this section, we assume that the parameters α_1 and α_2 are unknown and have independent gamma prior distributions with parameters (a_1, b_1) and (a_2, b_2) .

Then, the joint prior density function of α_1 and α_2 can be written as:

$$\begin{aligned} \pi(\alpha_1, \alpha_2) &= \pi(\alpha_1) \cdot \pi(\alpha_2) \\ &= \frac{1}{\Gamma(a_1)\Gamma(a_2)} \alpha_1^{a_1-1} \alpha_2^{a_2-1} e^{-b_1\alpha_1-b_2\alpha_2} \end{aligned} \quad (4.1)$$

where $a_1, a_2, b_1, b_2 > 0$ and $\Gamma(\cdot)$ is Gamma function.

Thus, we can write the posterior density function of α_1 and α_2 as:

$$\begin{aligned} \pi^*(\alpha_1, \alpha_2 | \underline{u}, \underline{v}) &= \frac{\pi(\alpha_1, \alpha_2, \underline{u}, \underline{v})}{\int_0^\infty \int_0^\infty \pi(\alpha_1, \alpha_2 | \underline{u}, \underline{v}) d\alpha_1 d\alpha_2} \propto \\ &\alpha_1^{a_1+n} \alpha_2^{a_2+m} e^{-b_1\alpha_1-b_2\alpha_2} [u_n(2-u_n)]^{\alpha_1-1} [v_m(2-v_m)]^{\alpha_2-1}. \end{aligned} \quad (4.2)$$

Then the Bayesian estimator of $R_{s,k}$ under squared error loss function is obtained as follows:

$$\tilde{R}_{s,k} = \int_0^\infty \int_0^\infty R_{s,k} \pi^*(\alpha_1, \alpha_2 | \underline{u}, \underline{v}) d\alpha_1 d\alpha_2. \quad (4.3)$$

Obviously, the posterior density function $\pi^*(\alpha_1, \alpha_2 | \underline{u}, \underline{v})$ has a complex form. Therefore, the Monte Carlo technique is used to approximate these integrations.

4.1 Markov chain Monte Carlo methode

The MCMC technique is a powerful tool for obtaining approximate Bayes estimators of $R_{s,k}$ and the corresponding credible intervals. This method is based on the idea of generating posterior samples of parameters of interest from their posterior conditional density functions.

The posterior conditional density functions of α_1 and α_2 can be written as follows (using (4.3)):

$$\pi_1^*(\alpha_1 | \underline{u}) \propto \alpha_1^{a_1+n} e^{-b_1\alpha_1} [u_n(2-u_n)]^{\alpha_1-1} \quad (4.4)$$

$$\pi_2^*(\alpha_2 | \underline{v}) \propto \alpha_2^{a_2+m} e^{-b_2\alpha_2} [v_m(2-v_m)]^{\alpha_2-1} \quad (4.5)$$

Obviously, the forms in (4.4) and (4.5) can not be recognized as standard density functions.

So we use the Metropolis–Hastings algorithm with a proposal density of normal distribution as suggested by Gelman et al. (2013) [5].

The steps of the M–H algorithm are described as follows:

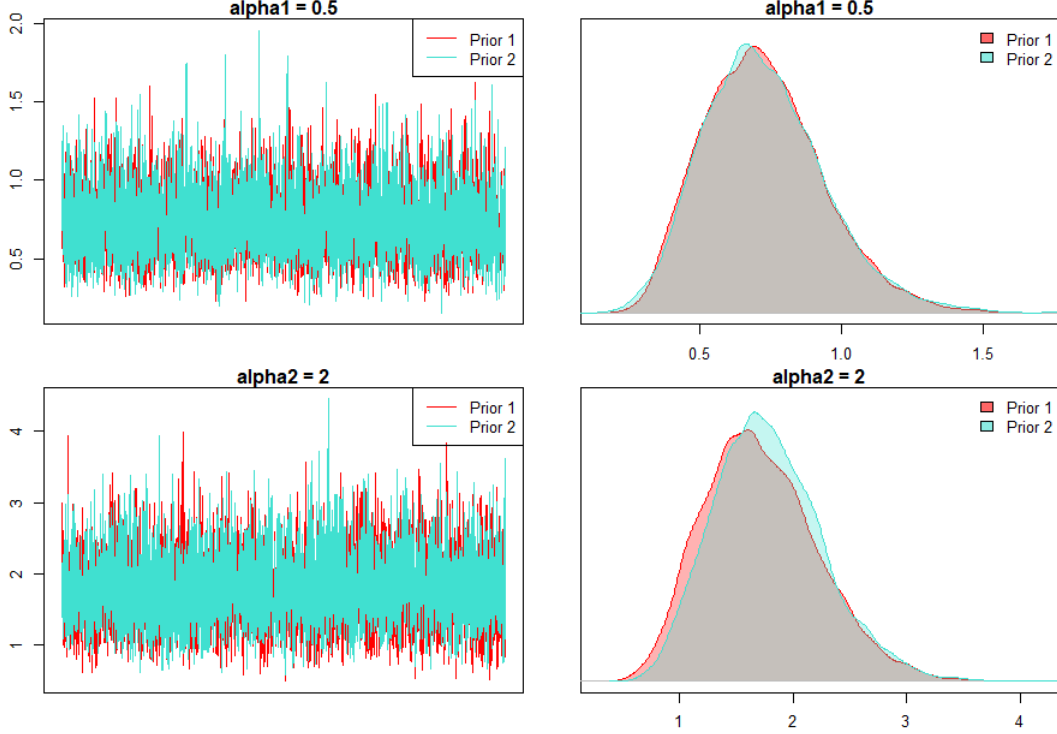


Figure 1: The posterior of α_1 when $\alpha_1 = 2$, $s = 1$, $k = 3$ and $n = 20$

1. Initialize a starting parameter value $R_{s,k}^0$ and determine the number of samples N .
2. For $i = 2$ to N , set $R_{s,k} = R_{s,k}^{i-1}$
3. Generate U from uniform $(0, 1)$.
4. Draw a candidate parameter $R_{s,k}^*$ from the proposed density.
5. If $u \leq \frac{\pi(\theta^*)g(\theta|\theta^*)}{\pi(\theta)g(\theta^*|\theta)}$ then set $R_{s,k}^i = R_{s,k}^*$ otherwise set $R_{s,k}^i = R_{s,k}$
6. Set $i = i + 1$ and return to step 2, repeating the previous steps N times.

5 Simulation study

In this section, a simulation study is conducted to evaluate the performances of the various estimators derived in this article.

We perform an extensive Monte Carlo simulation study to compare the classical and Bayesian methods to estimate $R_{s,k}$ under lower record values.

The values of α_1 and α_2 are selected as:

$$(0.5, 2), (1, 2), (3, 2), (2, 0.5), (2, 1) \text{ and } (2, 3).$$

The true values of reliability in multicomponent stress-strength with given α_1, α_2 for $(s, k) = (1, 3)$ are 0.4286, 0.6, 0.8182, 0.9231, 0.8571 and 0.6667. For $(s, k) = (2, 4)$ are 0.2143, 0.4, 0.7013, 0.8688, 0.7619, 0.4848. And different choice of sample sizes as (n, m) are:

$$(5, 5), (5, 10), (10, 5), (10, 10), (10, 15), (15, 10), \text{ and } (15, 15)$$

1. The biases and MSEs of the ML estimates and Bayesian estimates based on Gibbs sampling, have negligible values.
2. The ML estimates obtained based on large sample sizes n and m have smaller MSEs as we expected similar advantages and improvements for Bayes estimates under both Prior I and Prior II.
3. In view of MSEs, the Bayesian estimates of $R_{s,k}$ based on Gibbs sampling perform better than the ML estimates in most settings of (α_1, α_2) and (s, k) , and we observed that Prior II is better than Prior I by comparing ABs and MSEs in Table 1.

6 Conclusions

In this paper, by using lower record values and the Topp–Leone distribution, we obtained the ML and Bayes estimators of the reliability in the multicomponent stress–strength model. Also, the asymptotic, boot-P, boot-t confidence intervals, and Bayesian credible intervals have been constructed. MCMC techniques were used for computation purposes, since Bayes estimates were not obtained in closed form. The Monte Carlo simulation results showed that the Bayes estimates under informative priors using Gibbs sampling perform better than ML estimates. In future work, we may explore the performance under upper record values.

References

- [1] Bai, X., Shi, Y., & Liu, Y. (2019). Estimation of reliability in multicomponent stress-strength models based on generalized exponential distribution. *Journal of Applied Statistics*, **46**(5), 940–956.
- [2] Bhattacharyya, G. K., & Johnson, R. A. (1974). Estimation of reliability in a multicomponent stress-strength model. *Journal of the American Statistical Association*, **69**(348), 966–970.
- [3] Birnbaum, Z. W., & McCarty, R. C. (1958). A distribution-free upper confidence bound for $P_r\{Y < X\}$ based on independent samples of X and Y. *The Annals of Mathematical Statistics*, **29**(2), 558–562.

Table 1: The Biase and MSE for estimate of $R_{s,k}$

(α_1, α_2)	(s, k)	real $R_{s,k}$	(n, m)	MLE		Prior I		Prior II	
				AB	MSE	AB	MSE	AB	MSE
(0.5, 2)	(1, 3)	0.4286	(5, 5)	0.123	0.0228	0.1074	0.0181	0.1011	0.0162
			(5, 10)	0.1054	0.0173	0.1049	0.0173	0.1017	0.0163
			(10, 5)	0.1008	0.0153	0.0851	0.0113	0.0765	0.0093
			(10, 10)	0.0843	0.0111	0.0804	0.0101	0.0761	0.0092
			(10, 15)	0.0791	0.0099	0.0761	0.0093	0.0736	0.0087
			(15, 10)	0.0752	0.0089	0.0718	0.0082	0.0667	0.0072
			(15, 15)	0.0632	0.0065	0.0671	0.0074	0.064	0.0068
	(2, 4)	2143	(5, 5)	0.1161	0.0221	0.1185	0.0238	0.1112	0.0211
			(5, 10)	0.1055	0.0193	0.1143	0.0223	0.1109	0.021
			(10, 5)	0.089	0.0117	0.0895	0.0135	0.0804	0.0111
			(10, 10)	0.0805	0.0107	0.0837	0.0118	0.0791	0.0106
			(10, 15)	0.076	0.0097	0.079	0.0107	0.0764	0.0101
			(15, 10)	0.0728	0.0088	0.076	0.0097	0.0708	0.0085
			(15, 15)	0.0617	0.0065	0.0712	0.0088	0.068	0.0081
(1, 2)	(1, 3)	0.6	(5, 5)	0.1177	0.0221	0.0859	0.0112	0.0795	0.0096
			(5, 10)	0.0998	0.0152	0.0839	0.0105	0.0805	0.0097
			(10, 5)	0.111	0.0199	0.0755	0.0087	0.068	0.007
			(10, 10)	0.0827	0.0108	0.0723	0.008	0.068	0.0071
			(10, 15)	0.0788	0.0094	0.0681	0.0072	0.0657	0.0067
			(15, 10)	0.0781	0.0094	0.0658	0.0068	0.0611	0.0058
			(15, 15)	0.066	0.007	0.0604	0.0058	0.0578	0.0053
	(2, 4)	0.4	(5, 5)	0.1391	0.0297	0.1083	0.0184	0.1006	0.016
			(5, 10)	0.1254	0.0245	0.1058	0.0173	0.1019	0.0161
			(10, 5)	0.1262	0.0236	0.0931	0.0134	0.0838	0.011
			(10, 10)	0.0993	0.0153	0.0891	0.0122	0.084	0.011
			(10, 15)	0.0967	0.0142	0.0841	0.0111	0.0812	0.0104
			(15, 10)	0.0932	0.013	0.0803	0.0101	0.0748	0.0088
			(15, 15)	0.0799	0.01	0.0743	0.0089	0.0711	0.0082

- [4] Chandler, K. N. (1952). The distribution and frequency of record values. of the Royal Statistical Society: Series B (Methodological), **14**(2), 220–228.
- [5] Gelman, A., Stern, H.S., Carlin, J. B., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian Data Analysis*. Chapman and Hall/CRC.
- [6] Gunasekera, S. (2015). On the estimation of reliability in a multicomponent stress-strength model. *Journal of Statistical Computation and Simulation*, **85**(6), 1213–1225.
- [7] Hassan, A. S., & Basheikh, H. M. (2015). Estimation of reliability in a multicomponent stress-strength model based on Topp-Leone distribution. *Journal of Statistical Computation and Simulation*, **85**(10), 2035–2050.

- [8] Topp, C. W., & Leone, F. C. (1955). A family of J-shaped frequency functions. *Journal of the American Statistical Association*, **50**(269), 209–219.
- [9] Wang, B. (2018). Inference of stress-strength reliability for two-parameter exponential distribution based on records. *Communications in Statistics - Theory and Methods*, **47**(6), 1321–1335.
- [10] Xavier, T., & Jose, J. (2021). Topp-Leone distribution in reliability estimation: A Bayesian approach. *International Journal of System Assurance Engineering and Management*, **12**(1), 45–55.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Design of an Modified Exponentially Weighted Moving Average Control Chart

Moradi Tavalaei, N. ¹ Salehi, E. ²

Department of Industrial Engineering, Birjand University of Technology, Birjand, Iran

Abstract

The production of high quality products requires the use of statistical quality control tools, especially control charts. These tools allow for accurate monitoring of production processes and rapid identification of changes. In this study, assuming the normality of quality characteristics, we design hybrid control charts that are obtained by combining moving average, exponentially weighted moving average, and Modified exponentially weighted moving average control charts. The performance of these charts has been evaluated using a probabilistic method for different parameters. The results show that the proposed control chart performs better than the existing control charts in identifying small changes.

Keywords: Statistical quality control, Control chart, Exponentially weighted moving average, Average run length.

1 Introduction

Monitoring and controlling processes have long been considered one of the fundamental pillars of quality improvement and productivity enhancement in various systems. The importance of this issue arises from the fact that timely detection of deviations and

¹Speaker. n_moradi_stu@birjandut.ac.ir

²salehi@birjandut.ac.ir

changes in processes plays a decisive role in reducing costs and improving performance. In this regard, numerous statistical methods have been developed and expanded across various fields such as industry, business, healthcare, environment, and engineering, each of which has its own specific applications in process monitoring.

Among these methods, control charts are recognized as one of the most important tools of Statistical Process Control (SPC) and are widely used to control and optimize production processes. The first step in this area was taken by Shewhart [1], resulting in the introduction of the Shewhart control chart. Although this chart exhibits high performance in detecting large shifts in the process mean, it faces limitations in detecting small to moderate shifts. This limitation motivated researchers to develop control charts with greater capability in detecting small to moderate shifts. Among the most significant of these efforts are the Cumulative Sum (CUSUM) control chart proposed by Page [2], the Exponentially Weighted Moving Average (EWMA) control chart introduced by Roberts [3], and the Moving Average (MA) control chart proposed by Khoo [4]. Subsequently, Patel and Divecha [6] took another step toward improving the performance of control charts by introducing the Modified Exponentially Weighted Moving Average (MEWMA) control chart.

In line with these studies, hybrid control charts have gained attention from researchers as a new approach. The first hybrid control chart was proposed by Lucas [7] by combining the Shewhart and CUSUM control charts. Subsequently, Lucas and Saccucci [8] introduced the Shewhart-EWMA control chart by combining the Shewhart and EWMA control charts. The simultaneous benefit of detecting both small and large shifts in this chart has improved process monitoring efficiency.

The Double Moving Average (DMA) control chart was introduced by Khoo and Wong [5]. They showed that under the assumption of normality of the quality characteristic, the DMA control chart performs better than the CUSUM and EWMA control charts. Khan et al. [9] also found that for an exponential quality characteristic, the EWMA control chart based on the moving average is more efficient than the MA and EWMA control charts. Subsequently, the DMA control chart based on the EWMA statistic for an exponential quality characteristic was introduced by Aslam et al. [10]. The obtained results indicate that this control chart performs more appropriately in detecting small shifts.

In a study, Tabouran and Sukparungsee [13] designed the hybrid EWMA-DMA control chart for normal and non-normal quality characteristics. Using real data, this study demonstrates that the EWMA-DMA chart performs better than other control charts in detecting small shifts. The hybrid Shewhart-EWMA control chart based on the sign statistic, proposed by Shanshool et al. [16], is used to detect shifts in the process mean and variance. The results indicate that this hybrid chart improves the ability to detect process shifts with sensitivity to both small and large changes by creating a balance between the performance of the Shewhart and EWMA charts.

The hybrid Modified Exponentially Weighted Moving Average - Moving Average (MMEM) control chart for monitoring changes in the process mean was proposed by Talordphop et al. [14]. The results based on Monte Carlo simulation for normal and non-normal quality characteristic distributions, as well as an application to real data (chemical process temperature and survival of cancer patients), showed that this chart performs better than

other control charts in detecting small to moderate shifts. Subsequently, a new hybrid control chart called the Exponentially Weighted Moving Average - Modified Exponentially Weighted Moving Average (MEME) was proposed by Talordphop et al. [15], which is designed to monitor and quickly detect small to moderate shifts in the process mean. Monte Carlo simulation results and practical applications showed that this control chart has better performance for small to moderate shifts.

In this paper, the design of a new hybrid control chart is investigated under the condition that the quality characteristic follows a normal distribution. The structure of the paper is as follows: Section 2 is explained the control charts presented in existing studies; meanwhile, in this section, a new proposed control chart is introduced and examined. The performance of the proposed control chart compared to other existing control charts is evaluated based on the probability approach in Section 3. In Section 4, we evaluate the performance of the existing control charts against our proposed control chart using a numerical example with simulated data. Finally, the findings of this paper are summarized and presented in Section 5.

2 Control Charts

In this section, existing control charts including MA, EWMA, MEWMA, MMEM, and MEME are reviewed. Subsequently, a new proposed control chart will be introduced.

First, the design of the MA control chart is examined. The MA control chart is a type of weighted average control chart used for processes with subgroup data. Assume that the quality characteristic follows a normal distribution with mean μ_0 and variance σ^2 . Let

$$\bar{X}_i = (X_{i1} + X_{i2} + \dots + X_{in})/n,$$

be the mean of the i th subgroup of size n . Then, the MA statistic with width w at the i th stage is calculated as follows (Montgomery [11]):

$$\text{MA}_i = \frac{\bar{X}_i + \bar{X}_{i-1} + \dots + \bar{X}_{i-w+1}}{w}. \quad (2.1)$$

The control limits for the MA chart are obtained as:

$$\text{UCL/LCL} = \begin{cases} \mu_0 \pm H \frac{\sigma_{\bar{x}}}{\sqrt{i}}, & i < w; \\ \mu_0 \pm H \frac{\sigma_{\bar{x}}}{\sqrt{w}}, & i \geq w, \end{cases} \quad (2.2)$$

where H is the control constant of the MA chart and $\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{n}}$.

Next, the EWMA control chart is examined. The EWMA control chart was introduced by Roberts [3] and is suitable for detecting small shifts in the process. This chart is based on the following statistic:

$$Z_i = \lambda \bar{X}_i + (1 - \lambda)Z_{i-1}, \quad i = 1, 2, \dots, \quad (2.3)$$

where $0 < \lambda \leq 1$ is a constant, $Z_0 = \mu_0$ and $\bar{X}_i$ is the mean of the quality characteristic at time i . Assuming that the observations are independent and normally distributed, the control limits for the EWMA chart are calculated as:

$$\text{UCL/LCL} = \mu_0 \pm H \sigma_{\bar{x}} \sqrt{\left(\frac{\lambda}{2 - \lambda}\right)}, \quad (2.4)$$

where H is the control constant of the EWMA chart.

In the following, the MEWMA control chart proposed by Khan et al. [9] is introduced. Suppose that the observation X_i at time i follows a normal distribution. The MEWMA statistic is defined as:

$$M_i = \lambda X_i + (1 - \lambda)M_{i-1} + k(X_i - X_{i-1}), \quad i = 1, 2, \dots, \quad (2.5)$$

where M_i denotes the MEWMA statistic at time i , $0 < \lambda \leq 1$ is a constant, and k is set to $-\frac{\lambda}{2}$ in this paper. The control limits for the MEWMA chart are derived as follows:

$$\text{UCL/LCL} = \mu_0 \pm H \sigma_{\bar{x}} \sqrt{\frac{\lambda + 2\lambda k + 2k^2}{2 - \lambda}}, \quad (2.6)$$

where H is the control constant of the MEWMA chart.

We explain the hybrid MMEM control chart, which is a combination of the MEWMA and MA control charts in the following. This control chart was proposed by Talordphop et al. [14] and is suitable for detecting small shifts in the process. The statistic of the MMEM control chart is defined as:

$$S_i = \lambda M A_i + (1 - \lambda)S_{i-1} + k(M A_i - M A_{i-1}), \quad i = 1, 2, \dots, \quad (2.7)$$

where $M A_i$ is obtained from equation (2.1), $0 < \lambda \leq 1$ is a constant, and $k = -\frac{\lambda}{2}$. Therefore, the control limits of the MMEM chart are obtained as follows:

$$\text{UCL/LCL} = \mu_0 \pm H \sqrt{\left(\frac{\sigma_{\bar{x}}^2}{w}\right) \left(\frac{\lambda + 2\lambda k + 2k^2}{2 - \lambda}\right)}, \quad (2.8)$$

where H is the control constant of the MMEM chart.

Next, the MEME control chart proposed by Talordphop et al. [15] is presented. This control chart is a hybrid of the EWMA and MEWMA control charts. The statistic of the MEME control chart is defined as:

$$A_i = \lambda M_i + (1 - \lambda)A_{i-1}, \quad i = 1, 2, \dots, \quad (2.9)$$

where M_i is the MEWMA statistic. Therefore, the control limits of the MEME chart are calculated as:

$$\text{UCL/LCL} = \mu_0 \pm H \sigma_{\bar{x}} \sqrt{\left(\frac{\lambda + 2\lambda k + 2k^2}{2 - \lambda}\right) \left(\frac{\lambda}{2 - \lambda}\right)}, \quad (2.10)$$

where H is the control constant of the MEME chart.

As follow, the proposed hybrid control chart, called MAMEE, is introduced. This chart is a combination of the MA, MEWMA, and EWMA control charts. It is based on the following statistic:

$$D_i = \lambda S_i + (1 - \lambda)D_{i-1}, \quad i = 1, 2, \dots, \quad (2.11)$$

where S_i is obtained from equation (2.7) and $0 < \lambda \leq 1$ is a constant. Decision-making is performed by comparing each control chart's statistic with its own control limits. If the statistic falls between these limits, the process is considered in-control; otherwise, the process is deemed out-of-control and corrective action is required. To determine these limits, the mean and variance of the D_i statistic are first calculated under the in-control condition (assuming a mean of μ_0 for the data). The expectation and variance of the control chart are obtained as follows:

$$E(D_i) = \mu_0, \quad \text{Var}(D_i) = \frac{\sigma_x^2}{w} \left(\frac{\lambda}{2 - \lambda} \right) \left(\frac{\lambda + 2\lambda k + 2k^2}{2 - \lambda} \right),$$

where $k = -\frac{\lambda}{2}$. The control limits of the MAMEE chart are calculated as:

$$\text{UCL/LCL} = \mu_0 \pm H \sqrt{\left(\frac{\sigma_x^2}{w} \right) \left(\frac{\lambda}{2 - \lambda} \right) \left(\frac{\lambda + 2\lambda k + 2k^2}{2 - \lambda} \right)}. \quad (2.12)$$

The control constant H is determined based on the allowable Type I error or in-control average run length (ARL_0). The probability of correctly identifying the in-control process when the mean is μ_0 is:

$$P_{in}^0 = P(LCL \leq D_i \leq UCL | \mu_0) = \Phi(H) - \Phi(-H) = 2\Phi(H) - 1. \quad (2.13)$$

The ARL_0 depends on the constant H and is calculated as:

$$ARL_0 = \frac{1}{1 - P_{in}^0}. \quad (2.14)$$

If the process shifts to $\mu_1 = \mu_0 + c\sigma$, the probability of incorrectly identifying the shifted process as in-control is:

$$P_{in}^1 = P(LCL \leq D_i \leq UCL | \mu_1).$$

Assume that the process shifts from μ_0 to $\mu_1 = \mu_0 + c\sigma$. After calculating the expectation and variance of the statistic D_i , the probability that the shifted process is detected as in-control is $P_{in}^1 = P(LCL \leq D_i \leq UCL | \mu_1)$, which can be calculated based on the cumulative distribution function of the standard normal distribution. The out-of-control average run length for the shifted process is $ARL_1 = \frac{1}{1 - P_{in}^1}$.

3 Performance Evaluation of Control Charts

To properly evaluate the efficiency of control charts, a suitable and reliable criterion is needed. Therefore, the performance of the control charts in this paper is evaluated based

on the ARL criterion, which is the most widely used index for evaluating control charts (Montgomery [11]).

Table 1 presents the ARL_1 values for different values of widths $w = (5, 4, 3)$, constants $\lambda = (0.08, 0.15, 0.25)$, shift sizes $c = (0, 0.5)$ for $ARL_0 = 370$ and sample size $n = 1$ for the proposed control chart. For $ARL_0 = 370$, the control constant $H = 2.9997$, and for $ARL_0 = 300$, $H = 2.9352$ is determined. According to the data in these table, it is observed that smaller values of the constant λ detect process changes faster than larger values. Additionally, increasing the width w leads to faster detection of process changes. For example, when $ARL_0 = 370$, $\lambda = 0.1$, and $c = 0.01$, for $w = 3$, $ARL_1 = 20.7$, and for $w = 4$, $ARL_1 = 9.2$. To better understand this, the data from Table 1 are plotted in Figure 1, showing that appropriate parameter selection has a significant effect on reducing ARL_1 and consequently increasing the speed of detecting process changes.

Next, we intend to compare the performance of the proposed control chart with existing control charts. These comparisons are based on the probability approach, assuming that the data come from a normal distribution with parameters $\mu_0 = 0$ and $\sigma^2 = 1$, as shown in Table 2 and 3. These comparisons were performed for different values of $\lambda = 0.25$, $w = (5, 4)$, and $n = 1$.

The results in Table 2 and 3 show that for small shifts (e.g., $c = 0.001$ to $c = 0.01$), there is little difference between the control charts, and all have ARL_1 values close to 370. However, as the shift size increases, the performance of the control charts, especially the proposed MAMEE control chart, becomes more prominent in reducing ARL_1 compared to other control charts.

For shifts of $c = 0.5$ and more, the MAMEE control chart has the lowest ARL_1 values, indicating its high sensitivity to shifts. For example, for $c = 0.5$ with $w = 4$, the ARL_1 value for the MAMEE control chart is 1.00, while for the MA control chart it is approximately 43.86, for the EWMA control chart approximately 21.37, and for the MEWMA control chart approximately 17.72. This pattern is similar for $w = 5$. Hence, as the parameter w increases, the performance of the proposed control chart improves.

Furthermore, the comparison between $w = 4$ and $w = 5$ shows that, in general, increasing w reduces ARL_1 for all control charts, but this reduction is more noticeable for the proposed MAMEE control chart. For example, at $c = 0.1$ with $w = 4$, the ARL_1 value for the proposed MAMEE control chart is 15.06, which decreases to 10.83 when $w = 5$.

Overall, the results indicate that the proposed MAMEE control chart performs better than other existing control charts in detecting small shifts more quickly.

4 Numerical Example

In this section, the performance of the proposed control chart is evaluated using a numerical example consisting of 40 samples of size three ($n = 3$), the data of which are presented in Table 4. The first 20 samples in this table are in-control conditions following a standard normal distribution ($\mu_0 = 0$, $\sigma^2 = 1$). In contrast, for the next 20 samples (samples 21 to 40), a shift of 0.1σ is applied to the mean to simulate out-of-control conditions. To evaluate the performance of the control charts, the values $ARL_0 = 370$, $w = 3$, and

Table 1: The ARL_1 values of the proposed MAMEE control chart when $ARL_0 = 370$

Shift (c)	$w = 3$			$w = 4$			$w = 5$		
	$\lambda = 0.08$	$\lambda = 0.15$	$\lambda = 0.25$	$\lambda = 0.08$	$\lambda = 0.15$	$\lambda = 0.25$	$\lambda = 0.08$	$\lambda = 0.15$	$\lambda = 0.25$
0	370.03	370.03	370.03	370.03	370.03	370.03	370.03	370.03	370.03
0.001	366.78	369.14	369.73	365.71	368.84	369.63	364.64	368.54	369.52
0.005	301.68	348.73	362.52	283.70	342.11	360.08	267.53	335.71	357.66
0.008	231.27	319.57	351.33	203.92	305.40	345.47	181.69	292.30	339.78
0.01	188.17	296.29	341.54	159.19	277.27	332.91	137.02	260.31	324.65
0.015	108.76	234.65	311.05	84.35	207.55	294.98	67.73	185.42	280.33
0.02	63.33	178.93	275.73	46.02	149.98	253.28	35.14	128.12	233.84
0.03	23.75	100.63	205.46	15.87	77.24	176.77	11.43	61.52	154.29
0.04	10.38	57.39	147.85	6.70	41.27	119.98	4.78	31.27	99.85
0.05	5.27	33.95	105.51	3.44	23.27	81.50	2.52	17.04	65.23
0.07	2.05	13.47	54.80	1.50	8.76	39.23	1.25	6.23	29.62
0.09	1.26	6.31	29.94	1.09	4.09	20.32	1.03	2.96	14.79
0.1	1.12	4.58	22.60	1.03	3.01	15.06	1.01	2.24	10.83
0.25	1.00	1.00	1.01	1.00	1.00	1.00	1.00	1.00	1.00
0.5	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00

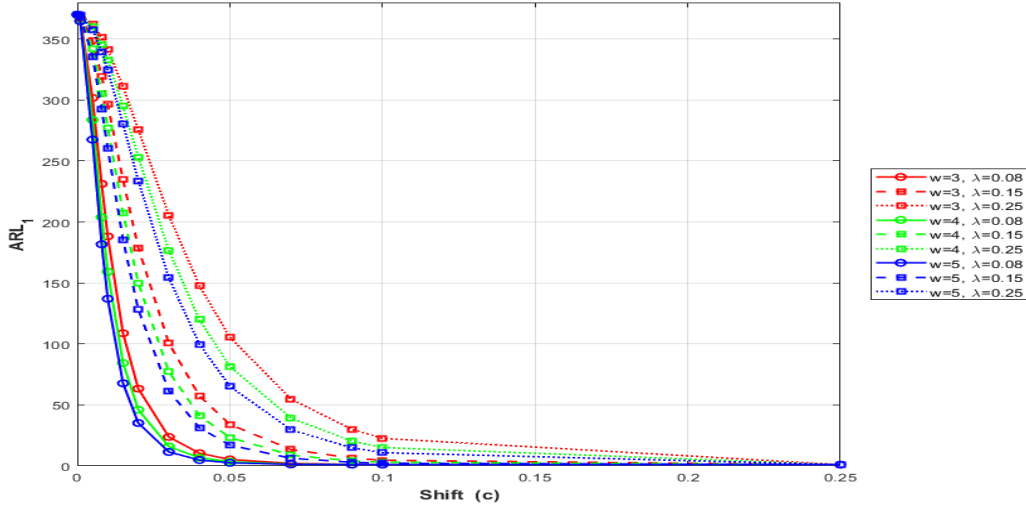


Figure 1: ARL_1 curves of the proposed MAMEE control chart for different values of parameters $\mu_0 = 0$, $\sigma^2 = 1$, $ARL_0 = 370$, $w = 3$, and $\lambda = 0.25$.

$\lambda = 0.25$ are considered. As observed in Figure 2, in response to the applied shift, the MA, EWMA, and MEWMA control charts were unable to detect it and all reported an in-control status. In contrast, the proposed MAMEE control chart demonstrated better performance. This chart successfully detected the shift in the process at sample 31 and issued an out-of-control signal; whereas the MMEM and MEME control charts detected this shift at samples 33 and 32, respectively. This result indicates that the proposed MAMEE control chart has higher sensitivity for detecting shifts in the process mean compared to other control charts under study, and is capable of detecting these shifts earlier. Given this feature, the MAMEE control chart can be used as an efficient and sensitive tool in quality control systems.

Table 2: ARL_1 values for the proposed MAMEE control chart with existing control charts for parameters $ARL_0 = 370$, $w = 4$, and $\lambda = 0.25$.

Shift (c)	$w = 4, \lambda = 0.25$					
	MA	EWMA	MEWMA	MMEM	MEME	MAMEE
0	370.03	370.03	370.03	370.03	370.03	370.03
0.001	370.03	370.02	370.02	369.98	369.93	369.63
0.005	369.85	369.72	369.67	368.58	367.50	360.08
0.008	369.57	369.22	369.10	366.34	363.61	345.47
0.01	369.31	368.76	368.58	364.28	360.08	332.91
0.015	368.40	367.18	366.78	357.32	348.29	294.98
0.02	367.14	364.99	364.28	347.96	332.91	253.28
0.03	363.58	358.87	357.32	323.50	294.98	176.77
0.04	358.69	350.59	347.96	294.01	253.28	119.98
0.05	352.59	340.44	336.56	262.49	212.94	81.50
0.07	337.18	315.79	309.19	201.89	145.80	39.23
0.09	318.41	287.49	278.34	151.48	98.78	20.32
0.1	308.13	272.74	262.49	130.76	81.50	15.06
0.25	155.09	101.91	90.57	17.72	7.72	1.30
0.5	43.86	21.37	17.72	2.31	1.30	1.00
0.75	14.96	6.45	5.27	1.12	1.00	1.00
1	6.30	2.76	2.31	1.00	1.00	1.00
1.25	3.24	1.61	1.42	1.00	1.00	1.00
1.5	2.00	1.20	1.12	1.00	1.00	1.00

Table 3: ARL_1 values for the proposed MAMEE control chart with existing control charts for parameters $ARL_0 = 370$, $w = 5$, and $\lambda = 0.25$.

Shift (c)	$w = 5, \lambda = 0.25$					
	MA	EWMA	MEWMA	MMEM	MEME	MAMEE
0	370.03	370.03	370.03	370.03	370.03	370.03
0.001	370.02	370.02	370.02	369.96	369.93	369.52
0.005	369.81	369.72	369.67	368.22	367.50	357.66
0.008	369.45	369.22	369.10	365.42	363.61	339.78
0.01	369.12	368.76	368.58	362.87	360.08	324.65
0.015	367.99	367.18	366.78	354.26	348.29	280.33
0.02	366.42	364.99	364.28	342.81	332.91	233.84
0.03	362.00	358.87	357.32	313.47	294.98	154.29
0.04	355.96	350.59	347.96	279.22	253.28	99.85
0.05	348.45	340.44	336.56	243.91	212.94	65.23
0.07	329.77	315.79	309.19	179.61	145.80	29.62
0.09	307.48	287.49	278.34	129.56	98.78	14.79
0.1	295.47	272.74	262.49	109.88	81.50	10.83
0.25	133.05	101.91	90.57	12.82	7.72	1.13
0.5	33.38	21.37	17.72	1.77	1.30	1.00
0.75	10.76	6.45	5.27	1.04	1.00	1.00
1	4.49	2.76	2.31	1.00	1.00	1.00
1.25	2.39	1.61	1.42	1.00	1.00	1.00
1.5	1.57	1.20	1.12	1.00	1.00	1.00

Figure 2: Performance comparison of detecting change using control charts with parameters $\mu_0 = 0$, $\sigma^2 = 1$, $ARL_0 = 370$, $w = 3$, and $\lambda = 0.25$.

Table 4: Simulated data from a standard normal distribution

i	x_1	x_2	x_3	X	MA_i	Z_i	M_i	S_i	A_i	D_i
1	-1.001	1.150	-0.343	-0.065	-0.065	-0.016	-0.008	-0.008	-0.002	-0.002
2	0.133	-1.243	0.427	-0.228	-0.146	-0.069	-0.043	-0.032	-0.012	-0.010
3	0.948	-0.666	1.802	0.695	0.134	0.122	0.026	-0.026	-0.003	-0.014
4	-0.958	-0.809	0.684	-0.361	0.035	0.001	0.062	0.002	0.013	-0.010
5	1.340	1.812	-0.974	0.726	0.353	0.182	0.092	0.050	0.033	0.005
6	0.265	-0.055	1.552	0.587	0.318	0.284	0.233	0.121	0.083	0.034
7	-0.992	0.874	-2.077	-0.731	0.194	0.030	0.157	0.155	0.101	0.064
8	-2.014	-0.806	-0.730	-1.183	-0.442	-0.273	-0.122	0.085	0.046	0.070
9	0.526	-0.058	1.123	0.531	-0.461	-0.072	-0.173	-0.049	-0.009	0.040
10	-0.367	0.280	0.376	0.096	-0.186	-0.030	-0.051	-0.118	-0.020	0.000
11	-0.816	-0.175	0.990	-0.001	0.209	-0.023	-0.027	-0.085	-0.021	-0.021
12	0.827	0.851	0.441	0.707	0.267	0.160	0.068	-0.004	0.001	-0.017
13	1.052	-0.262	-0.169	0.207	0.304	0.171	0.165	0.068	0.042	0.004
14	-0.424	2.630	-0.731	0.492	0.468	0.252	0.211	0.148	0.085	0.040
15	1.108	-0.323	-0.403	0.127	0.275	0.220	0.236	0.204	0.122	0.081
16	-0.032	-0.683	-0.115	-0.277	0.114	0.096	0.158	0.201	0.131	0.111
17	-0.829	-0.936	1.029	-0.245	-0.132	0.011	0.053	0.149	0.112	0.121
18	0.537	-0.354	-0.927	-0.248	-0.257	-0.054	-0.022	0.063	0.078	0.106
19	-1.372	0.101	-0.232	-0.501	-0.332	-0.166	-0.110	-0.026	0.031	0.073
20	-0.528	0.241	0.683	0.132	-0.206	-0.091	-0.129	-0.087	-0.009	0.033
21	-1.371	-1.036	0.572	-0.612	-0.327	-0.221	-0.156	-0.132	-0.046	-0.008
22	0.074	0.122	0.006	0.067	-0.137	-0.149	-0.185	-0.157	-0.080	-0.045
23	1.350	0.482	0.320	0.717	0.058	0.068	-0.041	-0.128	-0.071	-0.066
24	0.029	0.032	1.368	0.476	0.420	0.170	0.119	-0.036	-0.023	-0.058
25	-1.247	1.926	-1.027	-0.116	0.359	0.098	0.134	0.071	0.016	-0.026
26	0.408	-2.579	0.379	-0.597	-0.079	-0.076	0.011	0.088	0.015	0.002
27	0.212	1.298	1.231	0.914	0.067	0.172	0.048	0.064	0.023	0.018
28	-0.096	0.495	1.295	0.565	0.294	0.270	0.221	0.093	0.073	0.037
29	0.419	-0.200	-0.523	-0.101	0.459	0.177	0.224	0.164	0.110	0.069
30	-0.181	1.290	-0.175	0.311	0.258	0.211	0.194	0.213	0.131	0.105
31	2.360	0.909	0.379	1.216	0.475	0.462	0.336	0.251	0.182	0.141
32	-1.799	1.616	0.591	0.136	0.554	0.380	0.421	0.317	0.242	0.185
33	-1.111	0.840	-0.453	-0.241	0.370	0.225	0.303	0.353	0.257	0.227
34	0.403	0.338	0.268	0.336	0.077	0.253	0.239	0.321	0.253	0.251
35	-0.510	-1.961	1.187	-0.428	-0.111	0.083	0.168	0.237	0.231	0.247
36	1.269	-0.743	-1.129	-0.201	-0.098	0.012	0.047	0.151	0.185	0.223
37	-0.176	-0.124	0.437	0.046	-0.195	0.020	0.016	0.077	0.143	0.187
38	-1.009	-0.041	0.311	-0.247	-0.134	-0.046	-0.013	0.017	0.104	0.144
39	0.143	-0.567	1.056	0.211	0.003	0.018	-0.014	-0.004	0.074	0.107
40	-0.338	-0.621	-0.295	-0.418	-0.151	-0.091	-0.037	-0.021	0.047	0.075

5 Conclusion

In modern industry and factories, monitoring the quality of production processes is of particular importance. Control charts, as one of the main tools in this field, enable the rapid detection of changes and quality deviations. These tools play a significant role in reducing defective products and improving the efficiency of production lines. In this study, a new control chart for monitoring quality characteristics with a normal distribution has been introduced and evaluated. The proposed control chart is designed by combining three control charts: Moving Average (MA), Exponentially Weighted Moving Average (EWMA), and Modified Exponentially Weighted Moving Average (MEWMA). The control limits of this chart are determined based on the Shewhart method. The performance of the proposed control chart has been evaluated in comparison with existing control charts. The comparison criterion in this paper is the Average Run Length (ARL). The findings of the paper indicate a significant superiority of the proposed control chart over existing control charts. This control chart demonstrates better performance in the rapid and accurate detection of out-of-control conditions. The main advantage of the proposed control chart is its higher sensitivity to small shifts in the process. The accuracy and speed of detection of this control chart for small process shifts have been significantly improved.

Accordingly, the proposed control chart can be a suitable option for monitoring processes that are sensitive to small shifts.

References

- [1] Shewhart, W. A. (1931), *Economic control of quality of manufactured product*, D. Van Nostrand Company, New York.
- [2] Page, E. S. (1954), *Continuous inspection schemes*, *Biometrika*, **41**(1-2), 100–115.
- [3] Roberts, S. W. (1959), *Control chart tests based on geometric moving averages*, *Technometrics*, **1**(3), 239–250.
- [4] Khoo, M. B. C. (2004), *A moving average control chart for monitoring the fraction non-conforming*, *Quality and Reliability Engineering International*, **20**(6), 617–635.
- [5] Khoo, M. B. C. and Wong, V. H. (2008), *A double moving average control chart*, *Communications in Statistics - Simulation and Computation*, **37**(8), 1696–1708.
- [6] Patel, A. K. and Divecha, J. (2011), *Modified exponentially weighted moving average (EWMA) control chart for an analytical process data*, *Journal of Chemical Engineering and Materials Science*, **2**(1), 12–20.
- [7] Lucas, J. M. (1982), *Combined Shewhart-CUSUM quality control schemes*, *Journal of Quality Technology*, **14**(2), 51–59.
- [8] Lucas, J. M. and Saccucci, M. S. (1990), *Exponentially weighted moving average control schemes: Properties and enhancements*, *Technometrics*, **32**(1), 1–12.
- [9] Khan, N., Aslam, M. and Jun, C. H. (2016), *An EWMA control chart for exponentially distributed quality based on moving average statistics*, *Quality and Reliability Engineering International*, **32**, 1179–1190.
- [10] Aslam, M., Gui, W., Khan, N. and Jun, C. H. (2017), *Double moving average-EWMA control chart for exponentially distributed quality*, *Communications in Statistics - Simulation and Computation*, **46**(9), 7351–7364.
- [11] Montgomery, D. C. (2009), *Introduction to Statistical Quality Control*, John Wiley & Sons, Hoboken, NJ.
- [12] Aslam, M., Khan, N., Azam, M. and Jun, C. H. (2014b), *Designing of a new monitoring T-chart using repetitive sampling*, *Information Sciences*, **269**, 210–216.
- [13] Taboran, R. and Sukparungsee, S. (2022), *On designing of a new EWMA-DMA control chart for detecting mean shifts and its application*, *Thailand Statistician*, **21**(1), 148–164.
- [14] Talordphop, K., Sukparungsee, S. and Areepong, Y. (2022), *New modified exponentially weighted moving average-moving average control chart for process monitoring*, *Connection Science*, **34**(1), 1981–1998.

- [15] Talordphop, K., Sukparungsee, S. and Areepong, Y. (2023), *On designing new mixed modified exponentially weighted moving average - exponentially weighted moving average control chart*, Results in Engineering, **18**, 101152.
- [16] Shanshool, M. K., Mahadik, S. B., Godase, D. G. and Khoo, M. B. C. (2024), *The combined Shewhart-EWMA sign charts*, Quality and Reliability Engineering International, **40**(7), 4111–4122.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Comparative Evaluation of Group Replacement and Run to Failure Policies in Heterogeneous Multi Component Systems

Mostafavi Vala, F. S.¹ Rasay, H.² Najafi-Ghobadi, S.³

^{1,2,3} Department of Industrial Engineering, Kermanshah University of Technology,
Kermanshah, Iran

Abstract

In multi component maintenance planning, economies of scale can typically be achieved through group maintenance strategies. This study focuses on a heterogeneous multi unit system and formulates two cost minimizing maintenance policies. Using renewal theoretic reasoning, the expected number of failures over a finite interval is numerically estimated, forming a key element of the models. The analysis also shows how one policy reduces to alternative strategies under specific conditions. A numerical example is provided to demonstrate the models' performance and the resulting cost expressions.

Keywords: Maintenance, Renewal theory, Replacement, Group Replacement Policy, Run to Failure Policy.

¹lidamostafavivala@gmail.com

²Speaker. Hasan.Rasay@gmail.com

³s.najafighobadi@kut.ac.ir

1 Introduction and literature review

Maintenance comprises the technical and administrative activities required to restore or sustain system operability [4] , [1]. Maintenance strategies have evolved from run to failure to time based, condition based, and more recently predictive maintenance, which aims to intervene shortly before failure [9]. Maintenance costs include direct expenses such as labor and spare parts, and indirect costs associated with downtime and lost sales. Mathematical models seek to optimize objectives such as expected cost, availability, and reliability.

For multi component systems, single unit models are often insufficient due to economic, structural, stochastic, and resource dependencies among units [12], [3]. Economic dependencies arise mainly from setup costs -e.g., system shutdown and technician deployment- that motivate group maintenance policies. Structural dependencies occur when components must be disassembled to access others, as in machine tools. Stochastic dependencies arise when one unit's failure influences the failure behavior of other units. Group replacement/maintenance (GRP) consolidates actions to exploit economies of scale, such as in street lamp maintenance. Prior studies have analyzed various forms of these policies, load sharing systems under economic dependencies via predictive maintenance policies , clustering policies with significant cost reductions [6] , reinforcement learning based predictive group maintenance [5] , multi agent approaches for production lines [11], dynamic and health informed group policies [10], multi stage degradation models [8], and digital twin supported opportunistic group maintenance [7]. This paper examines two maintenance policies for a multi unit system— Group Replacement policy (GR policy) and Run-to-failure policy (RTF Policy). Section 2 formulates the problem and derives expected cost expressions; Section 3 provides a numerical example; Section 4 concludes the study.

2 Problem description and modeling

This section studies a system of N heterogeneous components and formulates two maintenance policies. GR policy applies group replacement at fixed intervals, renewing all components simultaneously. RTF policy relies on run to failure replacement, renewing each component only when it fails. Parameters and notations used in this study are as follows.

N : The number of the system units,

c_j^g : The cost of replacing unit j 'th under group replacement,

c_j^f : The cost of replacing unit j 'th under failure replacement,

T_g : The time required to replace all items under group replacement policy,

t_p : The interval of group replacement policy for replacing the components,

$H_j(t)$: The expected number of failures in interval $(0,t)$ for item j ,

$f_j(t)$: probability density function (p.d.f) of item j for failure time,

ECT_j : Expected cost per time unit under policy j (where j = GR or RTF).

For both Policies, cost rate expressions are developed. A prerequisite for these formulations is computing $H_j(t)$, the expected number of failures of component j over the interval $(0, t)$. In general, the calculation of $H_j(t)$ follows renewal theoretic principles. As shown in [2], applying Laplace transforms yields the following relation:

$$H_j^*(s) = \frac{f_j^*(s)}{s[1 - f_j^*(s)]} \quad (2.1)$$

Then $H_j(t)$ can be computed from $H_j^*(s)$. In practice, inversion of $H_j^*(s)$ to compute $H_j(t)$ can only be done for some simple cases. Accordingly, discrete approaches are commonly employed which leads to Equation 2.2

$$H_j(T) = \sum_{k=0}^{T-1} [1 + H_j(T - k - 1)] \int_k^{k+1} f_j(t) dt, \quad T \geq 1 \quad (2.2)$$

Since $H_j(0) = 0$. Eq 2.2 allows recursive evaluation, where $H_j(1)$ is obtained from $H_j(0)$, then $H_j(2)$ from $H_j(1)$, and so forth. In this discrete scheme, it is assumed that no more than one failure occurs within a single time unit. This assumption is not limited, as the time step can be chosen arbitrarily small.

2.1 GR Policy

As previously noted, a system with N heterogeneous items is considered where each item has its own failure rate and replacement cost. Under the policy examined here, failed items are replaced upon failure during the interval t_p , with item j incurring a corrective replacement cost c_j^f . At the end of each interval t_p , all items are preventively renewed simultaneously, and the preventive replacement cost for item j is c_j^p , with the typical assumption $c_j^p \leq c_j^f$. The goal is to determine an optimal value of t_p that minimizes the long run expected replacement cost per unit time. The system's cost rate function is computed using Equation 2.3.

$$ECT_{GR}(t_p) = \frac{\sum_{j=1}^N c_j^g + \sum_{j=1}^N H_j(t_p) c_j^f}{t_p + T_g} \quad (2.3)$$

2.2 RTF Policy

In this policy, each item is replaced immediately upon failure, representing a conventional run to failure approach. This policy can also be viewed as the limiting form of the GR policy when the preventive interval t_p grows without bound. As $t_p \rightarrow \infty$, the GR policy converges to the RTF policy. For large t (as $t \rightarrow \infty$), the expected number of failures $H_j(t)$ is obtained using the following expression from [2]:

$$H_j(t) \approx \frac{t}{\mu_j} + \frac{\sigma_j^2 - \mu_j^2}{2\mu_j^2} \quad (2.4)$$

Consequently, the following limit must be evaluated:

$$ECT_{GR}(t_p) = \lim_{t_p \rightarrow \infty} \frac{\sum_{j=1}^N c_j^g + \sum_{j=1}^N H_j(t_p) c_j^f}{t_p + T_g} \quad (2.5)$$

By inserting Equation 4 into Equation 5, we arrive at Equation 2.6:

$$ECT_{RTF} = \sum_{j=1}^N \frac{c_j^f}{\mu_j} \quad (2.6)$$

3 Numerical study

Consider a system composed of three components, whose parameter values are given in Table 1. Assume that the duration of a group replacement is negligible. All remaining inputs also appear in Table 1.

Table 1: The data of the units

unit	c_j^g	c_j^f	$f_j(t)$
unit1	20	150	Normal(20,5)
unit2	15	100	Normal(15,3)
unit3	10	80	Normal(10,3)

Table 2 summarizes the outcomes obtained for the two maintenance policies. Under the GR policy, the optimal preventive replacement interval is $t_p = 7$, which yields an expected cost rate (ECT_{GR}) of 8.397. If the RTF policy is used instead, Equation 6 indicates that the corresponding theoretical ECT_{RTF} is approximately $\frac{150+100+80}{20+15+10} \approx 22$.

Table 2: The data of the units

Policy	ECT_j	t_p
GR policy	8397	7
RTF policy	22	∞

These findings imply that, following the GR policy, all three components should be jointly replaced at time 7, whereas within the interval (0,7) any component that fails is replaced upon failure. As anticipated, the RTF policy results in the highest cost.

4 Conclusion

This study examined two maintenance policies for heterogeneous multi unit systems, focusing on the underlying concepts of group replacement. Closed form expressions were developed to estimate the expected cost rate associated with each policy, and it was shown that specific policy can simplify to others under particular conditions. Numerical example is solved to evaluate and contrast the performance of the considered policies. From a managerial standpoint, the results highlight that integrating group replacement maintenance can substantially improve cost efficiency.

References

- [1] Ahmad, R., and Kamaruddin, S. (2012). An overview of time-based and condition-based maintenance in industrial application. *Computers & industrial engineering*, **63**(1), 135-149.
- [2] Andrew K.S.Jardin and Albert H.c. tsang. (2013), *Maintenance, Replacement and Reliability; Theory and Applications, Second*, aylor & Francis Group.
- [3] Azizi, F., Rezvani, Z., Rasay, H., Safari, A., Salmani, M., and Naderkhani, F. (2025). Smart Maintenance Optimization for Large Scale Parallel Systems Using Deep Reinforcement Learning. *IEEE Transactions on Reliability*, **75**, 260-272.
- [4] Ding, S. H., and Kamaruddin, S. (2015). Maintenance policy optimization—literature review and directions. *The International Journal of Advanced Manufacturing Technology*, **76**(5), 1263-1283.
- [5] Fallahnezhad, M. S., Bazeli, S., and Rasay, H. (2022). Clustering of condition-based maintenance activities with imperfect maintenance and predication signals. *International Journal of Industrial Engineering*, **33**(4), 1-19.
- [6] Keizer, M. C. O., Teunter, R. H., Veldman, J., and Babai, M. Z. (2018). Condition-based maintenance for systems with economic dependence and load sharing. *International journal of production economics*, **195**, 319-327.
- [7] Liu, J. (2025). Group maintenance based on deep hierarchical reinforcement learning and multi-stage degradation analysis. *Applied Soft Computing*, **189**, 114516.
- [8] Ouahabi, N., Chebak, A., Kamach, O., and Zegrari, M. (2025). Dynamic production scheduling and maintenance planning under opportunistic grouping. *Computers & Industrial Engineering*, **199**, 110646.
- [9] Surucu, O., Gadsden, S. A., and Yawney, J. (2023). Condition monitoring using machine learning: A review of theory, applications, and recent advances. *Expert Systems with Applications*, **221**, 119738.
- [10] Yang, L., Zhou, S., Ma, X., Chen, Y., Jia, H., and Dai, W. (2024). Group machinery intelligent maintenance: Adaptive health prediction and global dynamic maintenance decision-making. *Reliability Engineering & System Safety*, **252**, 110426.
- [11] Zeng, J., and Liang, Z. (2024). Predictive group maintenance using probabilistic prognostics and deep reinforcement learning. *Computers & Industrial Engineering*, **212**, 111738.
- [12] Zhang, N., and Si, W. (2020). Deep reinforcement learning for condition-based maintenance planning of multi-component systems under dependent competing risks. *Reliability Engineering & System Safety*, **203**, 107094.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Optimal Weibull Reliability group test plans using truncated prior distribution

Naghizadeh Qomi, M.¹ Mahdizadeh, Z.² Pérez-González, Carlos J.³

^{1,2} Department of Statistics, University of Mazandaran, Babolsar, Iran

³ Investigación Operativa and the Instituto de Matemáticas y Aplicaciones (IMAULL), Universidad de La Laguna (ULL), 38271 La Laguna, Tenerife, Canary Islands, Spain

Abstract

This paper focuses on designing reliability group sampling plans based on Weibull failure count data. A linear combination of expected producer and consumer risks is minimized and the optimal number of groups are also determined by solving integer nonlinear programming problems where the prior information on failure probability is described by a truncated beta model. Numerical and graphical analyses are provided for illustrative and comparative purposes.

Keywords: Reliability testing, Expected producer and consumer risks, Lot sampling inspection, Weibull model.

1 Introduction

Reliability acceptance sampling plans in quality control are very important in reliability analysis and industrial manufacturing to decide the acceptability of submitted lots based on random samples and prior information. The construction of reliability sampling plans

¹Speaker. m.naghizadeh@umz.ac.ir

²mahdizade_zohre@yahoo.com

³cpgonzal@ull.edu.es

has been studied in a large variety of situations. Some recent references are [16], [12], [5], [8], [10] and [9].

Group acceptance sampling plans (*GASPs*) by attributes represent the most basic schemes in industrial quality control. The formulation of *GASPs* enable the simultaneous evaluation of multiple items and thereby achieve considerable reductions in both time and cost. These plans are constructed on g groups, yielding a total sample size of $n = kg$, where each group contains k items. The k items in a group are tested simultaneously on separate testers for a pre-specified time t . The experiment is terminated if, during this period, the number of failures in any group exceeds the allowable acceptance number c .

The methodology of *GASP* has been widely treated literature. The lifetime variable of interest follows generalized exponential distribution in [14] and [3]. More recently, we can cite the works by [15], [1], [11], [6] and [7] which have introduced *GASPs* based on inverse log-logistic, gamma-Lindley, type-I heavy-tailed Rayleigh, odd-Perks-Lomax and Weibull distributions.

In designing *GASPs* it is usual to assume a constant probability p of observing a defective or nonconforming unit in the lot under inspection. In many circumstances, however, this assumption is not reasonable and it should be convenient to adopt an adequate prior model on p . This paper deals with the determination of optimal *GASPs* with controlled expected weighted-average of risks and fixed acceptance numbers using Weibull failure count data and truncated beta distributions, which is essentially a constrained minimization problem. The optimal number of groups are found by solving a mixed integer nonlinear programming problem. Using available prior information on p in the sampling plan design can reduce the required sample size for testing.

The rest of this paper is structured as follows. The methodology of *GASP* is proposed in Section 2. Prior models on the defective probability, as well as classical and expected producer and consumer risks, are defined in section 3. A mixed integer nonlinear programming problem is formulated in Section 4 to find the optimal *GASPs* with minimal expected weighted-average risks.

2 Weibull reliability group sampling plans

Assume that the lifetime of a product follows a Weibull distribution with known shape parameter $s > 0$ and scale parameter $\theta > 0$ with cumulative distribution function (cdf)

$$F(t; s, \theta) = 1 - \exp \left(- \left(\frac{t}{\theta} \right)^s \right), \quad t > 0. \quad (2.1)$$

The mean life of a Weibull distribution product is given by $\mu = (\theta/s)\Gamma(1/s)$, where $\Gamma(\cdot)$ is the complete gamma function. Let μ_0 be the specified mean life and the termination time is given by $t_0 = q\mu_0$, where q is a positive constant known as termination ratio.

Following [2], the procedure of a *GASP* denoted by $\mathcal{S} = (g, c)$ is as follows:

1. Select the number of groups g and allocate pre-defined k units to each group so that the sample size for a lot will be $n = gk$.
2. Select the acceptance number c ($< k$) for a group and the fixed experiment time t_0 .

3. Perform the experiment for the g groups simultaneously and count the number of defectives ($\mathcal{D}_j$, $j = 1, 2, \dots, g$) for each group.
4. Accept the lot whenever $\mathcal{D}_j$ is at most equal to acceptance number c in each of all groups, otherwise, truncate the experiment and reject the lot.

The operating characteristic (OC) function gives the probability of accepting a lot versus the true proportion of defective items p . The OC function associated with the group sampling plan $\mathcal{S} = (g, c)$ is defined at p by $\mathcal{L}(p) \equiv \mathcal{L}(p; g, c)$ and is given by

$$\mathcal{L}(p) = Pr(\mathcal{D}_j \leq c, \text{ for all } g \text{ groups}) = \left[\sum_{i=0}^c \binom{k}{i} p^i (1-p)^{k-i} \right]^g, \quad (2.2)$$

where p denotes the probability that a product in a group fails before the time t_0 and is given by

$$p \equiv p(q, s, r) = 1 - \exp \left(- \left(\frac{q\Gamma(1/s)}{sr} \right)^s \right), \quad (2.3)$$

where the ratio $r = \mu/\mu_0$ is expressed as the quality level of a product.

3 Prior models and expected risks

In general, conventional risks are typically used in cases where the prior distribution of p is either unavailable or has been discarded. Suppose that there is some information about the defective rate, p .

Let r_0 and r_1 be the ratio of quality level associated with the producer and consumer, respectively. Moreover, $p_0 \equiv p(q, s, r_0)$ and $p_1 \equiv p(q, s, r_1)$ are the acceptable and rejectable defective rates corresponding to the producer and consumer, respectively. Conventionally, the lot is acceptable for producer if $p \leq p_0$ and it is rejectable for consumer if $p \geq p_1$. The producer and consumer risks (CPR and CCR) are defined respectively as

$$\sup_{p \leq p_0} \{1 - \mathcal{L}(p)\} \quad \text{and} \quad \sup_{p \geq p_1} \{\mathcal{L}(p)\}.$$

Since $\mathcal{L}(p)$ is a decreasing function of p , then the CPR and CCR are given by $CPR(g, c, p_0) = 1 - \mathcal{L}(p_0)$ and $CCR(g, c, p_1) = \mathcal{L}(p_1)$.

In the construction of *GASPs*, the frequentist approach relies solely on empirical observations and assumes a fixed probability p of encountering a defective or nonconforming unit within the lot under inspection. However, this assumption is often unrealistic, as prior information regarding the stochastic nature of the defective rate p is frequently available and can be incorporated into the design of sampling plans. According to [4], expected producer and consumer risks (EPR and ECR) are defined by

$$EPR(g, c, p_0) = E[1 - \mathcal{L}(p) | p \leq p_0] \quad \text{and} \quad ECR(g, c, p_1) = E[\mathcal{L}(p) | p \geq p_1],$$

respectively. The EPR is defined as $EPR = 1 - \mathcal{L}(p_0)$ when $P(p \leq p_0) = 0$ and the ECR is defined as $ECR = \mathcal{L}(p_1)$ when $P(p \geq p_1) = 0$, which coincides with the corresponding conventional producer and consumer risks. Otherwise, the EPR and ECR can be expressed respectively as

$$EPR(n, c, p_0) = 1 - \int_0^{p_0} \frac{\mathcal{L}(p)h_0(p)}{Pr(p \leq p_0)} dp,$$

and

$$ECR(n, c, p_1) = \int_{p_1}^1 \frac{\mathcal{L}(p)h_1(p)}{Pr(p \geq p_1)} dp,$$

where $h_0(\cdot)$ and $h_1(\cdot)$ are the prior producer and consumer density functions. That is, $h_0(\cdot)$ and $h_1(\cdot)$ are the corresponding prior densities of the conditional random variables $p|p \leq p_0$ and $p|p \geq p_1$. This paper considers the case in which both producer and consumer want to ensure the neutrality of the prior model and they agree with the selected prior distributions, $h(p) = h_0(p) = h_1(p)$.

The use of a Beta prior distribution for the proportion of defectives p in acceptance sampling procedures is well justified from both theoretical and practical perspectives. We assume truncated beta priors for the producer and consumer. In general, the defective rate p conditional to $p \in [l, u]$ follows a truncated beta $TB(a, b, l, u)$ prior distribution with pdf as

$$h(p) = \frac{p^{a-1}(1-p)^{b-1}}{\mathcal{B}(a, b, l, u)}, \quad a, b > 0, \quad 0 < l < u < 1, \quad l < p < u,$$

where $\mathcal{B}(a, b, l, u) = \int_l^u x^{a-1}(1-x)^{b-1} dx$ is a truncated beta function. Observe that the standard $Beta(a, b)$ distribution is a particular case of the TB distribution when $l = 0$ and $u = 1$. We consider the prior producer and consumer distributions are as $TB(a, b, 0, p_0)$ and $TB(a, b, p_1, 1)$, respectively. That is, $p|p \leq p_0 \sim TB(a, b, 0, p_0)$ and $p|p \geq p_1 \sim TB(a, b, p_1, 1)$.

4 Optimal *GASP* with minimal expected weighted-average of risks

The expected weighted-average risk (EWR) is defined as a linear combination of the expected producer and consumer risks, expressed through their respective weights. If λ_0 and λ_1 denote positive constants corresponding to the producer's and consumer's weights, respectively, with the constraint $\lambda_0 + \lambda_1 = 1$, then the EWR can be written as

$$EWR(g, c, p_0, p_1) = \lambda_0 EPR(g, c, p_0) + \lambda_1 ECR(g, c, p_1). \quad (4.1)$$

The ratio λ_0/λ_1 quantifies the relative importance assigned to producer error in comparison with consumer error. This formulation has the inherent advantage of adapting to the sample size, making it particularly relevant in acceptance sampling where the conflicting interests of producers and consumers must be balanced.

Given a fixed acceptance number $c \geq 0$, our interest is to determine the best number of groups, g^* , which can be obtained by minimizing (4.1). The constrained minimization problem can be stated as follows:

$$\begin{aligned} & \text{Minimize} && EWR(g, c, p_0, p_1) \\ & \text{Subject to} && (g, c) \in \Omega, \end{aligned} \tag{4.2}$$

where $\Omega = \{(g, c) : g \in \mathbb{N}, c \in \mathbb{N}_0\}$ is the feasible region. Therefore, we have to solve the following optimization problem

$$\min\{EWR(g, c, p_0, p_1) : (g, c) \in \Omega\}.$$

Clearly, $EWR(g^*, c, p_0, p_1) = \min\{EWR(g, c, p_0, p_1) : (g, c) \in \Omega\}$ and the inspection scheme with minimal EWR would be denoted by the $\mathcal{S}^* = (g^*, c)$.

Tables 1 and 2 presents the number of groups with minimum EWR , g^* , and their associated risks (EWR , EPR and ECR) for selected values of $c = 0(1)4$, $\lambda_0 = 0.2, 0.5, 0.8$, $k = 5$, $q = 0.5, 1.0$, $r_0 = 2$ and $r_1 = 1$ when $s = 1$ and $s = 2$, respectively. The prior distribution of the defective rate p is assumed to be a $Beta(a, b)$ with mode $5(p_0 + p_1)/6$ and $a + b = 5$. Since the mode of $Beta(a, b)$ is $(a - 1)/\{(a + b) - 2\}$, then we get

$$a \equiv a_{p_0, p_1} = 1 + 5(p_0 + p_1)/2, \text{ and } b \equiv b_{p_0, p_1} = 4 - 5(p_0 + p_1)/2. \tag{4.3}$$

In light of Tables 1 and 2, it is clear that:

1. The minimum- EWR number of groups, g^* , grows as acceptance number c increases while the EWR and EPR decrease. .
2. g^* reduces when the termination ratio q increases..
3. When λ_0 increases, g^* and the EPR decrease, whereas the ECR grows.

Figure 1 displays the EWR , EPR and ECR percentages versus λ_0 for $c = 2$, $r_0 = 2$, $r_1 = 1$, $k = 5$, $q = 0.5$ when $s = 1$. It is observed that the EPR reduces and the ECR increases when λ_0 increases while the EWR first increases and then decreases.

Figure 2 shows the minimum conventional weighted-average risks (CWR) and minimum EWR number of groups versus λ_0 for $c = 2$, $r_0 = 2$, $r_1 = 1$, $k = 5$, $q = 0.5$ and $s = 1$. In light of this figure, it is clear that the use of expected risks instead of conventional (classical) risks, in most cases, yields appreciate reductions in g^* .

5 Real data application

In this section, a data set is employed to illustrate the proposed *GASP* to assess the appropriateness of the Weibull distribution. The following data set reported in [13] represents the lifetimes (in minutes) of breakdown voltages of electrical cable insulation

0.19, 0.78, 0.96, 1.31, 2.78, 3.16, 4.15, 4.67, 4.85, 6.50,
7.35, 8.01, 8.27, 12.06, 31.75, 32.52, 33.91, 36.71, 72.89

Table 1: Optimal number of groups, g^* , with minimum-EWR and the corresponding risks (%) for selected values of c , $\lambda_0 = 0.2, 0.5, 0.8, r_0 = 2, r_1 = 1, k = 5$ and $q = 0.5, 1$ when $s = 1$.

q	c	$(\lambda_0, \lambda_1) = (0.2, 0.8)$				$(\lambda_0, \lambda_1) = (0.5, 0.5)$				$(\lambda_0, \lambda_1) = (0.8, 0.2)$			
		g^*	EWR	EPR	ECR	g^*	EWR	EPR	ECR	g^*	EWR	EPR	ECR
0.5	0	1	12.65	55.42	1.96	1	28.69	55.42	1.96	1	44.73	55.42	1.96
	1	2	8.32	32.25	2.34	1	14.78	18.10	11.46	1	16.77	18.10	11.46
	2	6	4.88	18.59	1.45	4	8.61	13.00	4.23	2	8.41	6.82	14.77
	3	29	2.52	9.90	0.68	18	4.47	6.35	2.60	10	4.51	3.61	8.13
	4	299	1.20	4.77	0.31	196	2.17	3.17	1.16	112	2.26	1.84	3.98
1.0	0	1	16.45	81.86	0.10	1	40.98	81.86	0.10	1	65.51	81.86	0.10
	1	1	10.60	47.80	1.30	1	24.55	47.80	1.30	1	38.50	47.80	1.30
	2	2	7.34	32.32	1.10	1	12.75	18.15	7.36	1	15.99	18.15	7.36
	3	5	4.38	17.86	1.01	3	7.78	11.30	4.25	2	8.12	7.74	9.62
	4	25	2.25	8.89	0.59	16	4.04	5.84	2.25	9	4.18	3.35	7.46

Table 2: Optimal number of groups, g^* , with minimum-EWR and the corresponding risks (%) for selected values of c , $\lambda_0 = 0.2, 0.5, 0.8, r_0 = 2, r_1 = 1, k = 5$ and $q = 0.5, 1$ when $s = 2$.

q	c	$(\lambda_0, \lambda_1) = (0.2, 0.8)$				$(\lambda_0, \lambda_1) = (0.5, 0.5)$				$(\lambda_0, \lambda_1) = (0.8, 0.2)$			
		g^*	EWR	EPR	ECR	g^*	EWR	EPR	ECR	g^*	EWR	EPR	ECR
0.5	0	2	7.24	24.82	2.84	1	13.18	13.47	12.89	1	13.35	13.47	12.89
	1	11	2.36	9.43	0.59	7	4.26	6.16	2.36	4	4.42	3.59	7.76
	2	79	0.65	2.67	0.15	56	1.23	1.91	0.55	35	1.37	1.20	2.04
	3	1011	0.17	0.70	0.03	760	0.33	0.53	0.13	528	0.39	0.37	0.47
	4	-	-	-	-	-	-	-	-	-	-	-	-
1.0	0	1	10.20	49.16	0.47	1	24.81	49.16	0.47	1	39.42	49.16	0.47
	1	2	5.20	24.68	0.34	1	8.80	13.40	4.21	1	11.56	13.40	4.21
	2	4	1.97	8.03	0.46	3	3.75	6.11	1.39	2	4.22	4.13	4.57
	3	15	0.62	2.60	0.13	11	1.22	1.91	0.53	8	1.44	1.40	1.63
	4	109	0.16	0.68	0.03	86	0.32	0.53	0.10	63	0.40	0.39	0.41

Dashes (-) are presented in the required cells for a large number of groups.

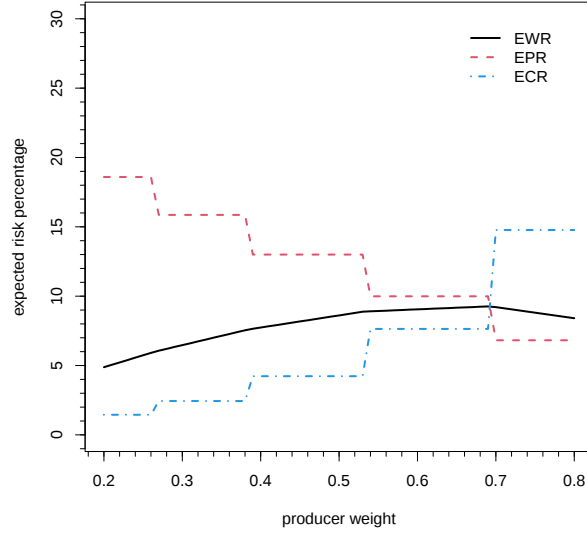


Figure 1: EWR, EPR and ECR percentages versus the producer weight λ_0 when $c = 2$, $s = 1$, $r_0 = 2$, $r_1 = 1$, $k = 5$ and $q = 0.5$.

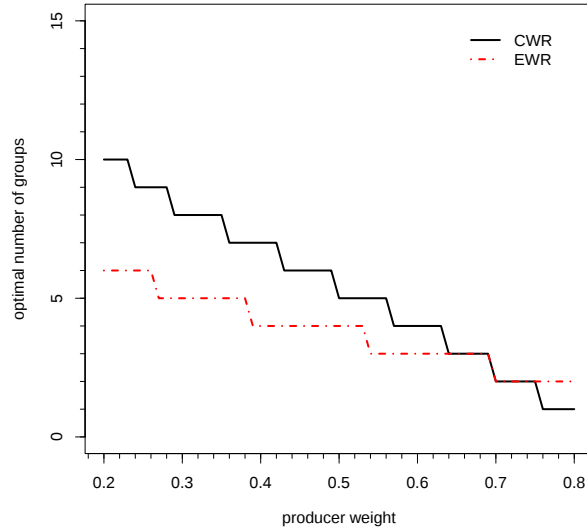


Figure 2: Optimal number of groups with minimum CWR and EWR the producer weight λ_0 when $c = 2$, $s = 1$, $r_0 = 2$, $r_1 = 1$, $k = 5$ and $q = 0.5$.

The maximum likelihood estimates of the parameters for this data set are $\hat{s} = 0.7708$ and $\hat{\theta} = 12.2222$. The suitability of data for the Weibull distribution was assessed using Kolmogorov–Smirnov (K-S) test and the corresponding p-value. The K-S test statistic 0.1613 with the associated significance value 0.6482 shows that the Weibull distribution has a good fit to considered data set.

Using the fact $\mu = (\theta/s)\Gamma(1/s)$, the estimated mean life of data is $\mu_0 = 14.24$. Assum-

Table 3: Optimal number of groups, g^* , for selected values of $c = 2$, $\lambda_0 = 0.5$, $r_0 = 2, 3, 4, 5$, $r_1 = 1$, $k = 5$ and $q = 0.5, 1$ when $s = 0.7708$.

r_0	$q = 0.5$					$q = 1.0$				
	p_0	p_1	a	b	g^*	p_0	p_1	a	b	g^*
2	0.3205	0.4828	3.0081	1.9919	2	0.4828	0.6753	3.8951	1.1049	1
3	0.2462	0.4828	2.8225	2.1775	3	0.3826	0.6753	3.6448	1.3552	1
4	0.2026	0.4828	2.7135	2.2865	3	0.3205	0.6753	3.4894	1.5106	1
5	0.1736	0.4828	2.6409	2.3591	4	0.2777	0.6753	3.3825	1.6175	1

ing a termination ratio as $q = 0.5$, the termination time is obtained as $t_0 = q\mu_0 = 7.12$ minutes.

Assume that the failure probability, p , has a $Beta(a, b)$ distribution with $a + b = 5$ and mode $5(p_0 + p_1)/6$. From Table 3, for $r_1 = 1$, $r_0 = 2$, $q = 0.5$, $k = 5$, $c = 2$ and $\lambda_0 = 0.5$, we obtain $p_0 = 0.3205$ and $p_1 = 0.4828$, $a = 3.0081$ and $b = 1.9919$. Moreover, the optimal number of groups is $g^* = 2$. According to these specifications, a total of 10 products are needed and that ten items will be allocated to each of the two groups. We will accept the lot if two failures occur before $t_0 = 7.12$ minutes in each of the two groups.

6 Concluding remarks

A Bayesian approach was presented in this paper in order to determine the optimal number of groups with fixed acceptance number to decide the acceptability of a submitted lot. Assuming a weighted-average of the expected producer and consumer risks (EWR) selected by the analyst, the optimal number of groups is obtained by minimizing the EWR. Constrained optimization problems have been formulated with the aim of determining the minimum-EWR c -defective plan (g^*, c) . The results show that using prior information for fraction defective increases the effectiveness of the proposed plans. A real data analysis is provided to illustrate the results by comparison of the existence method and the proposed method. The results derived for Weibull distribution are also valid for any other lifetime models.

Suppose that the analyst wants to control the risk incurred by the selected $GASP$ by considering $\varepsilon \in (0, 1)$ as the maximum risk tolerated, where $\varepsilon \leq \min\{\lambda_0, \lambda_1\}$. Proposed optimal $GASPs$ in this paper can be developed to determine the optimal number of groups and acceptance number when the maximum tolerated weighted-average risk is ε , i.e., $EW R(g, c, p_0, p_1) \leq \varepsilon$.

Acknowledgements

This research was financed by a research grant from the University of Mazandaran.

References

- [1] Al-Omari, A.I. and Ismail, M.T. (2024), Group acceptance sampling plans based on truncated life tests for gamma Lindley distribution with real data application, *Lobachevskii Journal of Mathematics*, **45**, 578-590.
- [2] Aslam, M. and Jun, C.-H. (2009), group acceptance sampling plan for truncated life test having Weibull distribution, *Journal of Applied Statistics*, **36(9)**, 1021-1027.
- [3] Aslam, M. Kundu, D., Jun, C.-H. and Ahmad, M. (2011), Time truncated group acceptance sampling plans for generalized exponential distribution, *Journal of Teting and Evaluation*, **39(4)**, 671-677.
- [4] Fernández, A.J., Correa-Álvarez, C.D. and Pericchi, L.R. (2020), Balancing producer and consumer risks in optimal attribute testing: A unified Bayesian/Frequentist design, *European Journal of Operational Research*, **211**, 576-587.
- [5] Fernández, A.J., Pérez-González, C.J., Carrión-García, A. and Giner-Bosch, V. (2023), Overdispersion effects on reliability test planning, *IEEE Transactions on Reliability*, **72(3)**, 1053-1063.
- [6] Naghizadeh Qomi, M., Al-Husseini, Z., and Aslam, M. (2025), Improved group sampling plans under odd-Perks-Lomax model with limited risks. *Electronic Journal of Applied Statistical Analysis*, **18(2)**, 344-358.
- [7] Naghizadeh Qomi, M. and Aslam, M. (2026), Developing optimal group acceptance sampling plans based on Weibull distribution with limited risks, *Journal of Applied Statistics*, **53(7)**, 1237-1252.
- [8] Naghizadeh Qomi, M. and Fernández, A.J. (2024), Optimal double acceptance sampling plans based on truncated life tests for Tsallis q-exponential distributions, *Journal of Applied Statistics*, **51**, 3456-3467.
- [9] Naghizadeh Qomi, M. and Mahdizadeh, Z. (2025), Time censored reliability sampling plans for inverted Nadarajah-Haghighi distribution with minimal and limited risks, *Quality and Reliability Engineering International*, **41(6)**, 2729-2739.
- [10] Naghizadeh Qomi, M., Piadeh, M., Pérez-González, C. J., and Fernández, A. J. (2025), Optimal acceptance sampling plans based on generalized half-normal distribution under time-censoring with conventional and expected limited risks, *Communications in Statistics-Simulation and Computation*, 1-20, doi:10.1080/03610918.2025.2512380.
- [11] Nwankwo, M.P., Alsadat, N., Kumar, A., Bahloul, M.M. and Obulezi, O.J. (2024), Group acceptance sampling plan based on truncated Llife tests for type-I heavy-tailed Rayleigh distribution, *Heliyon*, **10(19)**, e38150.

- [12] Pérez-González, C.J., Fernández, A.J., Giner-Bosch, V. and Carrión-García, A. (2023), Optimal repetitive reliability inspection of manufactured lots for lifetime models using prior information, *International Journal of Production Research*, **61**, 2214-2230.
- [13] Qutb, N. and Rajhi, E. (2016), Estimation of the parameters of compound Weibull distribution, *Journal of Mathematics*, **12(4)**, 11-18.
- [14] Rao, G.S. (2009), A group acceptance sampling plans for lifetimes following a generalized exponential distribution, *Economic Quality Control*, **24(1)**, 75-85.
- [15] Tripathi, H., Dey, S. and Saha, M. (2021), Double and group acceptance sampling plan for truncated life test based on inverse log-logistic distribution, *Journal of Applied Statistics*, **48(7)**, 1227-1242.
- [16] Tripathi, H., Saha, M. and Halder, S. (2023), Single acceptance sampling inspection plan based on transmuted Rayleigh distribution, *Life Cycle Reliability and Safety Engineering*, **12**, 111-123.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Stress vs. Strength: A Log-Uniform Model for Quantifying Reliability in Cognitive Performance and Battery Degradation

Rezaei, Amir ¹ Rezaei, Ahmad ²

¹ Department of Statistics, Faculty of Science, Bu-Ali Sina University, Hamedan, Iran

² Department of Educational Technology, Faculty of Educational Sciences and Psychology, Kharazmi University, Tehran, Iran

Abstract

This paper develops and compares estimators of the stress-strength parameter $R = P(X > Y)$ when both variables follow independent log-uniform distributions. We derive the maximum likelihood estimator (MLE), the uniformly minimum variance unbiased estimator (UMVUE), and a novel shrinkage estimator that optimally blends MLE and UMVUE. A simulation study confirms that the shrinkage estimator reduces bias and mean squared error compared to the MLE while approaching the UMVUE in large samples. The practical value of the methodology is then demonstrated through two real-world case studies. First, in an educational psychology context, we quantify the effect of coffee consumption on students' accuracy in a statistics exam, obtaining a reliability estimate of approximately 0.35. Second, in an electrical engineering setting, we analyse lithium-ion battery degradation under thermal stress, yielding a reliability estimate of about 0.56. In both applications, the log-uniform model provides an adequate fit, and the estimators deliver interpretable, actionable insights that align with domain-specific literature. The paper includes compact derivations of the key analytical forms, making it self-contained for applied researchers.

¹Speaker. a.rezaei@basu.ac.ir

²ahmadrezaei@khu.ac.ir

Keywords: Stress-strength model, Log-uniform distribution, Shrinkage estimator, Coffee and academic performance, Battery thermal reliability.

1 Introduction

The notion of “stress” and “strength” is not confined to mechanics. In psychology, stress refers to any cognitive or emotional demand that challenges an individual’s performance; in electrical engineering, it can be an elevated temperature that degrades a component. Across these vastly different domains, a central question recurs: What is the probability that the system’s inherent strength overcomes the applied stress? Answering this question in a principled, quantitative way is the task of the **stress-strength reliability model**.

The stress-strength model was originally formalized by Birnbaum [3] and defines the reliability parameter $R = P(X > Y)$, where X represents the “strength” random variable and Y the “stress” random variable. The parameter R has since found extensive application in reliability engineering, medicine, psychology, and education [8]. An important practical challenge is that X and Y are rarely known with precision, and the analyst must choose a probability distribution that reflects the available knowledge. When only vague information about the scale or range of the variables exists—often the case with behavioral scores or capacity measurements—the log-uniform distribution is an attractive candidate.

The log-uniform distribution, introduced by Hamming [5], has the probability density function (pdf)

$$f_Z(z; b) = \frac{1}{z \ln b}, \quad \frac{1}{b} \leq z \leq 1, \quad (1.1)$$

and the cumulative distribution function (cdf)

$$F_Z(z; b) = 1 + \frac{\ln z}{\ln b}, \quad \frac{1}{b} \leq z \leq 1, \quad (1.2)$$

where $b > 1$ is a scale parameter. It is continuous, right-skewed, and especially useful when the variables span several orders of magnitude or when a non-informative prior for a scale parameter is desired [2, 4, 9, 11, 12]. Its bounded support on $[\frac{1}{b}, 1]$ makes it particularly suitable for modelling rates, proportions, and normalized capacities, which appear in our applications.

A random variable Z is said to follow a log-uniform distribution with parameter b , denoted by $Z \sim LU(b)$, if its pdf is given by (1.1).

In this paper, we put the applications at the core while developing the necessary statistical machinery. Specifically, our contributions are threefold. First, we derive the explicit expression for the stress-strength parameter R when both $X \sim LU(b_1)$ and $Y \sim LU(b_2)$ are independent. Second, we construct three point estimators: the maximum likelihood estimator (MLE), the uniformly minimum variance unbiased estimator (UMVUE), and a novel shrinkage estimator that optimally blends MLE and UMVUE. Compact proofs of the key formulas are provided to keep the paper self-contained. Third, and most importantly, we apply these estimators to two real-world case studies that represent the model’s cross-disciplinary reach:

- **Educational Psychology:** We examine whether drinking two cups of coffee before a statistics exam improves students' accuracy. Here, the strength X is the accuracy score of coffee-drinking students, and the stress Y is the accuracy of non-drinking students. The estimated R yields a direct probabilistic answer to the question: "How likely is it that a coffee drinker outperforms a non-drinker?"
- **Electrical Engineering:** We analyze the degradation of lithium-ion batteries stored at 50°C (stress) relative to those stored at 25°C (strength). The resulting R quantifies the reliability penalty imposed by thermal stress, a critical metric for battery management systems.

2 The stress-strength parameter

In this section, we compute the stress-strength parameter for the log-uniform distributions. Let $X \sim LU(b_1)$ and $Y \sim LU(b_2)$ be two independent random variables, where b_1 and b_2 are unknown parameters. The stress-strength parameter can be expressed as

$$R = P(X > Y) = \int_{\max\{\frac{1}{b_1}, \frac{1}{b_2}\}}^1 F_Y(x; b_2) f_X(x; b_1) dx.$$

Define $a^* = \max\{\frac{1}{b_1}, \frac{1}{b_2}\}$. For $x \in [a^*, 1]$, substituting the pdf and cdf from (1.1) and (1.2), respectively, gives

$$R = \int_{a^*}^1 \left(1 + \frac{\ln x}{\ln b_2}\right) \frac{1}{x \ln b_1} dx = \frac{1}{\ln b_1} \int_{a^*}^1 \frac{1}{x} dx + \frac{1}{\ln b_1 \ln b_2} \int_{a^*}^1 \frac{\ln x}{x} dx.$$

The integrals evaluate to $-\ln a^*$ and $-\frac{1}{2}(\ln a^*)^2$, respectively. Hence

$$R = -\frac{\ln a^*}{\ln b_1} - \frac{(\ln a^*)^2}{2 \ln b_1 \ln b_2}.$$

Now if $\frac{1}{b_1} < \frac{1}{b_2}$ then $a^* = \frac{1}{b_2}$ and $-\ln a^* = \ln b_2$, yielding $R = \frac{\ln b_2}{2 \ln b_1}$. And, if $\frac{1}{b_1} \geq \frac{1}{b_2}$ then $a^* = \frac{1}{b_1}$ and $-\ln a^* = \ln b_1$, giving $R = 1 - \frac{\ln b_1}{2 \ln b_2}$. Therefore

$$R = \begin{cases} \frac{\ln b_2}{2 \ln b_1} & \text{if } \frac{1}{b_1} < \frac{1}{b_2}, \\ 1 - \frac{\ln b_1}{2 \ln b_2} & \text{if } \frac{1}{b_1} \geq \frac{1}{b_2}. \end{cases} \quad (2.1)$$

As derived in (2.1), the stress-strength parameter R is a function of the unknown parameters b_1 and b_2 , and is denoted by $R = R(b_1, b_2)$.

3 ESTIMATORS OF R

In this section, different estimators of R are developed using the methods of MLE, UMVUE, and shrinkage estimation. Assume independent samples $X_1, X_2, \dots, X_n \sim LU(b_1)$ and $Y_1, Y_2, \dots, Y_m \sim LU(b_2)$. The estimators are functions of the minimum order statistics $X_{(1)} = \min X_i$ and $Y_{(1)} = \min Y_j$.

3.1 MLE

The likelihood function for the LU sample $Z_1, Z_2, \dots, Z_N$ is

$$L(b) = \prod_{i=1}^N \frac{1}{z_i \ln b} I_{[\frac{1}{b}, 1]}(z_i) = \left(\frac{1}{\ln b}\right)^N \left(\prod_{i=1}^N \frac{1}{z_i}\right) I_{[\frac{1}{z_{(1)}}, +\infty]}(b) I_{[-\infty, 1]}(z_{(N)}),$$

where $z_{(j)}$ denotes the j th order statistic. Since $L(b)$ is decreasing in b on $[\frac{1}{z_{(1)}}, \infty)$, the MLE of b is $\hat{b} = \frac{1}{z_{(1)}}$. Applying this to both samples and using the invariance property of MLEs, we obtain the MLE of R as

$$\hat{R}_{MLE} = R(\hat{b}_1, \hat{b}_2) = \begin{cases} \frac{\ln Y_{(1)}}{2 \ln X_{(1)}} & \text{if } X_{(1)} < Y_{(1)}, \\ 1 - \frac{\ln X_{(1)}}{2 \ln Y_{(1)}} & \text{if } X_{(1)} \geq Y_{(1)}. \end{cases} \quad (3.1)$$

Exact moments of $\hat{R}_{MLE}$ can be derived and are used in the shrinkage estimator. The first and second moments of $\hat{R}_{MLE}$ are as follows:

$$E(\hat{R}_{MLE}) = \frac{nm}{2(n+m)(m+1)} \frac{\ln b_2}{\ln b_1} + \frac{m}{n+m} - \frac{nm}{2(n+m)(n+1)} \frac{\ln b_1}{\ln b_2}, \quad (3.2)$$

$$\begin{aligned} E(\hat{R}_{MLE}^2) &= \frac{nm}{4(n+m)(m+2)} \left(\frac{\ln b_2}{\ln b_1}\right)^2 + \frac{m}{n+m} - \frac{nm}{(n+m)(n+1)} \frac{\ln b_1}{\ln b_2} \\ &\quad + \frac{nm}{4(n+m)(n+2)} \left(\frac{\ln b_1}{\ln b_2}\right)^2. \end{aligned} \quad (3.3)$$

3.2 UMVUE

Minimum order statistics $X_{(1)}$ and $Y_{(1)}$ are complete sufficient for b_1 and b_2 . Let $h(X_{(1)}, Y_{(1)})$ be an unbiased estimator of R . Setting $E[h(X_{(1)}, Y_{(1)})] = R$ and using the joint density of the minima,

$$f_{X_{(1)}, Y_{(1)}}(x, y) = \frac{nm}{xy} \left(\frac{1}{\ln b_1}\right)^n \left(\frac{1}{\ln b_2}\right)^m \left(\ln \frac{1}{x}\right)^{n-1} \left(\ln \frac{1}{y}\right)^{m-1},$$

with $\frac{1}{b_1} \leq x \leq 1$ and $\frac{1}{b_2} \leq y \leq 1$, the unbiasedness condition yields, after applying Leibniz integral rule, the unique solution

$$\hat{R}_{UMVUE} = \begin{cases} \frac{(n-1)(m+1)}{2nm} \frac{\ln Y_{(1)}}{\ln X_{(1)}} & \text{if } X_{(1)} < Y_{(1)}, \\ 1 - \frac{(n+1)(m-1)}{2nm} \frac{\ln X_{(1)}}{\ln Y_{(1)}} & \text{if } X_{(1)} \geq Y_{(1)}. \end{cases} \quad (3.4)$$

It should be noted that as $n, m \rightarrow \infty$, the coefficients $\frac{(n-1)(m+1)}{2nm} \rightarrow \frac{1}{2}$ and $\frac{(n+1)(m-1)}{2nm} \rightarrow \frac{1}{2}$, so $\hat{R}_{UMVUE} \rightarrow \hat{R}_{MLE}$, confirming asymptotic equivalence.

3.3 Shrinkage Estimator

The shrinkage estimator of R blends its MLE and UMVUE as

$$\hat{R}_{\text{sh}} = \theta \hat{R}_{MLE} + (1 - \theta) \hat{R}_{UMVUE},$$

with the optimal weight θ minimizing the mean squared error (MSE) of $\hat{R}_{\text{sh}}$ given by

$$\theta_{opt} = \frac{(R - R_0)(E(\hat{R}_{MLE}) - R_0)}{E(\hat{R}_{MLE}^2) + R_0^2 - 2R_0E(\hat{R}_{MLE})}.$$

where $R_0 = \hat{R}_{UMVUE}$ and the explicit moments $E(\hat{R}_{MLE})$ and $E(\hat{R}_{MLE}^2)$ are available in (3.2) and (3.3), respectively. Replacing R with $\hat{R}_{MLE}$ in θ_{opt} yields a feasible weight $\hat{\theta}^*$ and the final estimator

$$\hat{R}_{\text{sh}}^* = \hat{\theta}^* \hat{R}_{MLE} + (1 - \hat{\theta}^*) \hat{R}_{UMVUE}. \quad (3.5)$$

It can be seen that for large sample sizes n and m , the estimators $\hat{R}_{UMVUE}$, $\hat{R}_{MLE}$, and $\hat{R}_{\text{sh}}^*$ coincide.

4 SIMULATION STUDY

To evaluate the performance of the three estimators, a Monte Carlo simulation was conducted. Independent samples $X_1, \dots, X_n$ and $Y_1, \dots, Y_m$ were generated from $LU(b_1 = 1.5)$ and $LU(b_2 = 3)$, respectively, with $n, m \in \{25, 40, 75, 116\}$. These parameter values represent distinct stress and strength conditions, and the chosen sample sizes cover a practical range from small to moderately large datasets. Each scenario was replicated 100,000 times to ensure stable Monte Carlo estimates. According to (2.1), the true value of R is 0.8154649.

Table 1 reports the average bias (A.Bias) and MSE of the MLE, UMVUE, and the proposed shrinkage estimator of R across different sample sizes.

The following conclusions emerge clearly:

- **Consistency:** As n and m increase, the absolute bias and MSE of all three estimators decrease toward zero, confirming their consistency.
- **Bias ordering:** The UMVUE consistently exhibits the smallest absolute bias, followed by the shrinkage estimator, and then the MLE. That is,

$$|\text{A.Bias}(\hat{R}_{UMVUE})| < |\text{A.Bias}(\hat{R}_{\text{sh}})| < |\text{A.Bias}(\hat{R}_{MLE})|.$$

- **MSE ordering:** Similarly,

$$\text{MSE}(\hat{R}_{UMVUE}) \leq \text{MSE}(\hat{R}_{\text{sh}}) < \text{MSE}(\hat{R}_{MLE}).$$

In several cases, the shrinkage estimator matches the UMVUE in MSE.

Table 1: A.Bias and MSE of different estimators of R .

n	m	$\hat{R}_{MLE}$		$\hat{R}_{UMVUE}$		$\hat{R}_{sh}$	
		A.Bias	MSE	A.Bias	MSE	A.Bias	MSE
25	25	-0.00027488	0.00011027	0.00002080	0.00010984	0.00002099	0.00010984
	40	0.00025408	0.00007897	-0.00001705	0.00007452	0.00001840	0.00007456
	75	0.00025427	0.00007470	-0.00001531	0.00005587	-0.00001818	0.00005621
	116	0.00024726	0.00007097	-0.00000752	0.00005336	0.00001271	0.00005323
40	25	-0.00296273	0.00009018	0.00003722	0.00007874	0.00005675	0.00007882
	40	-0.00009788	0.00004248	0.00001750	0.00004242	0.00001753	0.00004242
	75	0.00005837	0.00003029	-0.00001138	0.00001665	-0.00001507	0.00002664
	116	0.00002924	0.00003006	-0.00000969	0.00001283	0.00000973	0.00002281
75	25	-0.00021921	0.00009602	0.00002764	0.00006514	-0.00003160	0.00006536
	40	-0.00021222	0.00003373	0.00002745	0.00002814	0.00003013	0.00002815
	75	-0.00003131	0.00001206	0.00001490	0.00001205	0.00001951	0.00001205
	116	0.00003101	0.00000907	-0.00001125	0.00000845	-0.00001261	0.00000845
116	25	-0.00020586	0.00010298	-0.00001164	0.00006185	0.00007068	0.00006214
	40	-0.00011835	0.00003563	-0.00000497	0.00002503	0.00001576	0.00002505
	75	-0.00008901	0.00000958	0.00000427	0.00000871	0.00000653	0.00000871
	116	-0.00000385	0.00000507	0.00000385	0.00000507	0.00000385	0.00000507

- **Shrinkage benefit:** The shrinkage estimator $\hat{R}_{sh}$ effectively reduces both bias and MSE relative to $\hat{R}_{MLE}$, leveraging the UMVUE as a stable anchor.
- **Convergence:** For the largest sample sizes ($n = m = 116$), all three estimators yield nearly identical A.Bias and MSE, supporting their asymptotic equivalence established theoretically.

These results provide empirical evidence that the proposed estimators, particularly the UMVUE and the shrinkage estimator, perform reliably in finite samples and can be confidently applied to real data.

5 APPLICATION 1 – EDUCATIONAL PSYCHOLOGY

We now turn to the first real-world problem: quantifying the cognitive benefit of caffeine in an academic setting. A longstanding question in educational psychology is whether moderate coffee consumption before an exam improves performance. To address this, we analyse a publicly available dataset (<https://github.com/mwaskom/Psych252/blob/master/WWW/datasets/caffeine.csv>) containing accuracy scores (probability of answering a question correctly) on a 10-question statistics exam for two independent groups of students: those who consumed two cups of coffee prior to the test (strength, $n = 20$) and those who consumed none (stress, $m = 11$). The data are summarized in Table 2.

The log-uniform distribution is a natural first model for such bounded proportions, particularly when the underlying cognitive scale is unknown. Kolmogorov-Smirnov (KS)

Table 2: Accuracy scores (potential mediator) of students who consumed either 0 or 2 cups of coffee prior to the statistics exam.

Cups of coffee	Accuracy scores			
2	0.4212120	0.3528293	0.5564924	0.4788779
	0.4047098	0.5580818	0.3838812	0.5772318
	0.6154113	0.6771771	0.2402377	0.6036828
	0.3460965	0.6179977	0.5452024	0.4069703
	0.3953451	0.6337638	0.4618152	0.4105491
0	0.4498767	0.4995339	0.4985903	0.4543119
	0.6262978	0.3501472	0.3988599	0.7308104
	0.6786753	0.6320787	0.6273589	—

tests indeed fail to reject the log-uniform hypothesis at the 5% level: the coffee group yields $\hat{b}_1 = 4.1625$ with $p = 0.082$, and the no-coffee group gives $\hat{b}_2 = 2.8559$ with $p = 0.229$.

Using (3.1), (3.4) and (3.5), we obtain

$$\hat{R}_{MLE} = 0.3679204, \quad \hat{R}_{UMVUE} = 0.3177494, \quad \hat{R}_{sh} = 0.3588493.$$

The estimates converge on a probability of roughly 0.32-0.37 that a coffee-drinking student attains a higher accuracy than a non-drinker. This moderate positive effect is consistent with experimental findings that caffeine enhances memory consolidation, especially in morning sessions [14], and that energy drink consumption reduces reaction time in young adults [1]. Additionally, survey-based studies report a positive association between caffeine intake and academic performance among medical students [6, 7]. The stress-strength parameter thus translates a diffuse body of evidence into a single, easily interpretable metric: drinking two cups of coffee before an exam is beneficial, but the advantage is far from certain.

6 APPLICATION 2 – ELECTRICAL ENGINEERING

In a completely different domain, we demonstrate the same methodology on battery reliability data in <https://calce.umd.edu/battery-data>. Lithium-ion cells degrade faster at elevated temperatures, a phenomenon of great practical concern for electric vehicles and grid storage. We use charge capacity measurements (Ah) taken over three weeks from cells stored at a mild 25°C (strength, $n = 16$) and a harsh 50°C (stress, $m = 16$); the full dataset is listed in Table 3. The question is: what is the probability that a battery kept at room temperature retains a higher capacity than one exposed to thermal stress?

Once again, the log-uniform distribution provides an adequate description. KS tests yield $p = 0.103$ for the 25°C group ($\hat{b}_1 = 239.63$) and $p = 0.214$ for the 50°C group ($\hat{b}_2 = 479.33$). The computed estimates are

$$\hat{R}_{MLE} = 0.5561605, \quad \hat{R}_{UMVUE} = 0.5578942, \quad \hat{R}_{sh} = 0.5578908.$$

Table 3: Charge capacity (Ah) data of lithium-ion batteries stored at two different temperatures: 25°C and 50°C.

Temperature	Charge Capacity			
25°C	0.060511487	0.200307983	0.139800208	0.006259688
	0.022952455	0.131453684	0.166926078	0.137713617
	0.171099205	0.004173094	0.010432821	0.064684697
	0.039645317	0.116847193	0.062598126	0.075117785
50°C	0.002086253	0.168988979	0.181507477	0.039639947
	0.098051629	0.054244232	0.006258799	0.177334650
	0.014603951	0.018776534	0.208630818	0.010431352
	0.131433696	0.143952122	0.112656079	0.123088079

All three estimators agree on a value near 0.56, meaning that a battery stored at 25°C has roughly a 56 % chance of maintaining a higher capacity than its 50°C counterpart.

This moderate but unequivocal reliability deficit matches electrochemical knowledge. Li et al. [10] reported that even minor thermal inhomogeneities (3°C) can accelerate degradation by more than 300 %, and Shen et al. [13] linked elevated temperatures to accelerated solid-electrolyte interface growth and capacity fade. The parameter $R \approx 0.56$ thus serves as a compact summary of the thermal penalty, providing battery management engineers with a probability-based figure of merit for reliability-conscious design.

7 CONCLUSION

This paper has demonstrated that the stress-strength model under log-uniform distributions delivers actionable insights in both cognitive psychology and battery engineering. We derived the exact form of R , constructed and compared three point estimators through simulation, and applied them to authentic datasets. The shrinkage estimator improves upon the MLE in finite samples, while the UMVUE remains the gold standard for unbiasedness. The estimated probabilities – ≈ 0.35 for coffee enhancing exam performance and ≈ 0.56 for battery capacity retention under thermal stress – are consistent with domain-specific literature. The framework is readily extendable to Bayesian inference and can be applied to other fields where stress-strength reasoning is relevant.

References

- [1] Aniței, M., Schuhfried, G. and Chraif, M. (2011). The influence of energy drinks and caffeine on time reaction and cognitive processes in young Romanian students. *Procedia-Social and Behavioral Sciences*, **30**, 662-670.
- [2] Av, A. and Sebastian, R. (2025). Some Applications of Exponentiated Log-Uniform Distribution. *Reliability: Theory & Applications*, **20(1 (82))**, 124-134.
- [3] Birnbaum, Z. W. (1956). On a use of the Mann-Whitney statistic. Proceedings of the Third Berkeley Symposium Math, *Statistic Probability I*, 13-17.

- [4] Escobar, M.T and Moreno-Jiménez, J. M. (2000). Reciprocal distributions in the analytic hierarchy process, *European Journal of Operational Research*, **123**(1), 154-174.
- [5] Hamming, R. W. (1970), On the distribution of numbers, *The Bell System Technical Journal*, **49**(8), 1609-1625.
- [6] Ibn Idris, H. O., Abukanna, A. M., Ibn Idris, H. O., Rasheed, M. and Alanazi, Z. A. N. (2022). Effect of caffeine on medical student's performance during exams in Northern Border University (NBU), Saudi Arabia, *Medical Science*, **26**(125), ms291e2239. <https://doi.org/10.54905/disssi/v26i125/ms291e2239>.
- [7] Khan, M. S., Nisar, N. and Naqvi, S. A. A. (2017). Caffeine consumption and academic performance among medical students of Dow university of health science (DUHS), Karachi, Pakistan. *Annals of Abbasi Shaheed Hospital and Karachi Medical & Dental College*, **22**(3), 179-184.
- [8] Kotz, S., Lumelskii, Y. and Pensky, M. (2003). *The Stress-Strength Model and its Generalizations: Theory and Applications*, Singapore, World Scientific Press.
- [9] Kreinovich, V., Kosheleva, O., Nguyen, H. T. and Sriboonchitta, S. (2015). Why Some Families of Probability Distributions Are Practically Efficient: A Symmetry-Based Explanation, *Departmental Technical Reports (CS)*, Paper 917.
- [10] Li, S., Zhang, C., Zhao, Y., Offer, G. J. and Marinescu, M. (2023). Effect of thermal gradients on inhomogeneous degradation in lithium-ion batteries. *Communications Engineering*, **2**(1), 74.
- [11] Nguyen, H. T., Kreinovich, V. and Longpré, L. (2004). Dirty pages of logarithm tables, lifetime of the Universe, and (subjective) probabilities on finite and infinite intervals, *Reliable Computing*, **10**(2), 83-106.
- [12] Shaji, A. K. and Sebastian, R. (2023). Some Applications of Transmuted Log-Uniform Distribution. *Reliability: Theory & Applications*, **18** (4 (76)), 1046-1055.
- [13] Shen, W., Wang, N., Zhang, J., Wang, F. and Zhang, G. (2022). Heat generation and degradation mechanism of lithium-ion batteries during high-temperature aging. *ACS omega*, **7**(49), 44733-44742.
- [14] Sherman, S. M., Buckley, T. P., Baena, E. and Ryan, L. (2016). Caffeine enhances memory performance in young adults during their non-optimal time of day. *Frontiers in psychology*, **7**, 1764.



The 12th Seminar on
Reliability Theory and its Applications
15 & 16 July 2026



Ordering Results of Series Systems Consisting of Dependent and Heterogeneous Exponential Components under Archimedean Copula

Omid Shojaee ¹

Department of Statistics, University of Zabol, Zabol, Sistan and Baluchestan, Iran

Abstract

This paper investigates series systems comprising m dependent subsystems, each containing heterogeneous and dependent Exponential components. The dependence among components is modelled using Archimedean copulas. We establish sufficient conditions for the stochastic comparison of two such series systems when the vector of scale parameters of the first system majorizes that of the second, and when the vectors representing the number of components in each subsystem differ between the two systems. A numerical example is provided to illustrate the theoretical findings.

Keywords: Series system, Usual stochastic order, heterogeneous Exponential components, Archimedean Copula.

1 Introduction

In many real-world applications, including reliability engineering, lifetime populations are often heterogeneous and composed of multiple homogeneous sub-populations. This heterogeneity can arise from factors such as variations in production lines, differences in raw materials, or inconsistencies in work shifts and labor quality. For a detailed discussion

¹Speaker. o_shojaee@uoz.ac.ir

on heterogeneous components, we refer the interested readers to Shojaee et al. [5] and Shojaee et al. [6].

This paper derives sufficient conditions for the stochastic comparison of two series systems with dependent and heterogeneous Exponential-distributed components. Specifically, we establish ordering results under the usual stochastic order when (i) the vector of scale parameters of the first system majorizes that of the second, and (ii) the vector representing the number of components in each subsystem also follows a majorization relationship.

The rest of the paper is organized as follows. Section 2 provides some basic definitions and lemma. Section 3 is devoted to the usual stochastic order of series systems. Section 4 concludes the paper.

2 Preliminaries

Let the random variables X and Y have survival functions $\bar{F}$ and $\bar{G}$, respectively.

Definition 2.1. It is said to be the random variable X is smaller than Y in the usual stochastic order if $\bar{F}(x) \leq \bar{G}(x)$ for all x , and denoted by $X \leq_{st} Y$ or $\bar{F} \leq_{st} \bar{G}$.

Definition 2.2. (Marshall et al. [2]). Let $a_{(1)} \leq \dots \leq a_{(n)}$ and $b_{(1)} \leq \dots \leq b_{(n)}$ denote the increasing arrangements of $\mathbf{a} = (a_1, \dots, a_n)$ and $\mathbf{b} = (b_1, \dots, b_n)$, respectively.

- (i) If $\sum_{j=1}^i a_{(j)} \leq \sum_{j=1}^i b_{(j)}$ for $i = 1, \dots, n-1$, and $\sum_{j=1}^n a_{(j)} = \sum_{j=1}^n b_{(j)}$, then $\mathbf{a}$ is said to majorize $\mathbf{b}$ and denoted by $\mathbf{a} \succeq^m \mathbf{b}$.
- (ii) If $\sum_{j=1}^i a_{(j)} \leq \sum_{j=1}^i b_{(j)}$ for $i = 1, \dots, n$, then $\mathbf{a}$ is said to weakly supermajorize $\mathbf{b}$, and denoted by $\mathbf{a} \succeq^w \mathbf{b}$.
- (iii) If $\sum_{j=i}^n a_{(j)} \geq \sum_{j=i}^n b_{(j)}$ for $i = 1, \dots, n$, then $\mathbf{a}$ is said to weakly submajorize $\mathbf{b}$, denoted by $\mathbf{a} \succeq_w \mathbf{b}$.

Let us consider the definition of Archimedean copula from Nelsen [4].

Definition 2.3. The copula function

$$C_\Psi(v_1, \dots, v_n) = \Psi\left(\sum_{i=1}^n \varphi(v_i)\right), \quad (v_1, \dots, v_n) \in [0, 1]^n,$$

where $\Psi : [0, +\infty) \mapsto [0, 1]$ is $(n-2)$ th differentiable, with $\Psi(0) = 1$ and $\lim_{t \rightarrow +\infty} \Psi(t) = 0$ is called an Archimedean survival copula with generator Ψ if $(-1)^s \Psi^s(t) \geq 0$, $s = 0, 1, \dots, n-2$ and $(-1)^{n-2} \Psi^{n-2}(t)$ is decreasing and convex for all $t \geq 0$. In this copula $\varphi = \Psi^{-1}$ is the inverse of Ψ and is right continuous $\varphi(v) = \Psi^{-1}(v) = \sup\{t \in \mathcal{R} : \Psi(t) > v\}$.

Definition 2.4. (Marshall et al. [2]). Consider a real-valued function ϕ defined on a set $\mathbb{A} \subseteq \mathbb{R}^n$. If $\mathbf{a} \succeq^m \mathbf{b}$ implies $\phi(\mathbf{a}) \geq (\leq) \phi(\mathbf{b})$ for any $\mathbf{a}, \mathbf{b} \in \mathbb{A}$, then ϕ is Schur-convex (Schur-concave) on $\mathbb{A}$.

Characterizations of Schur-convex (Schur-concave) functions are provided in the next lemma.

Lemma 2.5. (Marshall et al. [2]). Suppose that $\phi : I^n \rightarrow \mathbb{R}$ be a real-valued, continuously differentiable function, where $I \subseteq \mathbb{R}$ is an open interval. Then, ϕ is Schur-convex (Schur-concave) on I^n if and only if ϕ is symmetric on I^n , and for all $i \neq j$ and all $\mathbf{x} \in I^n$, $(a_i - a_j) \left(\frac{\partial \phi}{\partial a_i}(\mathbf{a}) - \frac{\partial \phi}{\partial a_j}(\mathbf{a}) \right) \geq (\leq) 0$, where $\frac{\partial \phi}{\partial a_i}$ is the partial derivative of ϕ with respect to its i -th argument.

Lemma 2.6. (Marshall et al. [2]). Let ϕ be a real-valued function defined on a set $\mathbb{A} \subseteq \mathbb{R}^n$. Then,

(i) ϕ is increasing and Schur-convex on $\mathbb{A}$ if and only if $\mathbf{a} \succeq_w \mathbf{b}$ implies $\phi(\mathbf{a}) \geq \phi(\mathbf{b})$;

(ii) ϕ is decreasing and Schur-convex on $\mathbb{A}$ if and only if $\mathbf{a} \stackrel{w}{\succeq} \mathbf{b}$ implies $\phi(\mathbf{a}) \geq \phi(\mathbf{b})$.

Lemma 2.7. (Li and Fang [1]). Let C_Φ and C_{Φ^*} are two n -dimensional Archimedean copulas. If $\varphi^* \circ \Phi$ is super-additive, then $C_\Phi(\mathbf{v}) \leq C_{\Phi^*}(\mathbf{v})$, for all $\mathbf{v} \in [0, 1]^n$ (a function g is super-additive, if $g(p) + g(q) \leq g(p + q)$, for all p and q belong to the support of g).

Before giving the main results of the paper, we need the following notation:

$$\mathcal{S}_n = \{(\mathbf{a}, \mathbf{b}) : a_i, b_i \geq 0 \quad (a_i - a_j)(b_i - b_j) \geq 0, \quad i, j = 1, \dots, n\}.$$

3 Ordering Results

This section examines the usual stochastic ordering of reliability functions for two series systems, each comprising m dependent subsystems. The components within each subsystem are assumed to be dependent and heterogeneous, with their dependence structure modeled using Archimedean copulas. Here, m represents the number of subsystems in a series system, while n_j (where $j = 1, \dots, m$) denotes the number of components in the j -th subsystem. The total number of components in the system is given by $n = \sum_{j=1}^m n_j$.

Theorem 3.1. Let $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)$ and $E_{1:n}^*(m, \mathbf{n}, \boldsymbol{\gamma}, \Psi^*)$ are the lifetimes random variables of two series systems with the respective parameters $\boldsymbol{\lambda} = (\lambda_{1,1}, \dots, \lambda_{1,n_1}, \dots, \lambda_{m,1}, \dots, \lambda_{m,n_m})$ and $\boldsymbol{\gamma} = (\gamma_{1,1}, \dots, \gamma_{1,n_1}, \dots, \gamma_{m,1}, \dots, \gamma_{m,n_m})$, the generators log-convex Φ and Φ^* , $j = 1, \dots, m$, and the respective generators Ψ and Ψ^* of Archimedean copulas with reliability functions $\bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x) = \Psi \left(\sum_{j=1}^m \psi \left[\Phi \left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x)) \right) \right] \right)$ and $\bar{F}_{E_{1:n}^*(m, \mathbf{n}, \boldsymbol{\gamma}, \Psi^*)}(x) = \Psi^* \left(\sum_{j=1}^m \psi^* \left[\Phi^* \left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x)) \right) \right] \right)$, respectively. Let $\varphi^* \circ \Phi$ and $\psi^* \circ \Psi$ are super-additive. Then $(\lambda_{j,1}, \dots, \lambda_{j,n_j}) \succeq_w (\gamma_{j,1}, \dots, \gamma_{j,n_j})$, $j = 1, \dots, m$ implies $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:n}^*(m, \mathbf{n}, \boldsymbol{\gamma}, \Psi^*)$.

Proof. We must to show that $\bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x) \leq \bar{F}_{E_{1:n}^*(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x)$. By Theorem 3.1 of Nadeb et al. [3], if $(\lambda_{j,1}, \dots, \lambda_{j,n_j}) \succeq_w (\gamma_{j,1}, \dots, \gamma_{j,n_j})$, $j = 1, \dots, m$, we have

$$\Phi\left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x))\right) \leq \Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right), \quad j = 1, \dots, m$$

and thus, from decreasingness property of ψ , we have

$$\psi\left(\Phi\left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x))\right)\right) \geq \psi\left(\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right). \quad (3.1)$$

Then,

$$\sum_{j=1}^m \psi\left[\Phi\left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x))\right)\right] \geq \sum_{j=1}^m \psi\left[\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right]. \quad (3.2)$$

Hence, from decreasingness property of Ψ , we have

$$\Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x))\right)\right]\right) \leq \Psi\left(\sum_{j=1}^m \psi\left[\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right]\right). \quad (3.3)$$

From assumption $\psi^* \circ \Psi$ is super-additive together with Lemma 2.7, one can obtain

$$\Psi\left(\sum_{j=1}^m \psi\left[\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right]\right) \leq \Psi^*\left(\sum_{j=1}^m \psi^*\left[\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right]\right). \quad (3.4)$$

By combining inequalities (3.3) and (3.4), we get

$$\Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(\sum_{i=1}^{n_j} \varphi(\exp(-\lambda_{j,i}x))\right)\right]\right) \leq \Psi^*\left(\sum_{j=1}^m \psi^*\left[\Phi^*\left(\sum_{i=1}^{n_j} \varphi^*(\exp(-\gamma_{j,i}x))\right)\right]\right).$$

That means $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:n}^*(m, \mathbf{n}, \boldsymbol{\gamma}, \Psi^*)$, and this completing the proof of the theorem. $\square$

Theorem 3.1 gives the following upper bound for the SF of the random variable $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)$.

Corollary 3.2. Set $(\gamma_1, \dots, \gamma_n) = (\bar{\lambda}_1, \dots, \bar{\lambda}_j)$, where $\bar{\lambda} = \frac{1}{n_j} \sum_{i=1}^{n_j} \lambda_{j,i}$. One can see that $\boldsymbol{\lambda}_j \succeq_w \boldsymbol{\gamma}_j$. Let $(\mathbf{n}, \boldsymbol{\lambda}) \in \mathcal{S}_m$ and Ψ and Φ be log-convex. For the SF of the random variable $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)$ an upper bound is as follows:

$$\bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x) \leq \Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j \varphi(\exp(-\bar{\lambda}_j x))\right)\right]\right).$$

Remark 3.3. The assumption of sub-additivity or super-additivity of $\psi^* \circ \Psi$ in Theorems 3.1 is not restrictive. Many Archimedean copulas satisfy this condition. We present some examples below. The super-additivity of $\psi^* \circ \Psi$ is easy to verify.

- (i) Let $\Psi(u) = e^{\frac{1-e^u}{\theta_1}}$ and $\Psi^*(u) = e^{\frac{1-e^u}{\theta_2}}$, where $0 < \theta_1, \theta_2 \leq 1$. Thus, $\psi^* \circ \Psi$ is super-additive for $\theta_1 \geq \theta_2$.
- (ii) Let $\Psi(u) = \frac{\theta}{\log(e^\theta + u)}$ and $\Psi^*(u) = \log(e^\theta + u)^{-\frac{1}{\theta}}$, where $\theta > 1$. Thus, $\psi^* \circ \Psi$ is super-additive.

Remark 3.4. The assumption of log-convexity of the generators of Archimedean copulas in Theorems 3.1 is satisfied in the following generators.

- (i) The Gumbel-Hougaard copula with generator $\Psi(u) = e^{-u^{\frac{1}{\theta}}}$.
- (ii) $\Psi(u) = \frac{\theta}{\log(e^\theta + u)}$, $\theta \in (0, +\infty)$.
- (iii) $\Psi(u) = \log(e^\theta + u)^{-\frac{1}{\theta}}$, $\theta \in (0, +\infty)$.
- (iv) The Clayton copula with generator $\Psi(u) = (1 + \theta u)^{-\frac{1}{\theta}}$, $\theta \in (0, +\infty)$.

Now, let the components in each subsystem have the same SFs, $\bar{F}_j(x) = \exp(-\lambda_j x)$, $j = 1, \dots, m$, but the number of components in each subsystem of the first system is different from the second one. Also, suppose that the subsystems have the same dependency structure.

Theorem 3.5. Suppose that $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)$ and $E_{1:n^*}^*(m, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)$ are the lifetimes random variables of two series systems with the common parameters $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_m)$, the common log-concave generators Φ , and the respective log-convex generators Ψ and Ψ^* of Archimedean copulas with SF's

$$\bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x) = \Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j \varphi(\exp(-\lambda_j x))\right)\right]\right)$$

and

$$\bar{F}_{E_{1:n^*}^*(m, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)}(x) = \Psi^*\left(\sum_{j=1}^m \psi^*\left[\Phi\left(n_j^* \varphi(\exp(-\lambda_j x))\right)\right]\right),$$

where $\sum_{j=1}^m n_j = n$ and $\sum_{j=1}^m n_j^* = n^*$. For $(\mathbf{n}, \boldsymbol{\lambda}) \in \mathcal{S}_m$ and $(\mathbf{n}^*, \boldsymbol{\lambda}) \in \mathcal{S}_m$, if $\mathbf{n} \succeq_w \mathbf{n}^*$ and $\psi^* \circ \Psi$ be super-additive, then $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:n^*}^*(m, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)$.

Proof. We must to show that $\bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x)$ is decreasing and Schur-concave with respect to $\mathbf{n}$. By considering $\zeta(\mathbf{n}) = \bar{F}_{E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x)$, the partial derivative of $\xi(\mathbf{n})$ with respect to n_j is

$$\begin{aligned} \frac{\partial \zeta(\mathbf{n})}{\partial n_j} &= \Psi'\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j \varphi(\exp(-\lambda_j x))\right)\right]\right) \times \psi'\left[\Phi\left(n_j \varphi(\exp(-\lambda_j x))\right)\right] \\ &\quad \times \Phi'\left(n_j \varphi(\exp(-\lambda_j x))\right) \times \varphi(\exp(-\lambda_j x)). \end{aligned}$$

The decreasingness property of generators of Archimedean copulas yields: $\Psi' \leq 0$, $\psi' = (\Psi^{-1})' \leq 0$ and $\Phi' \leq 0$. Thus, $\frac{\partial \zeta(\mathbf{n})}{\partial n_j} \leq 0$, and $\zeta(\mathbf{n})$ is decreasing in n . In addition, for any pair (u, s) , similar to the proof of Theorem 3.1, we have

$$\begin{aligned} & \left[\frac{\partial \zeta(\mathbf{n})}{\partial n_u} - \frac{\partial \zeta(\mathbf{n})}{\partial n_s} \right] \stackrel{\text{sign}}{=} - \left[\psi' \left[\Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right) \right] \times \Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right) \times \right. \\ & \frac{\Phi' \left(n_u \varphi(\exp(-\lambda_u x)) \right)}{\Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right)} \times \varphi(\exp(-\lambda_u x)) - \psi' \left[\Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right) \right] \times \varphi(\exp(-\lambda_s x)) \\ & \left. - \psi' \left[\Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right) \right] \times \Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right) \times \frac{\Phi' \left(n_s \varphi(\exp(-\lambda_s x)) \right)}{\Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right)} \right. \\ & \left. \times \varphi(\exp(-\lambda_s x)) \right]. \end{aligned}$$

Now, for $u \leq s$ ($u > s$) and from the assumption that $(\mathbf{n}, \boldsymbol{\lambda}) \in \mathcal{S}_m$, $n_u \leq (\geq) n_s$ and $\lambda_u \leq (\geq) \lambda_s$. Thus, $\exp(-\lambda_u x) \geq (\leq) \exp(-\lambda_s x)$. On the one hand, from the decreasing property of φ , we have

$$\varphi(\exp(-\lambda_u x)) \leq (\geq) \varphi(\exp(-\lambda_s x)),$$

and then

$$n_u \varphi(\exp(-\lambda_u x)) \leq (\geq) n_s \varphi(\exp(-\lambda_s x)).$$

On the other hand, from log-concavity and decreasingness of Φ , respectively, we get

$$\begin{aligned} 0 & \geq \frac{\Phi' \left(n_u \varphi(\exp(-\lambda_u x)) \right)}{\Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right)} \geq (\leq) \frac{\Phi' \left(n_s \varphi(\exp(-\lambda_s x)) \right)}{\Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right)} (\leq 0), \\ & \Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right) \geq (\leq) \Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right). \end{aligned}$$

Since ψ is log-convex, one can obtain

$$\begin{aligned} 0 & \geq \Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right) \psi' \left[\Phi \left(n_u \varphi(\exp(-\lambda_u x)) \right) \right] \\ & \geq (\leq) \Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right) \psi' \left[\Phi \left(n_s \varphi(\exp(-\lambda_s x)) \right) \right] (\leq 0). \end{aligned}$$

Consequently,

$$(n_u - n_s) \left[\frac{\partial \zeta(\mathbf{n})}{\partial n_u} - \frac{\partial \zeta(\mathbf{n})}{\partial n_s} \right] \leq 0.$$

Thus, using Lemma 2.5, $\zeta(\mathbf{n})$ is Schur-concave with respect to $\mathbf{n}$. In addition, one can see that $\zeta(\mathbf{n})$ is symmetric with respect to n . Therefore, from Lemma 2.6, if $\mathbf{n} \succeq_w \mathbf{n}^*$, then

$$\Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j \varphi(\exp(-\lambda_j x))\right)\right]\right) \leq \Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j^* \varphi(\exp(-\lambda_j x))\right)\right]\right). \quad (3.5)$$

Similar to the proof of Theorem 3.1, from assumption $\psi^* \circ \Psi$ is super-additive, we get

$$\Psi\left(\sum_{j=1}^m \psi\left[\Phi\left(n_j \varphi(\exp(-\lambda_j x))\right)\right]\right) \leq \Psi^*\left(\sum_{j=1}^m \psi^*\left[\Phi\left(n_j^* \varphi(\exp(-\lambda_j x))\right)\right]\right).$$

That means $E_{1:n}(m, \mathbf{n}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:n^*}^*(m, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)$. $\square$

3.1 Examples

Consider the following example to illustrate the theoretical findings.

Example 3.6. Let $\Psi(v) = \frac{2}{\log(e^2 + v)}$ and $\Psi^*(v) = \log(e^{1.8} + v)^{-\frac{1}{1.8}}$. Clearly, Ψ and Ψ^* are log-convex, and $\psi^* \circ \Psi$ is super-additive. Set $n = 12$, $m = 3$, $\mathbf{n} = (4, 4, 4)$, and $\Phi(u) = \frac{1.5}{\log(e^{1.5} + u)}$. Put the parameters of the systems

$$\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_{12}) = (0.3, 0.36, 0.42, 0.48, 0.54, 0.55, 0.57, 0.66, 0.74, 0.79, 0.85, 0.9),$$

and

$$\boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_{12}) = (0.2, 0.24, 0.28, 0.32, 0.36, 0.4, 0.45, 0.48, 0.55, 0.58, 0.66, 0.77).$$

It is easy to see that $\boldsymbol{\lambda}_j \succeq_w \boldsymbol{\gamma}_j$. Thus, all conditions of Theorem 3.1 are met. Figure 1 depicts the plots of $\bar{F}_{E_{1:12}(3, \mathbf{4}, \boldsymbol{\lambda}, \Psi)}(x)$ and $\bar{F}_{E_{1:12}^*(3, \mathbf{4}, \boldsymbol{\gamma}, \Psi^*)}(x)$. As can be seen from the plot $E_{1:12}(3, \mathbf{4}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:12}^*(3, \mathbf{4}, \boldsymbol{\gamma}, \Psi^*)$.

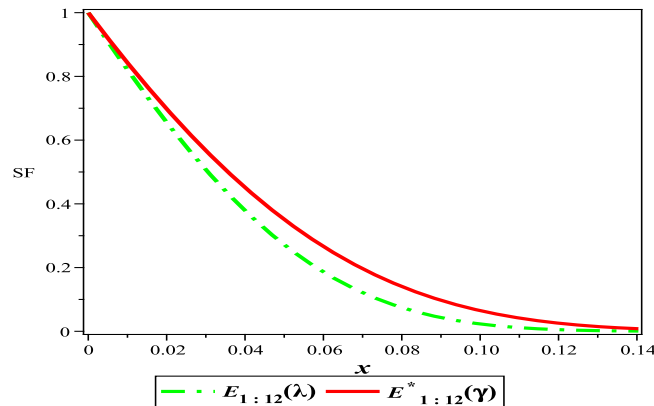


Figure 1: $\bar{F}_{E_{1:12}(3, \mathbf{4}, \boldsymbol{\lambda}, \Psi)}(x)$ (dash dot) and $\bar{F}_{E_{1:12}^*(3, \mathbf{4}, \boldsymbol{\gamma}, \Psi^*)}(x)$ (solid) in Example 3.6.

Example 3.7. Let $\Psi(v) = e^{\frac{1-e^v}{0.75}}$ and $\Psi^*(v) = e^{\frac{1-e^v}{0.55}}$. Clearly, $\psi^* \circ \Psi$ is super-additive. Set $n = 9$, $n^* = 7$, $m = 3$, $\mathbf{n} = (n_1, n_2, n_3) = (4, 3, 2)$, $\mathbf{n}^* = (n_1^*, n_2^*, n_3^*) = (3, 2, 2)$, $\Phi(u) = e^{\frac{1-e^u}{0.8}}$, and $\boldsymbol{\lambda} = (\lambda_1, \lambda_2, \lambda_3) = (0.6, 0.5, 0.4)$. It is easy to see that $\mathbf{n} \succeq_w \mathbf{n}^*$, $(\mathbf{n}, \boldsymbol{\lambda}) \in \mathcal{T}_3$ and $(\mathbf{n}^*, \boldsymbol{\lambda}) \in \mathcal{T}_3$. Thus, all conditions of Theorem 3.5 are met. Figure 2 depicts the plots of $\bar{F}_{E_{1:9}(3, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x)$ and $\bar{F}_{E_{1:7}^*(3, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)}(x)$, which indicates that $E_{1:9}(3, \mathbf{n}, \boldsymbol{\lambda}, \Psi) \leq_{st} E_{1:7}^*(3, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)$. This coincides with the result of Theorem 3.5.

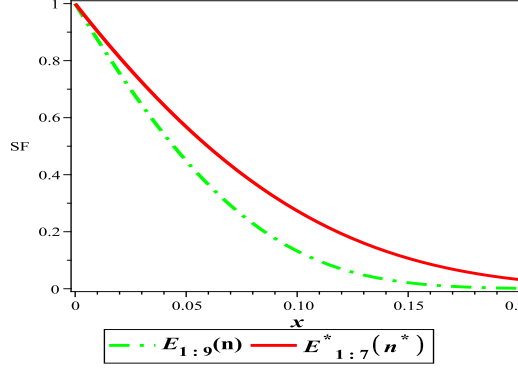


Figure 2: $\bar{F}_{E_{1:9}(3, \mathbf{n}, \boldsymbol{\lambda}, \Psi)}(x)$ (dash dot) and $\bar{F}_{E_{1:7}^*(3, \mathbf{n}^*, \boldsymbol{\lambda}, \Psi^*)}(x)$ (solid) in Example 3.7.

4 Conclusions

This paper has analyzed series systems with dependent and heterogeneous Exponential components, where dependence is modeled via Archimedean copulas. Our results demonstrate that when the first system's vector of subsystem component counts and scale parameters majorize those of the second system, the first system is less reliable than the second under the usual stochastic order. These findings provide a theoretical foundation for comparing system reliability in scenarios involving complex dependencies and heterogeneity, which are common in real-world applications such as manufacturing, telecommunications, and power grids.

Future research could extend this work in several directions such as investigating non-Exponential (e.g., Weibull, Gamma or general case) component lifetimes while preserving dependence structures and expanding the analysis to parallel, series-parallel, or k -out-of- n system configurations.

References

- [1] Li, X. and Fang, R. (2015). Ordering properties of order statistics from random variables of Archimedean copulas with applications. *Journal of Multivariate Analysis*, **133**, 304-320.
- [2] Marshall, A. W., Olkin, I. and Arnold, B. C. (2011). *Inequalities: Theory of Majorization and its Applications*. Springer, New York.

- [3] Nadeb, H., Torabi, H., and Dolati, A. (2021). Some general results on usual stochastic ordering of the extreme order statistics from dependent random variables under Archimedean copula dependence. *Journal of the Korean Statistical Society*, 1-17.
- [4] Nelsen, R. B. (2007). *An introduction to copulas*. New York: Springer Science & Business Media.
- [5] Shojaee, O., Asadi, M. and Finkelstein, M. (2022). Stochastic properties of generalized finite α -mixtures. *Probability in the Engineering and Informational Sciences*, **36(4)**, 1055-1079.
- [6] Shojaee, O., Mohammadi, S. M. and Momeni, R. (2024). Ordering results for the smallest (largest) and the second smallest (second largest) order statistics of dependent and heterogeneous random variables, *Metrika*, **87(3)** 325-347.